

Springer-Lehrbuch

Thomas Westermann

Mathematik für Ingenieure

Ein anwendungsorientiertes Lehrbuch

5., neu bearbeitete Auflage

 Springer

Professor Dr. Thomas Westermann
Hochschule Karlsruhe
Technik und Wirtschaft
Postfach 2440
76012 Karlsruhe

e-mail: Thomas.Westermann@hs-karlsruhe.de
<http://www.home.hs-karlsruhe.de/~weth0002>

Homepage zum Buch
<http://www.home.hs-karlsruhe.de/~weth0002/buecher/mathe/start.htm>
<http://www.springer.com/978-3-540-77730-4>

Vorauslagen erschienen 2bändig unter dem Titel: Mathematik für Ingenieure mit Maple

ISBN 978-3-540-77730-4

e-ISBN 978-3-540-77731-1

DOI 10.1007/978-3-540-77731-1

Springer Lehrbuch ISSN 0937-7433

Bibliografische Information der Deutschen Nationalbibliothek

Die Deutsche Bibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie;
detaillierte bibliografische Daten sind im Internet über <http://dnb.d-nb.de> abrufbar.

© 2008, 2006, 2003, 1996 Springer-Verlag Berlin Heidelberg

Dieses Werk ist urheberrechtlich geschützt. Die dadurch begründeten Rechte, insbesondere die der Übersetzung, des Nachdrucks, des Vortrags, der Entnahme von Abbildungen und Tabellen, der Funksendung, der Mikroverfilmung oder der Vervielfältigung auf anderen Wegen und der Speicherung in Datenverarbeitungsanlagen, bleiben, auch bei nur auszugsweiser Verwertung, vorbehalten. Eine Vervielfältigung dieses Werkes oder von Teilen dieses Werkes ist auch im Einzelfall nur in den Grenzen der gesetzlichen Bestimmungen des Urheberrechtsgesetzes der Bundesrepublik Deutschland vom 9. September 1965 in der jeweils geltenden Fassung zulässig. Sie ist grundsätzlich vergütungspflichtig. Zuwiderhandlungen unterliegen den Strafbestimmungen des Urheberrechtsgesetzes.

Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen usw. in diesem Werk berechtigt auch ohne besondere Kennzeichnung nicht zu der Annahme, dass solche Namen im Sinne der Warenzeichen- und Markenschutz-Gesetzgebung als frei zu betrachten wären und daher von jedermann benutzt werden dürften.

Einbandgestaltung: WMXDesign, Heidelberg

Herstellung: le-tex publishing services oHG, Leipzig

Satz: Reproduktionsfähige Vorlage des Autors

Gedruckt auf säurefreiem Papier

9 8 7 6 5 4 3 2 1

springer.com

Vorwort zur 5. Auflage

In der vorliegenden 5. Auflage wurde die Umstellung von Diplom- zu **Bachelor**-Studiengängen berücksichtigt. Diese Umstellung bedeutet in der Regel, dass weniger Stunden für die Mathematik zur Verfügung stehen. Daher wurde das ursprünglich zweibändige Lehrbuch nun in einem Band zusammengefasst, der den vollständigen Stoffumfang der Mathematikausbildung für Ingenieure an Berufsakademien und technischen Hochschulen beinhaltet. Die grundlegenden Kapitel wurden erweitert und an den Bachelor-Standard mit einem größeren Übungsanteil angepasst.

In Erweiterung der *Bachelor*-Studieninhalten wurden zusätzlich für die **Master**-Ausbildung die folgenden Themen

- Numerisches Lösen von Gleichungen
- Numerisches Differenzieren und Integrieren
- Numerisches Lösen von gewöhnlichen Differenzialgleichungen
- Diskrete Fourier-Transformation und deren Anwendungen
- Signal- und Systemtheorie für lineare Systeme
- Partielle Differenzialgleichungen
- Vektoralgebra

sowie viele ergänzende und weiterführende Abschnitte mit aufgenommen aber auf die CD-Rom ausgelagert. Die pdf-Version des Buches auf der CD-Rom beinhaltet auch diese weiterführenden Themen, die im Inhaltsverzeichnis mit der Seitenangabe „cd“ gekennzeichnet sind.

Mit der Umstellung auf einen Band wurde ein komplett neues Layout eingeführt, das eine übersichtlichere Darstellung der Lehrinhalte ermöglicht aber auch die Möglichkeit bietet, die pdf-Version auf der CD-Rom als elektronisches, interaktives E-Book zu benutzen. Das neue Layout verbessert die bisherige Darstellung, indem

- mit dem Symbol „ **Achtung:**“ auf Stellen besonders hingewiesen, die man anfänglich oftmals falsch bearbeitet, übersieht oder nicht beachtet,
- durch Markierungen am Seitenrand Gliederungspunkte und Orientierungshilfen gegeben werden,
- Definitionen und wichtige Sätze in blau unterlegten Boxen markiert werden,
- die zahlreichen Zusammenfassungen farblich hervorgehoben werden,
- wichtige Formel und Ergebnisse gekennzeichnet werden,
- Musterbeispiele und Anwendungsbeispiele übersichtlich aus dem Text hervorgehen;
- 380 ausführlich durchgerechnete Beispiele,
- über 360 Aufgaben mit Lösungen
- und mehr als 200 Abbildungen und Skizzen zum **Selbststudium** und zur **Prüfungsvorbereitung** dienen.

Der Grundidee folgend, mathematische Begriffe zu visualisieren, um sie greifbarer zu machen und den interaktiven Gebrauch des Buches zu fördern, wurde die CD-Rom völlig neu und benutzerfreundlicher gestaltet. Die Interaktivität der CD-Rom-Version wird unterstützt, indem

① Inhaltsverzeichnis, Index und Beispiele verlinkt sind, so dass eine bequeme Navigation innerhalb der pdf-Version möglich ist;



② Animationen, die im gif-Format vorliegen, durch Anklicken des nebenstehenden Symbols abgespielt werden können;



③ das CD-Rom-Symbol auf [MAPLE-Beschreibungen](#) hinweist. Durch Anklicken des Symbols oder der rot unterlegten Textstelle startet man das zugehörige MAPLE-Worksheet.

Sämtlichen MAPLE-Beschreibungen sowie MAPLE-Worksheets sind auf der CD-Rom enthalten. Die pdf-Version des Buches ist verlinkt, so dass die elektronischen Arbeitsblätter direkt aus pdf gestartet werden können. Voraussetzung ist, dass MAPLE auf dem Rechner installiert und die Erweiterung *.mws* mit einer MAPLE-Version verknüpft ist. Die zugehörigen Links sind Rot gekennzeichnet; bei Beispielen steht zusätzlich „(Mit MAPLE-Worksheet)“. Die MAPLE-Prozeduren sind zusätzlich über den Prozedurnamen verlinkt und am Ende jeden Kapitels ist eine verlinkte Liste der zum Kapitel gehörenden Worksheets.

Für die vorliegende 5. Auflage wurden die MAPLE-Beschreibungen an MAPLE 11 angepasst. Um auch zukünftig mit neuen MAPLE-Versionen Schritt halten zu können, werden Updates der MAPLE-Worksheets unter

<http://www.home.hs-karlsruhe.de/~weth0002/buecher/mathe/start.htm> abrufbar sein.

Auf der CD-Rom sind weitere Informationen und Ergänzungen zugänglich:

- alle Worksheets, die im Text beschrieben sind, inclusive vieler zusätzlicher MAPLE-Prozeduren zur Visualisierung mathematischer Begriffe;
- die Lösungen der Aufgaben;
- zusätzliche Kapitel und Ergänzungen, die in der Buchform des Gesamtumfangs wegen nicht mehr eingebunden werden konnten;
- eine Einführung in die Benutzeroberfläche von MAPLE 11.

Mein Dank gilt Herrn Richard von Scientific Computers und Waterloo Maple Inc., die mir MAPLE 11 zur Verfügung gestellt haben sowie Frau Hestermann-Beyerle vom Springer-Verlag für die gute und angenehme Zusammenarbeit. Dieses Buch ist meinen Töchtern Veronika (zum Nachschlagen von längst Vergessenem) und Juliane (zur Vorbereitung der Matheklausuren ihres Studiums) gewidmet.

Vorwort zur 1. Auflage

Dieses zweibändige Lehrbuch entstand aus Vorlesungen und Übungen zur Mathematik und Physikalischen Simulation für Ingenieure an der Hochschule Karlsruhe. Es wendet sich aber an alle Studenten der Natur- und Ingenieurwissenschaften, da auch Themengebiete einbezogen sind, die nicht bzw. nicht in der vorliegenden Tiefe in der Vorlesung behandelt wurden.

...

Die stürmische Entwicklung von Computersoftware im Bereich der Mathematik erfordert eine Erweiterung der Ingenieur-Ausbildung, indem nicht nur praxisorientiertes mathematisches Wissen, sondern auch das Rüstzeug vermittelt wird, mit diesen Systemen erfolgreich arbeiten zu können. Die Computeralgebra-Systeme werden zum numerischen Rechnen genauso verwendet wie zum Manipulieren von Formeln sowie der graphischen Darstellung komplizierter Sachverhalte. Die Rechentechnik tritt in den Hintergrund; die interessante Modellierung und das systematische Vorgehen gewinnt an Bedeutung. In diesem Lehrbuch wird dieser neue spannende Aspekt aufgegriffen und das Computeralgebra-System MAPLE in die Mathematikausbildung mit einbezogen.

Mathematische Begriffe werden anschaulich motiviert, systematisch anhand praxisbezogener Beispiele verdeutlicht und mit MAPLE-Worksheets umgesetzt, was sich in vielen Animationen niederschlägt. Auf mathematische Beweise wird fast gänzlich verzichtet und einer anschaulich prägnanten Sprechweise den Vorzug gegenüber einer mathematisch exakten Formulierung gegeben.

Um den ständig wachsenden Gebrauch von Rechnern und numerischen Problemlösungen zu berücksichtigen, wurden zahlreiche Abschnitte zur rechnerischen Lösung von Standard-Problemen in dieses Mathematikbuch aufgenommen. Die numerischen Algorithmen sind als MAPLE-Prozeduren auf der beigelegten CD-Rom enthalten, können aber von etwas geübten Programmierern leicht in jede andere höhere Sprache umgesetzt werden.

...

Besonders bedanken möchte ich mich bei Herrn F. Wohlfarth und Frau Raviol für die präzise und fehlerfreie Erstellung des $\text{\LaTeX}$ -Quelltextes mit all den vielen Formeln, den Herren M. Baus und F. Loeffler für die exzellente Erstellung der meisten Skizzen und Bilder unter CorelDraw, so wie der Autor sie sich vorgestellt hat. Mein Dank gilt auch dem Springer-Verlag für die angenehme und reibungslose Zusammenarbeit, speziell Herrn Dr. Merkle.

Zuletzt möchte ich mich bei meiner Familie (Ulrike, Veronika, Juliane) bedanken, die mit viel Verständnis meine Arbeit an diesem Buch mitgetragen und tatkräftig unterstützt hat.

Karlsruhe, im Juni 1996

Thomas Westermann

Hinweise zum Gebrauch dieses Buches

Die einzelnen Kapitel fassen mehrere Aspekte einer Thematik zusammen. Nicht immer ließ es sich vermeiden, Teilergebnisse aus späteren Kapiteln vorwegzunehmen und zu verwenden. Dem didaktischen Anliegen, Themenbereiche geschlossen in einem Block zu bearbeiten, wurde dabei stärkere Priorität als der mathematischen Strenge beigemessen. Die Reihenfolge innerhalb eines Vorlesungszyklus muss sich nicht an die im Buch gewählte Reihenfolge halten, einzelne Kapitel können auch aufgesplittet werden.

Dieses Buch ist ein Lehrbuch über Mathematik und kann **ohne** Rechner zum Erlernen von mathematischem Grundwissen oder zur Prüfungsvorbereitung herangezogen werden. Um den vollen Umfang und die ganze Schönheit der Mathematik und der Anwendungen zu erleben, sind die Animationen und Ausarbeitungen mit dem Computeralgebra-System MAPLE unverzichtbar. Nur wenn eine Animation als Animation erlebt wird, kommt die volle Erkenntnis zum Tragen.

Darstellung: Neu eingeführte Begriffe werden *kursiv* im Text markiert und zumeist in einer Definition **fett** spezifiziert. Lehrsätze, wichtige Formeln und Zusammenfassungen sind durch Umrahmungen besonders gekennzeichnet. Am Ende eines jeden Kapitels befinden sich Aufgaben, deren Lösungen auf der CD-Rom angegeben sind. Bei der Erarbeitung der Themengebiete wird eine anwendungsorientierte Problemstellung vorangestellt und anschließend auf die allgemeine mathematische Struktur übergegangen. Die Thematik wird dann innerhalb der Mathematik bearbeitet und anhand von mathematischen Beispielen erläutert. Neben der Behandlung der Problemstellungen mit MAPLE werden aussagekräftige Anwendungsbeispiele diskutiert.

Beispiele: Die zahlreichen Beispiele sind für den Zugang zu den Themengebieten unverzichtbar. Beim Selbststudium und zur Prüfungsvorbereitung sollten möglichst die mathematischen Beispiele eigenständig bearbeitet werden. Wer dieses Werk als Nachschlagewerk benutzt, kann sich an den durchgerechneten Beispielen sowie an den eingerahmten Definitionen, Sätzen und Zusammenfassungen orientieren.

Aufgaben: Alle Übungsaufgaben sind soweit nicht speziell gekennzeichnet mit den Hilfsmitteln der einzelnen Paragraphen zu bearbeiten. Die Lösungen zu den Aufgaben befinden sich als *pdf*-File auf der CD-Rom.

MAPLE: Die CD-Rom-Version dieses Buchs kann als eine themengebundene Einführung in die Anwendung von MAPLE in der Mathematik gesehen werden, da sämtliche Themengebiete des Buches mit MAPLE bearbeitet werden. Alle MAPLE-Befehle sind im Text **fett** hervorgehoben; die MAPLE-Syntax erkennt man an der Eingabeaufforderung `">"` zu Beginn einer Zeile. Diese MAPLE-

Zeilen sind im Textstil **sans serif** angegeben und können direkt in MAPLE eingegeben werden. Die MAPLE-Ausgabe erscheint im Formelmodus. Am Ende jedes Kapitels steht eine Zusammenfassung der verwendeten Befehle.

CD-Rom: Alle MAPLE-Ausarbeitungen sind auf der CD-Rom als elektronische Arbeitsblätter (Worksheets) enthalten, so dass der interessierte Leser die im Text entwickelten Methoden umsetzen bzw. an abgeänderten Beispielen erproben kann. Es wird besonders auf die vielen Animationen und Prozeduren hingewiesen, welche die elementaren Begriffe visualisieren und die mathematischen Zusammenhänge aufzeigen: Durch eine benutzerfreundliche Menüführung soll die interaktive Benutzung der Worksheets sowohl zum Lösen von mathematischen Problemen als auch zum experimentieren mit mathematischen Begriffen gefördert werden.

Aufbau der CD-Rom: Die Struktur der Dateien und Verzeichnisse ist:

<code>buch.pdf</code>	enthält den Inhalt des vorliegenden, gedruckten Buches.
<code>buch_CD.pdf</code>	enthält den gesamten Inhalt des erweiterten Buches mit allen zusätzlichen Kapiteln und Abschnitten. Zum Navigieren innerhalb des Textes verwendbar sowie zum direkten Starten der zugehörigen MAPLE-Worksheets.
<code>index.mws</code>	Inhaltsverzeichnis der MAPLE-Worksheets.
<code>\wrksheet\</code>	enthält alle Worksheets im <i>mws</i> -Format.
<code>\animationen\</code>	enthält alle Animationen im <i>gif</i> -Format.
<code>readme.wri</code>	letzte Änderungen, die nicht mehr im Text aufgenommen werden konnten.

Arbeiten mit der mws-Datei: Durch Doppelklicken der Datei *index.mws* öffnet man das MAPLE-Inhaltsverzeichnis. Durch anschließendes Anklicken des gewünschten Abschnitts wird das zugehörige MAPLE-Worksheet gestartet und ist dann interaktiv bedienbar. Mit der Zurück-Taste der oberen Taskleiste kommt man vom Worksheet wieder zum Inhaltsverzeichnis zurück.

Systemvoraussetzungen: MAPLE 11 ist auf dem Rechner installiert (empfohlen), mindestens jedoch MAPLE 6. *mws* ist je nach Version mit dem ausführbaren Programm *cwmaple.exe* bzw. *maplew.exe* im Maple-bin-Verzeichnis verknüpft. Acrobat-Reader steht zur Verfügung; er kann ansonsten von der CD-Rom installiert werden.

Die Worksheets auf der CD-Rom sind unter der Extension *.mws* abgespeichert und unter beiden Benutzeroberflächen (Classic Worksheet und Standard Worksheet) uneingeschränkt lauffähig. Bis auf kleine Einschränkungen sind die Worksheets unter allen MAPLE-Versionen ab MAPLE 6 lauffähig. Alleine die auf dem lokalen Rechner spezifizierte Verknüpfung entscheidet, welche MAPLE-Variante gestartet wird.

Inhaltsverzeichnis

1	Zahlen, Gleichungen und Gleichungssysteme	1
1.1	Mengen	3
1.2	Natürliche Zahlen	5
1.3	Reelle Zahlen	13
1.4	Gleichungen und Ungleichungen	19
1.5	Lineare Gleichungssysteme	26
1.6	Aufgaben zu Zahlen, Gleichungen und Gleichungssystemen	36
2	Vektoren und Vektorrechnung	39
2.1	Vektoren im $\mathbb{R}^2$	42
2.2	Vektoren im $\mathbb{R}^3$	50
2.3	Geraden und Ebenen im $\mathbb{R}^3$	61
2.4	Vektorräume	76
2.5	Aufgaben zur Vektorrechnung	92
3	Matrizen und Determinanten	97
3.1	Matrizen	99
3.2	Determinanten	114
3.3	Lösbarkeit von linearen Gleichungssystemen	124
3.4	Aufgaben zu Matrizen und Determinanten	135
4	Elementare Funktionen	137
4.1	Allgemeine Funktionseigenschaften	139
4.2	Polynome	152
4.2.1	Festlegung von Polynomen durch gegebene Wertepaare	153
4.2.2	Koeffizientenvergleich	154
4.2.3	Teilbarkeit durch einen Linearfaktor	155
4.2.4	Nullstellenproblem	156
4.2.5	Interpolationspolynome mit dem Newton-Algorithmus	159
4.3	Rationale Funktionen	162
4.4	Potenz- und Wurzelfunktionen	167
4.5	Exponential- und Logarithmusfunktion	170
4.6	Trigonometrische Funktionen	175
4.7	Arkusfunktionen	182
4.8	Aufgaben zu elementaren Funktionen	188
5	Komplexe Zahlen	191
5.1	Darstellung komplexer Zahlen	194
5.2	Komplexe Rechenoperationen	200
5.2.1	Addition	200
5.2.2	Subtraktion	200
5.2.3	Multiplikation	201
5.2.4	Division	203
5.2.5	Potenz	205

5.2.6	Wurzeln	206
5.2.7	Fundamentalsatz der Algebra.....	207
5.3	Anwendungen.....	209
5.4	Aufgaben zu komplexen Zahlen	219
6	Grenzwert und Stetigkeit	221
6.1	Reelle Zahlenfolgen	223
6.2	Funktionsgrenzwert	229
6.3	Stetigkeit einer Funktion	235
6.4	Intervallhalbierungs-Methode	237
6.5	Aufgaben zu Grenzwert und Stetigkeit	240
7	Differenzialrechnung	241
7.1	Einführung	243
7.2	Rechenregeln bei der Differenziation	249
7.2.1	Faktorregel	249
7.2.2	Summenregel	249
7.2.3	Produktregel	250
7.2.4	Quotientenregel	251
7.2.5	Kettenregel.....	252
7.2.6	Begründung der Formeln 7.2.1 - 7.2.5.....	254
7.2.7	Ableitung der Umkehrfunktion	255
7.2.8	Logarithmische Differenziation	258
7.2.9	Implizite Differenziation	260
7.3	Anwendungsbeispiele aus Physik und Technik	262
7.4	Differenzial einer Funktion	265
7.5	Anwendungen in der Mathematik	270
7.6	Extremwertaufgaben (Optimierungsprobleme)	277
7.7	Sätze der Differenzialrechnung	281
7.8	Spektrum eines strahlenden schwarzen Körpers	287
7.9	Newton-Verfahren	289
7.10	Aufgaben zur Differenzialrechnung	293
8	Integralrechnung	295
8.1	Das Riemann-Integral	297
8.2	Fundamentalsatz der Differenzial- und Integralrechnung... ..	302
8.3	Grundlegende Regeln der Integralrechnung	311
8.4	Integrationsmethoden	313
8.4.1	Partielle Integration	313
8.4.2	Integration durch Substitution	315
8.4.3	Partialbruchzerlegung	321
8.5	Uneigentliche Integrale.....	327
8.6	Anwendungen der Integralrechnung	329
8.7	Aufgaben zur Integralrechnung.....	339

9	Funktionenreihen	341
9.1	Zahlenreihen	344
9.2	Potenzreihen	355
9.3	Taylor-Reihen	361
9.4	Anwendungen	371
9.5	Komplexwertige Funktionen	376
9.6	Aufgaben zu Funktionenreihen	385
10	Differenzialrechnung bei Funktionen mit mehreren Variablen	387
10.1	Funktionen mit mehreren Variablen	389
10.2	Stetigkeit	398
10.3	Differenzialrechnung	400
10.3.1	Partielle Ableitung	400
10.3.2	Totale Differenzierbarkeit	408
10.3.3	Gradient und Richtungsableitung	410
10.3.4	Der Taylorsche Satz	416
10.4	Anwendungen der Differenzialrechnung	423
10.4.1	Das Differenzial als lineare Näherung	423
10.4.2	Fehlerrechnung	428
10.4.3	Lokale Extrema bei Funktionen mit mehreren Variablen	432
10.4.4	Ausgleichen von Messfehlern; Regressionsgerade	439
10.5	Aufgaben zur Differenzialrechnung	446
11	Integralrechnung bei Funktionen mit mehreren Variablen	449
11.1	Doppelintegrale (Gebietsintegrale)	451
11.2	Dreifachintegrale	464
11.3	Aufgaben zur Integralrechnung	471
12	Gewöhnliche Differenzialgleichungen	473
12.1	Differenzialgleichungen erster Ordnung	476
12.1.1	Einleitende Problemstellungen	476
12.1.2	Lösen der homogenen Differenzialgleichung	480
12.1.3	Lösen der inhomogenen Differenzialgleichung	482
12.1.4	Lineare Differenzialgleichungen mit konstantem Koeffizient	490
12.1.5	Nichtlineare Differenzialgleichungen 1. Ordnung	494
12.1.6	Numerisches Lösen von DG 1. Ordnung	497
12.2	Lineare Differenzialgleichungssysteme	501
12.2.1	Einführung	501
12.2.2	Homogene lineare Differenzialgleichungssysteme	504
12.2.3	Eigenwerte und Eigenvektoren	508
12.2.4	Lösen von homogenen LDGS mit konstanten Koeffizienten	513

12.3	Lineare Differenzialgleichungen n -ter Ordnung	523
12.3.1	Einleitende Beispiele.....	523
12.3.2	Reduktion einer linearen DG n -ter Ordnung auf ein System	526
12.3.3	Homogene DG n -ter Ordnung mit konst. Koeffizienten....	531
12.3.4	Inhomogene DG n -ter Ord. mit konstanten Koeffizienten .	541
12.4	Aufgaben zu Differenzialgleichungen	555
13	Laplace-Transformation	559
13.1	Die Laplace-Transformation.....	563
13.2	Inverse Laplace-Transformation	568
13.3	Zwei grundlegende Eigenschaften.....	569
13.4	Methoden der Rücktransformation	574
13.5	Anwendungen der Laplace-Transformation	577
13.6	Aufgaben zur Laplace-Transformation	581
14	Fourier-Reihen	583
14.1	Einführung	585
14.2	Bestimmung der Fourier-Koeffizienten.....	587
14.3	Fourier-Reihen für 2π -periodische Funktionen	590
14.4	Fourier-Reihen für p -periodische Funktionen	597
14.5	Fourier-Reihen für komplexwertige Funktionen	605
14.6	Aufgaben zu Fourier-Reihen	610
15	Fourier-Transformation	611
15.1	Fourier-Transformation und Beispiele.....	613
15.2	Eigenschaften der Fourier-Transformation	623
15.3	Fourier-Transformation der Deltafunktion.....	637
15.4	Aufgaben zur Fourier-Transformation.....	645
	Literaturverzeichnis	647
	Sachverzeichnis	649

Zusätzliche Kapitel auf der CD-Rom

12*	Numerisches Lösen von Differenzialgleichungen	
12.5	Numerisches Lösen von Differenzialgleichungen 1. Ordnung	cd
12.5.1	Streckenzugverfahren von Euler	cd
12.5.2	Verfahren höherer Ordnung	cd
12.5.3	Vergleich der numerischen Verfahren mit MAPLE	cd
12.5.4	Numerisches Lösen von DG 1. Ordnung mit MAPLE: dsolve	cd
12.6	Beschreibung elektrischer Filterschaltungen	cd
12.6.1	Physikalische Gesetzmäßigkeiten der Bauelemente	cd
12.6.2	Aufstellen der DG für elektrische Schaltungen	cd
12.6.3	Aufstellen und Lösen der DG für Filterschaltungen	cd
16	Numerisches Lösen von Gleichungen	
16.1	Intervallhalbierungs-Methode	cd
16.2	Pegasus-Verfahren	cd
16.3	Banachsches Iterationsverfahren	cd
16.4	Anwendung des Banach-Verfahrens	cd
16.5	Newton-Verfahren	cd
16.6	Regula falsi	cd
17	Numerische Differenziation und Integration	
17.1	Numerische Differenziation	cd
17.2	Numerische Integration	cd
18	Diskrete Fourier-Transformation und Anwendungen	
18.1	Diskrete Fourier-Transformation	cd
18.2	Diskrete Fourier-Transformation mit MAPLE	cd
18.3	Anwendungsbeispiele zur DFT mit MAPLE	cd
19	Partielle Differenzialgleichungen	
19.1	Einführung	cd
19.2	Die Wellengleichung	cd
19.3	Die Wärmeleitungsgleichung	cd
19.4	Die Laplace-Gleichung	cd
19.5	Die zweidimensionale Wellengleichung	cd
19.6	Die Biegeschwingungsgleichung	cd
20	Vektoranalysis und Integralsätze	
20.1	Divergenz und Satz von Gauß	cd
20.2	Rotation und Satz von Stokes	cd
20.3	Rechnen mit Differenzialoperatoren	cd
20.4	Anwendung: Die Maxwellschen Gleichungen	cd
	Lösungen zu den Übungsaufgaben	cd

Kapitel 1

Zahlen, Gleichungen und Gleichungssysteme

1

1	Zahlen, Gleichungen und Gleichungssysteme	1
1.1	Mengen	3
1.2	Natürliche Zahlen	5
1.2.1	Peanosche Axiome	6
1.2.2	Vollständige Induktion	7
1.2.3	Geometrische Summenformel	10
1.2.4	Permutationen	10
1.2.5	Der binomische Lehrsatz	11
1.3	Reelle Zahlen	13
1.3.1	Zahlenmengen und Operationen	13
1.3.2	Die Rechengesetze für reelle Zahlen	14
1.3.3	Potenzrechnen	15
1.3.4	Logarithmen	16
1.3.5	Anordnung der reellen Zahlen	17
1.4	Gleichungen und Ungleichungen	19
1.4.1	Gleichungen	19
1.4.2	Ungleichungen	23
1.5	Lineare Gleichungssysteme	26
1.5.1	Einführung	26
1.5.2	Begriffsbildung und Notation	28
1.5.3	Das Lösen von linearen Gleichungssystemen	29
1.6	Aufgaben zu Zahlen, Gleichungen und Gleichungssystemen	36
Zusätzliche Abschnitte auf der CD-Rom		
1.7	Mathematische Beweismethoden	cd
1.8	MAPLE: Zahlen, Gleichungen und Gleichungssysteme	cd

1 Zahlen, Gleichungen und Gleichungssysteme

Zahlen und Mengen gehören zu den wichtigsten Grundbegriffen der Mathematik, auf denen alle weiteren Gebilde und Konstruktionen aufbauen. In diesem Kapitel werden die Grundlagen sowohl über Mengen als auch über die natürlichen Zahlen gelegt. Zur Beschreibung der natürlichen Zahlen werden die Peanoschen Axiome eingeführt und das Prinzip der vollständigen Induktion an vielen Beispielen demonstriert. Die reellen Zahlen und elementaren Rechengesetze werden angegeben; die Grundgesetze zu den Potenzen und Logarithmen werden wiederholt.

Zu den elementaren Aufgaben der Mathematik gehört das Lösen von Gleichungen. In diesem Kapitel werden auch einfache Gleichungen sowie die für die Anwendungen wichtigen linearen Gleichungssysteme behandelt und der Gauß-Algorithmus eingeführt. Da nur wenige Typen von Gleichungen explizit lösbar sind, werden wir nicht systematisch auf das Lösen von Gleichungen eingehen, sondern exemplarisch zeigen, wie man grundlegende Gleichungen bearbeitet.

Hinweis: Auf der CD-Rom befindet sich ein zusätzlicher Abschnitt über elementare mathematische [Beweismethoden](#) sowie das Lösen von [Gleichungen mit MAPLE](#) inklusive der graphischen Darstellung von [Ungleichungen](#).

1.1 Mengen

1.1

”Unter einer *Menge* M verstehen wir jede Zusammenfassung von bestimmten wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens zu einem Ganzen”; diese Festlegung (=Definition) des Mengenbegriffs stammt von G. Cantor (1895). Diese Definition des Mengenbegriffs reicht für unsere Zwecke vollständig aus. Mengen bezeichnen wir im Folgenden immer mit Großbuchstaben. Die Objekte einer Menge A heißen *Elemente* von A und werden mit Kleinbuchstaben bezeichnet.

$a \in A$ heißt: a ist Element der Menge A .

$a \notin A$ heißt: a ist nicht Element der Menge A .

Mengen werden üblicherweise angegeben durch

- das Auflisten der Elemente in einer Mengenklammer
 $\{a_1, a_2, a_3, a_4, \dots\}$,
- eine Aussageform
 $\{a \in A : a \text{ hat die Eigenschaft } E\}$.

Die *leere Menge* $\emptyset$ bzw. $\{\}$ enthält keine Elemente. B heißt *Teilmenge* von A ($B \subset A$), wenn jedes Element von B auch Element von A ist.

Beispiele 1.1 (Mengen):

$\mathbb{N}$	= Menge der natürlichen Zahlen	$= \{1, 2, 3, 4, \dots\}$.
$\mathbb{N}_0$	= Menge der natürlichen Zahlen mit Null	$= \{0, 1, 2, 3, 4, \dots\}$.
$\mathbb{Z}$	= Menge der ganzen Zahlen	$= \{0, \pm 1, \pm 2, \pm 3, \dots\}$.
$\mathbb{Q}$	= Menge der rationalen Zahlen	$= \{\frac{p}{q} : p \in \mathbb{Z}, q \in \mathbb{N}\}$.
$\mathbb{R}$	= Menge der reellen Zahlen	

Es gilt: $\mathbb{N} \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R}$.

□

Bemerkungen:

- (1) Die Reihenfolge der Elemente einer Menge spielt keine Rolle. Es ist daher $\{a, b, c, d\} = \{d, c, a, b\}$.
- (2) Jedes Element einer Menge wird nur einmal aufgezählt, d.h. $\{a, a, a, b, d\} = \{a, b, d\}$.

② **Mengenoperationen**

Für zwei Mengen A und B sind der *Durchschnitt* $A \cap B$, die *Vereinigung* $A \cup B$ und das *Komplement* $A \setminus B$ definiert durch

$$A \cap B := \{x : x \in A \text{ und } x \in B\},$$

$$A \cup B := \{x : x \in A \text{ oder } x \in B\},$$

$$A \setminus B := \{x : x \in A \text{ und } x \notin B\}.$$

Hierbei bedeutet ":", dass das Symbol auf der linken Seite durch die rechte Seite der Gleichung festgelegt (= definiert) wird. Ähnlich ist das Zeichen ":", als logische Äquivalenz nach Definition dessen, was auf Seiten des Doppelpunktes steht.

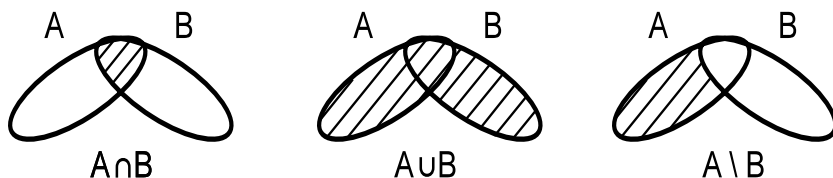


Abb. 1.1. Venn-Diagramme

Durch die sog. *Venn-Diagramme* (siehe Abb. 6.2) lassen sich Mengen und Mengenoperationen schematisch darstellen. Links ist der Schnitt zweier Mengen $A \cap B$, in der Mitte die Vereinigung der Mengen $A \cup B$ und rechts das Komplement der Menge A zur Menge B aufgezeigt. Mit den Venn-Diagrammen lassen sich die folgenden Rechenregeln für Mengen leicht veranschaulichen:

1. $A \cup B = B \cup A$
2. $A \cup A = A$
3. $A \cup (B \cup C) = (A \cup B) \cup C$
4. $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$
5. $(A \setminus B) \setminus C = A \setminus (B \cup C)$
6. $A \subset B \Leftrightarrow A \cap B = A \Leftrightarrow A \cup B = B \Leftrightarrow A \setminus B = \emptyset$

Das *kartesische Produkt* von zwei Mengen M_1 und M_2 ist die Menge, die aus allen Paaren (x, y) besteht, wobei $x \in M_1$ und $y \in M_2$:

$$M_1 \times M_2 := \{(x, y) : x \in M_1 \text{ und } y \in M_2\}.$$

Beispiel 1.2: $\mathbb{R} \times \mathbb{R}$ besteht aus allen Paaren von reellen Zahlen. Dies ist nichts anderes als die Zahlenebene; (x, y) ist jeweils ein Punkt in dieser Ebene. Statt $\mathbb{R} \times \mathbb{R}$ schreibt man kurz $\mathbb{R}^2$. □

1.2 Natürliche Zahlen

1.2

Zu den einfachsten Gegenständen der Arithmetik gehören die natürlichen Zahlen. Sie bilden das Fundament unseres Zahlengebäudes. Die Gesamtheit aller natürlichen Zahlen nennen wir die "natürliche Zahlenmenge" $\mathbb{N}$. Der Ausdruck "natürliche" Zahlen für $\mathbb{N} = \{1, 2, 3, \dots\}$ ist sicherlich gut gewählt, denn Kinder beginnen so zu zählen und in allen Kulturen beginnt das mathematische Denken mit diesen Zahlen. Die Null wurde erst sehr spät, nämlich etwa 870 n. Chr. von den Indern *erfunden* und wird heute zu den natürlichen Zahlen hinzugenommen: $\mathbb{N}_0$.

Erst durch die Entdeckung der Zahl Null war die indische Mathematik erstmals in der Lage, ein heutzutage in der ganzen Welt übernommenes Stellenwertsystem zu schaffen, das nur mit zehn Ziffern (einschließlich der Null) auskommt. Das Stellenwertsystem ist schon bei A. Ries (1492-1559) in seinem ersten Rechenbuch 1422 beschrieben. Darin befindet sich auch die Würdigung der Zahl Null! Das Fundamentalprinzip der natürlichen Zahlen geht auf den Mathematiker Peano (1858-1939, 1889) zurück.

➤ 1.2.1 Peanosche Axiome

- (1) 1 ist eine natürliche Zahl.
- (2) Zu jeder natürlichen Zahl n existiert genau ein Nachfolger n' , der ebenfalls der natürlichen Zahlenmenge angehört.
- (3) Es gibt keine natürliche Zahl, deren Nachfolger 1 ist.
- (4) Die Nachfolger zweier verschiedener natürlicher Zahlen sind voneinander verschieden.
- (5) Eine Teilmenge der natürlichen Zahlen enthält alle natürlichen Zahlen, wenn 1 zur Menge gehört und mit einer natürlichen Zahl n stets auch der Nachfolger n' zur Menge gehört.

Mit den Peanoschen Axiomen ist man in der Lage, die natürliche Zahlenmenge aufzubauen, denn man erhält sofort die folgenden Konsequenzen aus den Axiomen:

Folgerungen:

- (1) Die natürliche Zahlenmenge hat unendlich viele verschiedene Elemente: Wegen (A1) gibt es mindestens eine natürliche Zahl: 1. Wegen (A2) gibt es zu 1 einen Nachfolger, der nach (A3) $\neq 1$: 2. Wegen (A2) gibt es zu 2 einen Nachfolger, der $\neq 1$ (A3) und $\neq 2$ (A4): 3. usw.
- (2) Die Elemente der natürlichen Zahlenmenge lassen sich in einer bestimmten Reihenfolge anordnen, wobei schrittweise alle natürlichen Zahlen erfasst werden:

$$1; 2; 3; 4; 5; \dots, n; n + 1; \dots$$

Hierbei bedeutet $n + 1$ der Nachfolger von n . Durch diese Reihenfolge wird auf natürliche Weise die Addition von natürlichen Zahlen festgelegt.

- (3) Jede Teilmenge M der natürlichen Zahlen $M \subset \mathbb{N}$, welche die 1 und mit $n \in M$ auch stets den Nachfolger $n + 1$ enthält, ist gleich der Menge aller natürlichen Zahlen.

Aus der Folgerung (3) erhalten wir ein Beweisprinzip, welches zu den wichtigsten Beweismethoden der Analysis gehört, nämlich die *vollständige Induktion*.

► 1.2.2 Vollständige Induktion

Um eine Aussage $A(n)$ für alle natürlichen Zahlen zu beweisen, genügt es nach Folgerung (3) zu zeigen:

- (I) **Induktionsanfang:** Für $n = 1$ ist die Aussage richtig.
- (II) **Induktionsschluss** von n_0 auf $n_0 + 1$: Ist die Aussage für eine beliebige natürliche Zahl n_0 gültig, dann muss sie auch für den Nachfolger $n_0 + 1$ richtig sein.

Können beide Schritte durchgeführt werden, dann gilt die Aussage für alle $n \in \mathbb{N}$. Denn nach (I) ist die Aussage für 1 richtig. Nach (II) ist dann die Aussage auch für den Nachfolger 2 richtig. Nach (II) ist dann die Aussage auch für den Nachfolger, also 3, richtig. usw.

Beispiel 1.3.

$$1 + 2 + \cdots + n = \frac{n(n+1)}{2} \quad (n \in \mathbb{N})$$

Beweis mit vollständiger Induktion. Der Induktionsanfang besteht darin, dass man explizit nachprüfen muss, dass die Formel für $n = 1$ richtig ist. Wir setzen daher $n = 1$ sowohl in die linke wie auch rechte Seite der Gleichung ein: Für $n = 1$ ist die linke Seite der Gleichung 1 und die rechte Seite $\frac{1 \cdot 2}{2} = 1$. Damit stimmt die Formel für $n = 1$.

Induktionsschluss von n_0 auf $n_0 + 1$: Sei $n_0 \in \mathbb{N}$ beliebig und die Formel sei richtig für dieses n_0 , d.h. es gilt $1 + 2 + \cdots + n_0 = \frac{n_0(n_0+1)}{2}$. Unter dieser Voraussetzung müssen wir zeigen, dass die Formel dann für $n_0 + 1$ gilt. Wir starten mit der linken Seite der Gleichung; die Summe geht nun bis $n_0 + 1$. Wir setzen die Induktionsvoraussetzung ein und formen den Term so lange um, bis die rechte Seite für $n_0 + 1$ heraus kommt:

$$\begin{aligned} 1 + 2 + \cdots + n_0 + (n_0 + 1) &= (1 + 2 + \cdots + n_0) + (n_0 + 1) \\ &= \frac{n_0(n_0 + 1)}{2} + (n_0 + 1) \\ &= \frac{n_0(n_0 + 1) + 2(n_0 + 1)}{2} \\ &= \frac{(n_0 + 1)(n_0 + 2)}{2}. \end{aligned}$$

Dies ist die zu beweisende Formel für $n_0 + 1$.

□

Anmerkung: Man sagt, dass diese Summenformel auf F. Gauß (1777-1855) zurück geht, der einer Anekdote zufolge die Summe der ersten 100 Zahlen dadurch berechnete, dass er die Summe der ersten mit der letzten, der zweiten mit der zweitletzten, der dritten mit der drittletzten usw. bildete:

$1 + 2 + 3 + \dots + 98 + 99 + 100 = (1 + 100) + (2 + 99) + (3 + 98) + \dots$.
Somit erhält man von den 100 Summanden nur noch $\frac{100}{2}$, jeder mit dem Wert 101, also $1 + 2 + 3 + \dots + 100 = \frac{100 \cdot 101}{2}$. Tatsächlich war sowohl die Formel als auch der Rechenweg schon Adam Ries (1492-1559; 1422) bekannt. $\square$

Als Abkürzung für Summen und Produkte führt man folgende Notationen ein:

Definition:

(1) *Summen:* Für die Summe der Zahlen $a_l, a_{l+1}, \dots, a_n \in \mathbb{R}$ schreibt man

$$\sum_{k=l}^n a_k := a_l + a_{l+1} + \dots + a_n.$$

(2) *Produkt:* Für das Produkt der Zahlen $a_l, a_{l+1}, \dots, a_n \in \mathbb{R}$ schreibt man

$$\prod_{k=l}^n a_k := a_l \cdot a_{l+1} \cdot \dots \cdot a_n.$$

(3) *Fakultät:* Zu jedem $n \in \mathbb{N}$ definiert man

$$n! := 1 \cdot 2 \cdot \dots \cdot n \quad (\textbf{Fakultät von } n) \quad \text{und} \quad 0! := 1.$$

Übrigens wächst $n!$ sehr schnell. Z.B. $13! \approx 6 \cdot 10^9$; um diese Zahl zu zählen, benötigte man 100 Jahre, wenn man in einer Minute bis 100 zählen könnte!

Beispiele 1.4.

$$\textcircled{1} \quad \sum_{i=5}^{10} i^2 = 5^2 + 6^2 + 7^2 + 8^2 + 9^2 + 10^2 = 355.$$

$$\textcircled{2} \quad \sum_{i=1}^5 \frac{1}{2i} = \frac{1}{2 \cdot 1} + \frac{1}{2 \cdot 2} + \frac{1}{2 \cdot 3} + \frac{1}{2 \cdot 4} + \frac{1}{2 \cdot 5} = \frac{137}{120}.$$

$$\textcircled{3} \quad \prod_{i=3}^6 (2i-1)^2 = 5^2 \cdot 7^2 \cdot 9^2 \cdot 11^2 = 12006225.$$

$$\textcircled{4} \quad 5! = 1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 = 120. \quad \square$$

Beispiel 1.5. $1 + 3 + 5 + \cdots + (2n - 1) = \sum_{k=1}^n (2k - 1) = n^2 \quad (n \in \mathbb{N})$

Beweis mit vollständiger Induktion. Wir beginnen beim Induktionsanfang $n = 1$, indem wir $n = 1$ sowohl in die linke wie auch rechte Seite der Gleichung einsetzen: $1 = 1^2$. Damit ist die Formel richtig für $n = 1$.

Induktionsschluss von n auf $n + 1$: Sei n beliebig und die Formel richtig für dieses n , dann ist zu zeigen, dass die Formel auch für $n + 1$ gilt:

$$\begin{aligned} 1 + 3 + 5 + \cdots + (2n - 1) + (2n + 1) &= [1 + 3 + \cdots + (2n - 1)] + (2n + 1) \\ &= n^2 + 2n + 1 = (n + 1)^2. \end{aligned}$$

Dies ist die Formel für $n + 1$. □

Achtung: Zum Beweisprinzip der **vollständigen** Induktion gehören sowohl der Induktionsanfang als auch der Induktionsschluss. Fehlt einer dieser beiden Beweisteile, ist der Beweis nicht vollständig und die Aussage nicht bewiesen, wie die beiden folgenden Beispiele zeigen:

⚠ **Beispiel 1.6.** Nach L. Euler (1707-1783) liefert der Ausdruck

$$p = n^2 - n + 41$$

für $n = 1, 2, 3, \dots, 40$ *Primzahlen*, nämlich $p = 41, 43, 47, \dots, 1601$ wie man durch Einsetzen explizit nachrechnet. Dies reicht aber nicht für einen allgemeinen Beweis aus. Für $n = 41$ folgt $p = 41^2 - 41 + 41 = 41^2$. Dies ist keine Primzahl. Die Primzahlberechnung ist zwar für viele einzelne n richtig; sie gilt aber **nicht** allgemein! □

⚠ **Beispiel 1.7.** Wir betrachten die *falsche* Formel

$$1 + 2 + 3 + \cdots + n = \frac{n(n+1)}{2} + 1$$

(vgl. Beispiel 1.3) und zeigen, dass dennoch der Induktionsschluss durchführbar ist: Wir nehmen also an, die Formel sei für ein n richtig und zeigen, dass sie dann auch für $n + 1$ gültig ist.

$$\begin{aligned} 1 + 2 + 3 + \cdots + n + (n + 1) &= \frac{n(n+1)}{2} + 1 + (n + 1) \\ &= \frac{n(n+1) + 2(n+1)}{2} + 1 \\ &= \frac{(n+1)(n+2)}{2} + 1. \end{aligned}$$

Dies ist die Formel für $n + 1$. Obwohl der Induktionsschluss durchführbar ist, gibt es **keine** natürliche Zahl n , für welche die Formel richtig ist. Der Induktionsschluss verliert also seinen Sinn, wenn der Nachweis für $n = 1$ oder für einen anderen festen Zahlenwert nicht erbracht werden kann. □

1.2.3 Geometrische Summenformel

Für viele Anwendungen wichtig ist die geometrische Summenformel:

Satz (Geometrische Summenformel)

Für jede reelle Zahl $q \neq 1$ gilt:

$$\sum_{i=0}^n q^i = \frac{1 - q^{n+1}}{1 - q} \quad (n \in \mathbb{N}_0)$$

Beweis durch vollständige Induktion.

Induktionsanfang $n = 0$: $\sum_{i=0}^0 q^i = q^0 = 1 = \frac{1 - q^1}{1 - q}$.

Induktionsschluss von n auf $n + 1$: Sei n beliebig und die Formel richtig für n , dann gilt für $n + 1$:

$$\begin{aligned} \sum_{i=0}^{n+1} q^i &= \sum_{i=0}^n q^i + q^{n+1} = \frac{1 - q^{n+1}}{1 - q} + q^{n+1} \\ &= \frac{1 - q^{n+1} + (1 - q)q^{n+1}}{1 - q} = \frac{1 - q^{n+2}}{1 - q}. \end{aligned}$$

□

1.2.4 Permutationen

Als Permutation einer Menge versteht man alle möglichen Anordnungen der Elemente der Menge. Ist $A = \{a_1, a_2, a_3, \dots, a_n\}$, so kommt auf jede Position in der Menge genau ein Element. Eine andere Anordnung der Menge ist z.B.

$$\{a_2, a_1, a_3, \dots, a_n\}.$$

Der folgende Satz gibt Aufschluss darüber, wie groß die Anzahl aller verschiedener Anordnungen einer n -elementigen Menge ist:

Satz: Die Anzahl aller möglichen Anordnungen einer n -elementigen Menge

$$\{a_1, \dots, a_n\} \text{ ist gleich } n! = 1 \cdot 2 \cdot 3 \cdot \dots \cdot n.$$

Beweis durch vollständige Induktion. Der Induktionsanfang ist wieder bei $n = 1$: Die Anzahl aller Anordnungen der 1-elementigen Menge $\{a_1\}$ ist 1. Wegen $1! = 1$ gilt die zu beweisende Formel also für $n = 1$.

Induktionsschluss von n auf $n + 1$: Gesucht ist die Anzahl aller Anordnungen einer $(n + 1)$ -elementigen Menge $\{a_1, a_2, a_3, \dots, a_{n+1}\}$. Dazu betrachten wir das Element a_1 und dessen Plätze (Positionen) in dieser Menge. a_1 kann an 1. Stelle stehen; dann gibt es nach Induktionsvoraussetzung für die restlichen n Elemente $n!$ Anordnungen. a_1 kann auch an 2. Stelle stehen; dann gibt es

nach Induktionsvoraussetzung für die restlichen n Elemente $n!$ Anordnungen. a_1 kann auch an 3. Stelle stehen; wieder gibt es dann für die restlichen n Elemente $n!$ Anordnungen. usw. a_1 kann also an $n + 1$ verschiedenen Positionen stehen, und die verbleibenden n Elemente haben dann noch $n!$ unterschiedliche mögliche Anordnungen. Insgesamt gibt es also $n! \cdot (n + 1) = (n + 1)!$ Möglichkeiten. $\square$

Folgerung: Die Anzahl der k -elementigen Teilmengen einer n -elementigen Menge $\{a_1, \dots, a_n\}$ ist gleich

$$\frac{n!}{k!(n-k)!}.$$

Begründung: Alle k -elementigen Teilmengen der Menge $M = \{a_1, \dots, a_n\}$ findet man, indem aus allen $n!$ Anordnungen nur die ersten k Elemente genommen werden. Dabei tritt jede k -elementige Teilmenge $k!$ -mal auf und die verbleibenden $(n - k)!$ -mal. Damit sind die k -elementigen Teilmengen gleich $\frac{n!}{k!(n-k)!}$. $\square$

Anwendung: Die Chance, beim Lotto-Spiel "6 aus 49" die richtige Kombination zu erraten, ist etwa 1:14 Millionen. Denn die Anzahl der 6-elementigen Teilmengen einer 49-elementigen Menge ist gleich $\frac{49!}{6!43!} = \frac{44 \cdot 45 \cdot 46 \cdot 47 \cdot 48 \cdot 49}{1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 \cdot 6} = 13.983.816$.

► 1.2.5 Der binomische Lehrsatz

Für zwei natürliche Zahlen n und k mit $0 \leq k \leq n$ bezeichnet man die Zahl

$$\binom{n}{k} := \frac{n!}{k!(n-k)!}$$

(man spricht n über k) als *Binomialkoeffizient*. Die Binomialkoeffizienten bestimmen sich entweder durch die obige Formel oder durch das nach Pascal benannte Schema, dem sog. **Pascalschen-Dreieck**: Beginnend mit 1 wird die unten angegebene Pyramide in jeder Stufe um eine 1 rechts und links erweitert. Die Zahlen im Schema ergeben sich aus der Summe der beiden darüber stehenden Zahlen.

$\binom{0}{k} :$					1												
$\binom{1}{k} :$					1		1										
$\binom{2}{k} :$					1		2		1								
$\binom{3}{k} :$					1		3		3		1						
$\binom{4}{k} :$					1		4		6		4		1				
$\binom{5}{k} :$					1		5		10		10		5		1		
$\binom{6}{k} :$					1		6		15		20		15		6		1
$\vdots$																	

Bemerkung: Mit den Binomialkoeffizienten können wir die letzte Folgerung kurz formulieren: Es gibt $\binom{n}{k}$ Möglichkeiten aus n Objekten genau k auszuwählen. Aus dieser Aussage erhalten wir die binomische Formel:

Satz: (Binomischer Lehrsatz). Für beliebige Zahlen $a, b \in \mathbb{R}$ und jede natürliche Zahl $n \geq 0$ gilt:

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k$$

Beweis: Multipliziert man die rechte Seite aus, so kommt der Term b^k so oft vor, wie man k Faktoren aus n Faktoren wählen kann, also $\binom{n}{k}$ -mal (siehe obige Bemerkung). Die restlichen $(n - k)$ Faktoren tragen zu a^{n-k} bei. $\square$

Beispiele 1.8.

① $(x + y)^0 = 1$

$$(x + y)^1 = x + y$$

$$(x + y)^2 = x^2 + 2xy + y^2$$

$$(x + y)^3 = x^3 + 3x^2y + 3xy^2 + y^3$$

② Wir berechnen den Wert der Potenz $(104)^3$ mit dem Lehrsatz:

$$\begin{aligned} (104)^3 &= (100 + 4)^3 = 100^3 + 3 \cdot 100^2 \cdot 4 + 3 \cdot 100 \cdot 4^2 + 4^3 \\ &= 1\,000\,000 + 120\,000 + 4\,800 + 64 = 1.124.864. \end{aligned}$$

$\square$

1.3 Reelle Zahlen

Wir stellen uns auf den Standpunkt, dass uns die reellen Zahlen zur Verfügung stehen und gehen nicht auf den axiomatischen Aufbau ein. Für physikalische Messungen würden die rationalen Zahlen ausreichen, für die höhere Analysis weisen die rationalen Zahlen "zu viele Löcher" auf. Erst ihre Erweiterung zu den reellen Zahlen macht die Differenzial- und Integralrechnung möglich.

1.3.1 Zahlenmengen und Operationen

Auf den natürlichen Zahlen $\mathbb{N}$ gibt es als Grundrechenarten $+$ und $\cdot$. Die Gleichung $x + 1 = 0$ ist innerhalb $\mathbb{N}$ formulierbar, aber nicht lösbar. Man erweitert daher den Zahlenbereich um all die Lösungen der Gleichungen

$$x + n = 0,$$

wenn $n \in \mathbb{N}_0$. Die Lösungen sind $0, -1, -2, -3, \dots$ und der erweiterte Zahlenbereich nennt man $\mathbb{Z}$, die ganzen Zahlen. In $\mathbb{Z}$ lässt sich für jedes $n \in \mathbb{Z}$ die Gleichung $x + n = 0$ lösen. Nicht lösbar ist die Gleichung $2x = 1$. Man erweitert nun $\mathbb{Z}$ um all die Lösungen von Gleichungen der Form

$$q \cdot x = p$$

mit $p, q \in \mathbb{Z}$ und $q \neq 0$. Somit erhält man die Zahlenmenge der rationalen Zahlen $\mathbb{Q}$. In diesem Zahlenbereich sind alle Gleichungen obiger Form lösbar. Aber die Gleichung

$$x^2 = 2$$

besitzt in $\mathbb{Q}$ keine Lösung. Also erweitert man die rationalen Zahlen um all die Lösungen von Gleichungen obiger Bauart und kommt so zu den reellen Zahlen $\mathbb{R}$. In den reellen Zahlen sind noch die sog. *transzendenten* Zahlen, wie e und π enthalten, die wir im Kapitel über Folgen noch genauer untersuchen. In Tabelle 1 sind die Zahlenbereiche mit den zugehörigen Rechenoperationen nochmals aufgelistet.

Tabelle 1: Zahlenmengen und Grundrechenoperationen

Mengen	Grundoperationen	nicht lösbar
$\mathbb{N}_0$ natürliche Zahlen	$+$ $\cdot$	$x + 1 = 0$
$\mathbb{Z}$ ganze Zahlen	$+$ $-$ $\cdot$	$2 \cdot x = 1$
$\mathbb{Q}$ rationale Zahlen	$+$ $-$ $\cdot$ $\backslash$	$x^2 = 2$
$\mathbb{R}$ reelle Zahlen	$+$ $-$ $\cdot$ $\backslash$	$x^2 + 1 = 0$

Darstellung reeller Zahlen. Zur Veranschaulichung der reellen Zahlen dient die bekannte von $-$ nach $+$ gerichtete *Zahlengerade*. Jeder Punkt auf der Zahlengeraden entspricht genau einer reellen Zahl.

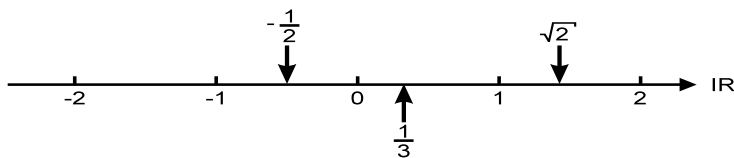


Abb. 1.2. Reelle Zahlengerade

➤ 1.3.2 Die Rechengesetze für reelle Zahlen

In $\mathbb{R}$ sind zwei Verknüpfungen gegeben, nämlich $+$ und $\cdot$. Addition und Multiplikation zweier reeller Zahlen liefern wieder reelle Zahlen. Formal hat man hiermit zwei Abbildungen $+$ und $\cdot$ definiert:

$$+ : \mathbb{R} \times \mathbb{R} \longrightarrow \mathbb{R} \quad \text{mit} \quad (x, y) \longmapsto x + y$$

$$\cdot : \mathbb{R} \times \mathbb{R} \longrightarrow \mathbb{R} \quad \text{mit} \quad (x, y) \longmapsto x \cdot y.$$

Es gelten die **Rechengesetze der Addition**

- | | | |
|------|--|--------------------------|
| (A1) | $x + (y + z) = (x + y) + z$ | <i>Assoziativgesetz</i> |
| (A2) | $x + y = y + x$ | <i>Kommutativgesetz</i> |
| (A3) | $x + 0 = x$ | <i>Existenz der Null</i> |
| (A4) | Zu jedem x gibt es ein $(-x) \in \mathbb{R}$ mit
$x + (-x) = 0$ | <i>Inverses Element</i> |

Es gelten die **Rechengesetze der Multiplikation**

- | | | |
|------|---|--------------------------|
| (M1) | $x \cdot (y \cdot z) = (x \cdot y) \cdot z$ | <i>Assoziativgesetz</i> |
| (M2) | $x \cdot y = y \cdot x$ | <i>Kommutativgesetz</i> |
| (M3) | Es gibt eine Zahl $1 \in \mathbb{R}$ mit $1 \neq 0$, so dass
$1 \cdot x = x$ | <i>Existenz der Eins</i> |
| (M4) | Zu jedem $x \in \mathbb{R} \setminus \{0\}$ gibt es ein $x^{-1} \in \mathbb{R}$ mit
$x \cdot x^{-1} = 1$ | <i>Inverses Element</i> |

Es gilt das **Distributivgesetz**

$$(D) \quad x \cdot (y + z) = x \cdot y + x \cdot z$$

Alle weiteren Rechengesetze der reellen Zahlen lassen sich auf diese elementaren Gesetze zurückführen. Da diese Rechengesetze nicht nur für die Menge der reellen Zahlen gelten, sondern auch für andere Zahlengebilde, führt man den Begriffs des Körpers ein und verallgemeinert: Eine Menge K zusammen mit

zwei Verknüpfungen

$$+ : K \times K \longrightarrow K \quad \text{mit} \quad (x, y) \longmapsto x + y$$

$$\cdot : K \times K \longrightarrow K \quad \text{mit} \quad (x, y) \longmapsto x \cdot y,$$

die den Axiomen (A1)-(A4), (M1)-(M4) und (D) genügen, nennt man **Körper**.

Beispiele 1.9:

- ① Sowohl $(\mathbb{R}, +, \cdot)$ als auch $(\mathbb{Q}, +, \cdot)$ bilden Körper.
- ② $(\mathbb{Z}, +, \cdot)$ ist kein Körper, da (M4) verletzt ist: z.B. $2 \in \mathbb{Z}$ besitzt bezüglich der Multiplikation kein Inverses, so dass $2 \cdot x = 1$.
- ③ $(\mathbb{N}, +, \cdot)$ ist kein Körper, da z.B. (A4) verletzt ist.
- ④ $(F_2, +, \cdot)$ mit $F_2 = \{0, 1\}$ und den Verknüpfungen

+	0	1
0	0	1
1	1	0

·	0	1
0	0	0
1	0	1

ist ein Körper. Man rechnet die Rechengesetze direkt nach. F_2 ist der kleinste Körper; denn jeder Körper muss mindestens zwei Elemente enthalten: 0 und 1. □

► 1.3.3 Potenzrechnen

Wir definieren zu jeder reellen Zahl $a \in \mathbb{R}$ die **Potenz** von a durch

$$a^0 := 1, \quad a^1 := a, \quad a^n := \underbrace{a \cdot \dots \cdot a}_{n\text{-mal}} \quad (n \in \mathbb{N}).$$

Definition: Die n -te **Wurzel** einer Zahl $a \geq 0$

$$b := \sqrt[n]{a} := a^{\frac{1}{n}} \quad (n \in \mathbb{N})$$

ist definiert als diejenige positive reelle Zahl b mit der Eigenschaft $b^n = a$.

Mit der Schreibweise $\sqrt[n]{a} = a^{\frac{1}{n}}$ lassen sich die n -ten Wurzeln einer Zahl als Potenzen mit rationalem Exponenten interpretieren. Für die Potenzen von Produkten bzw. Quotienten von reellen Zahlen gelten die allgemeinen Potenzrechenregeln, die in der folgenden Übersicht zusammengestellt sind:

Es gelten die **Potenzrechenregeln**

$$(1) a^n \cdot b^n = (a \cdot b)^n$$

$$(2) \frac{a^n}{b^n} = \left(\frac{a}{b}\right)^n \text{ für } (b \neq 0)$$

$$(3) \frac{a^n}{a^m} = a^{n-m} \text{ für } (a \neq 0)$$

$$(4) (a^m)^n = a^{n \cdot m}$$

$$(5) \sqrt[n]{a^m} = a^{m/n} \text{ für } (a \geq 0)$$

$$(n, m \in \mathbb{N})$$

Beispiele 1.10.

$$\textcircled{1} \quad \frac{a^{5x-2y}}{b^{6m-1}} : \frac{a^{4x+y}}{b^{m-2}} = \frac{a^{5x-2y-4x-y}}{b^{6m-1-m+2}} = \frac{a^{x-3y}}{b^{5m+1}}.$$

$$\textcircled{2} \quad \frac{(a^2 b)^2}{2a\sqrt{ab}} = a^4 b^2 \frac{1}{2} a^{-1} a^{-\frac{1}{2}} b^{-\frac{1}{2}} = \frac{1}{2} a^{\frac{5}{2}} b^{\frac{3}{2}}.$$

$$\textcircled{3} \quad \frac{(8a^3 b^{-3})^{-2}}{(12a^{-2} b^{-4})^{-3}} = \frac{8^{-2} a^{-6} b^6}{12^{-3} a^6 b^{12}} = \frac{3^3}{2^6 a^{12} b^6}.$$

□

1.3.4 Logarithmen

Definition: Gegeben ist die Gleichung $a = b^x$ ($a, b > 0$). Gesucht ist bei gegebenem a und b der Exponent x . Wir nennen

$$x = \log_b a$$

den **Logarithmus von a zur Basis b** .

Für feste Basis b gelten die **Logarithmenrechenregeln**

$$(1) \log(u \cdot v) = \log(u) + \log(v) \quad (u, v > 0)$$

$$(2) \log\left(\frac{u}{v}\right) = \log(u) - \log(v) \quad (u, v > 0)$$

$$(3) \log(u^n) = n \cdot \log(u) \quad (u > 0)$$

Spezielle Logarithmen sind der Logarithmus zur Basis 10

$$\log a := \log_{10} a \text{ (10er Logarithmus),}$$

der Logarithmus zur Basis 2 (Logarithmus dualis)

$$\lg a := \log_2 a \text{ (2er Logarithmus)}$$

und der Logarithmus zu Basis e

$$\ln a := \log_e a \quad (\text{natürlicher Logarithmus}).$$

Zwischen unterschiedlichen Logarithmen besteht der Zusammenhang

$$\log_b y = \frac{\log_c y}{\log_c b} \quad (b, c, y > 0).$$

Dadurch ist es ausreichend einen Logarithmus (z.B. den natürlichen Logarithmus) berechnen zu können. Die Logarithmen zu anderen Basen ergeben sich dann durch obige Formel.

Beweis der Logarithmenformel: Aus $b^x = y$ folgt per Definition des Logarithmus zur Basis b , dass $x = \log_b y$. Andererseits gilt für den Logarithmus zur Basis c nach der Logarithmusregel (3):

$$\log_c y = \log_c b^x = x \cdot \log_c b \Rightarrow x = \frac{\log_c y}{\log_c b}.$$

Hieraus folgt die behauptete Formel. □

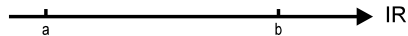
Beispiele 1.11:

- ① $2^x = \frac{1}{8} \Rightarrow x = \log_2 \frac{1}{8} = -\log_2 8 = -3.$
- ② $10^x = 0.0001 \Rightarrow x = \log_{10} 10^{-4} = -4 \log_{10} 10 = -4.$
- ③ $\ln \frac{\sqrt{ab^{-2}}}{\sqrt[3]{cd^{-3}}} = \ln \sqrt{a} + \ln b^{-2} - \ln c^{\frac{1}{3}} - \ln d^{-3} = \frac{1}{2} \ln a - 2 \ln b - \frac{1}{3} \ln c + 3 \ln d.$
- ④ $\log \sqrt[3]{a^2 b \sqrt[4]{a c^2}} = \log((a^2 b a^{\frac{1}{4}} c^{\frac{1}{2}})^{\frac{1}{3}})^{\frac{1}{2}} = \log a^{\frac{1}{3}} b^{\frac{1}{6}} a^{\frac{1}{24}} c^{\frac{1}{12}}$
 $= \frac{9}{24} \log a + \frac{1}{6} \log b + \frac{1}{12} \log c.$ □

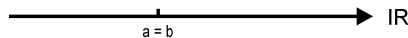
► 1.3.5 Anordnung der reellen Zahlen

Unter den reellen Zahlen herrscht eine bestimmte *Anordnung*: Zwei reelle Zahlen $a, b \in \mathbb{R}$ stehen stets in genau einer der drei folgenden Beziehungen zueinander:

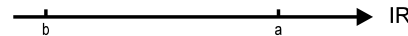
$a < b$ (a liegt links von b),



$a = b$ (a identisch mit b),



$a > b$ (a liegt rechts von b).



Unter dem **Betrag** einer reellen Zahl a wird der Abstand von a zum Nullpunkt verstanden. Er wird durch das Symbol $|a|$ gekennzeichnet:

$$|a| := \begin{cases} a & \text{für } a > 0 \\ 0 & \text{für } a = 0 \\ -a & \text{für } a < 0 \end{cases}$$

Beispiele 1.12: $|3| = 3$; $|-5| = 5$; $|\frac{1}{2}| = \frac{1}{2}$; $|\sqrt{2}| = \sqrt{2}$. □

Der **Abstand** zweier Zahlen x und a auf der Zahlengerade ist $|x - a|$.

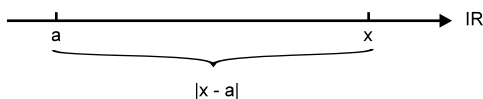


Abb. 1.3. Abstand zweier Zahlen

Nach der Definition des Betrags ist also

$$|x - a| = x - a \text{ für } x - a \geq 0 \text{ bzw. } x \geq a,$$

$$|x - a| = -(x - a) \text{ für } x - a < 0 \text{ bzw. } x < a.$$

Diese Fallunterscheidung wird beim Lösen von Betragsgleichungen und Betragsungleichungen eine wichtige Rolle spielen.

Es gelten die **Gesetze**:

(1) $x > 0, y > 0 \Rightarrow x + y > 0$.

(2) $x > 0, y > 0 \Rightarrow x \cdot y > 0$.

(3) Sind $x > 0, y > 0$, dann gibt es immer eine natürliche Zahl $n \in \mathbb{N}$ mit

$$n x > y$$

(Archimedes Axiom).

Folgerung: (Bernoullische Ungleichung)

$$(1 + x)^n \geq 1 + n x \quad (x \geq -1 \text{ und } n \in \mathbb{N})$$

Beweis durch vollständige Induktion. Für $n = 1$ gilt sogar die Gleichheit. Induktionsschluss von n auf $n + 1$: Wegen $1 + x > 0$ folgt durch Multiplikation der Induktionsvoraussetzung $(1 + x)^n \geq 1 + n x$ mit $(1 + x)$:

$$(1 + x)^{n+1} \geq (1 + n x)(1 + x) = 1 + (n + 1)x + n x^2 \geq 1 + (n + 1)x. \quad \square$$

Intervalle: Teilmengen der reellen Zahlen nennt man Intervalle. Dabei unterscheidet man zwischen den endlichen und den unendlichen Intervallen. Zur Beschreibung dieser Teilmengen von $\mathbb{R}$ führen wir folgende Notationen ein:

(1) **Endliche Intervalle** ($a < b$)

$[a, b]$	$:= \{x : a \leq x \leq b\}$	abgeschlossenes Intervall
$[a, b)$	$:= \{x : a \leq x < b\}$	halb offene Intervalle
$(a, b]$	$:= \{x : a < x \leq b\}$	
(a, b)	$:= \{x : a < x < b\}$	offenes Intervall

(2) **Unendliche Intervalle**

$\mathbb{R}_{\geq a}$	$:= [a, \infty)$	$:= \{x : a \leq x < \infty\}$
$\mathbb{R}_{> a}$	$:= (a, \infty)$	$:= \{x : a < x < \infty\}$
$\mathbb{R}_{\leq a}$	$:= (-\infty, a]$	$:= \{x : -\infty < x \leq a\}$
$\mathbb{R}_{< a}$	$:= (-\infty, a)$	$:= \{x : -\infty < x < a\}$

1.4 Gleichungen und Ungleichungen

1.4

Die Lösungsmethoden beim Lösen von Gleichungen sind so vielfältig wie es Gleichungstypen gibt. Wir zeigen exemplarisch, wie einfache Gleichungen und Ungleichungen zu lösen sind. Allerdings werden wir nicht systematisch auf das Lösen von Gleichungen und Ungleichungen eingehen, da sie in vielen Fällen nicht exakt gelöst werden können und man daher auf numerische Verfahren angewiesen ist (siehe z.B. das Bisektionsverfahren §6.4 oder das Newton-Verfahren §7.9).

Hinweis: Auf der CD-Rom befindet sich ein zusätzliches Kapitel über das [numerische Lösen von Gleichungen](#) sowie ein Abschnitt über das Lösen von [Gleichungen](#) und [Ungleichungen](#) mit MAPLE.

➤ 1.4.1 Gleichungen

Jede Beziehung zwischen (reellen) Größen, in der ein Gleichheitszeichen auftritt, nennt man **Gleichung**. Treten die Größen der Gleichung nur als Summen und Produkte von Potenzen auf, so spricht man von **algebraischen Gleichungen**. Die größte Summe aller Exponenten eines Terms gibt den *Grad* der Gleichung an.

② 1. Quadratische Gleichungen

Die Gleichung

$$x^2 + px + q = 0$$

heißt quadratische Gleichung oder Gleichung zweiten Grades. Liegt die Gleichung in der Form $ax^2 + bx + c = 0$ mit $a \neq 0$ vor, so dividiert man durch a , um sie in der p/q -Form zu erhalten.

Für die quadratische Gleichung $x^2 + px + q = 0$ gibt es die p/q -Lösungsformel

$$x_{1/2} = -\frac{p}{2} \pm \sqrt{\frac{p^2}{4} - q}.$$

Setzt man $D := \frac{p^2}{4} - q$ (*Diskriminante*), so hat die Gleichung für $D > 0$ zwei verschiedene reelle Lösungen, für $D = 0$ eine doppelte reelle Lösung und für $D < 0$ keine reelle (aber zwei verschiedene komplexe) Lösungen.

Beispiele 1.13:

- ① $x^2 + 2x - 3 = 0$ hat zwei reelle Lösungen:
 $x_{1/2} = -1 \pm \sqrt{(-1)^2 + 3} = -1 \pm 2$. Damit ist $x_1 = 1$, $x_2 = -3$.
- ② $x^2 + 4x + 4 = 0$ hat eine doppelte reelle Lösung:
 $x_{1/2} = -2 \pm \sqrt{2^2 - 4} = -2$. Damit ist $x_1 = -2$ eine doppelte Lösung.
- ③ $x^2 - 4x + 13 = 0$ hat keine reellen (aber zwei komplexe) Lösungen:
 $x_{1/2} = 2 \pm \sqrt{4 - 13} = 2 \pm \sqrt{-9} = 2 \pm 3i$. Damit gibt es keine reellen Lösungen. Hierbei bedeutet i die imaginäre Einheit (siehe Kapitel 5, Komplexe Zahlen). □

② 2. Gleichungen höheren Grades

Gleichungen höheren Grades sind mathematisch nur zum Teil exakt lösbar, da in den seltensten Fällen eine geschlossene Lösung existiert. Gelegentlich führt eine Substitution zu einer quadratischen Gleichung, wie das folgende Beispiel zeigt:

Beispiel 1.14. Gesucht sind Lösungen der Gleichung vierten Grades

$$x^4 - 5x^2 + 4 = 0.$$

Wir substituieren $z = x^2$ und erhalten die quadratische Gleichung

$$z^2 - 5z + 4 = 0.$$

Wenden wir die p/q -Formel an, ist

$$z_{1/2} = \frac{5}{2} \pm \sqrt{\frac{25}{4} - 4} = \frac{5}{2} \pm \sqrt{\frac{25 - 16}{4}} = \frac{5}{2} \pm \frac{3}{2}.$$

Daher ist $z_1 = 4$, $z_2 = 1$. Wegen $z = x^2$ und damit $x = \pm\sqrt{z}$, gilt

$$x_{1/2} = \pm\sqrt{4} = \pm 2$$

$$x_{3/4} = \pm\sqrt{1} = \pm 1.$$

Die Lösungen der Gleichung sind $x = -2, -1, 1, 2$. □

Manchmal kann man auch eine Nullstelle erraten und anschließend durch Polynomdivision oder durch Anwenden des Horner-Schemas das Problem um einen Grad erniedrigen. Die Rechentechniken hierfür werden wir im Kapitel über Polynome §4.2 ausführlich beschreiben.

➤ 3. Wurzelgleichungen

Einfache Wurzelgleichungen werden gelöst, indem man die Wurzel isoliert, z.B. auf die linke Seite der Gleichung bringt, die restlichen Terme auf die rechte Seite. Anschließend wird bei Quadratwurzelgleichungen quadriert und nach der gesuchten Variablen aufgelöst.

⚠ Achtung: Durch das Quadrieren der Gleichungen verändert man die Lösungsmenge: Die Gleichung $x = 1$ hat als Lösung die Zahl 1. Quadriert man die Gleichung jedoch, erhält man $x^2 = 1$. Diese Gleichung hat als Lösungen sowohl $x = 1$ als auch $x = -1$. Die quadrierten Gleichungen können mehr Lösungen als die ursprünglichen besitzen. Man muss daher immer eine Probe durchführen, ob die gefundenen Kandidaten auch wirklich Lösungen der Wurzelgleichung darstellen.

Beispiel 1.15. Gesucht sind Lösungen der Wurzelgleichung

$$\sqrt{5-x} + 3 - x = 0.$$

Wir isolieren die Wurzel

$$\sqrt{5-x} = x - 3,$$

quadrieren

$$5 - x = x^2 - 6x + 9$$

und formen um

$$x^2 - 5x + 4 = 0.$$

Dies ist eine quadratische Gleichung, die mit der p/q -Formel gelöst werden kann. Beispiel 1.14 liefert die Lösungen

$$x_1 = 4, \quad x_2 = 1.$$

Nun werden wir nachprüfen, ob es bei den beiden Werten um Lösungen der Wurzelgleichung handelt. Dazu setzen wir die Werte in die Wurzelgleichung ein:

$x_1 = 4$: $\sqrt{1} - 1 = 0$. Der Wert $x_1 = 4$ erfüllt die Wurzelgleichung.

$x_2 = 1$: $\sqrt{4} + 2 = 4 \neq 0$. Der Wert $x_1 = 1$ erfüllt die Wurzelgleichung **nicht**.

Damit ist die Lösungsmenge der Wurzelgleichung $\mathbb{L} = \{4\}$. $\square$

④ 4. Betragsgleichungen

Da der Betrag $|a|$ einer reellen Zahl a definiert ist durch den Abstand zum Nullpunkt

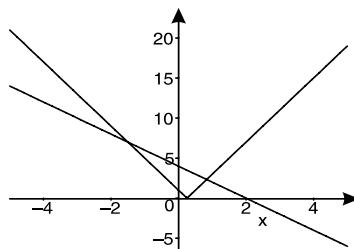
$$|a| := \begin{cases} a & \text{für } a \geq 0 \\ -a & \text{für } a < 0 \end{cases},$$

führt dies bei Betragsgleichungen immer zu einer Fallunterscheidung, je nachdem ob in den Betragszeichen eine positive Zahl steht, dann wird das Betragszeichen durch eine Klammer ersetzt, oder ob in den Betragszeichen eine negative Zahl steht, dann wird das Betragszeichen durch $-()$ ersetzt.

Beispiel 1.16 (Mit MAPLE-Worksheet). Gesucht sind die Lösungen der Betragsgleichung

$$|4x - 1| = -2x + 4.$$

Um sich einen Überblick über die beiden Funktionen zu verschaffen, zeichnen wir die linke und die rechte Seite der Gleichung:



Man erkennt, dass es zwei Schnittpunkte der Graphen gibt, die es zu bestimmen gilt.

1. Fall: $4x - 1 \geq 0$, d.h. $x \geq \frac{1}{4}$:

In den Betragszeichen stehen für $x \geq \frac{1}{4}$ nichtnegative Zahlen. Unter dieser Voraussetzung können die Betragszeichen durch eine einfache Klammer

ersetzt werden. In diesem Fall folgt

$$4x - 1 = -2x + 4.$$

Das Auflösen nach x liefert

$$6x = 5 \quad \text{bzw.} \quad x = \frac{5}{6}.$$

$x = \frac{5}{6}$ erfüllt die Bedingung $x \geq \frac{1}{4}$, unter der wir das Ergebnis berechnet haben.

2. Fall: $4x - 1 < 0$, d.h. $x < \frac{1}{4}$:

In den Betragszeichen stehen für $x < \frac{1}{4}$ nur negative Zahlen. Unter dieser Voraussetzung werden die Betragszeichen durch $-()$ ersetzt. In diesem Fall folgt

$$-(4x - 1) = -2x + 4 \quad \Rightarrow \quad -4x + 1 = -2x + 4.$$

Das Auflösen nach x liefert

$$2x = -3 \quad \text{bzw.} \quad x = -\frac{3}{2}.$$

$x = -\frac{3}{2}$ erfüllt die Bedingung $x < \frac{1}{4}$, unter mit der wir das Ergebnis berechnet haben.

Die Lösungsmenge ist damit $\mathbb{L} = \{\frac{5}{6}, -\frac{3}{2}\}$. □

Hinweis: Durch den verbreiteten Einsatz von Computern zur Lösung von mathematischen Fragestellungen, insbesondere dem Lösen von Gleichungen, kommen numerischen Methoden immer mehr Bedeutung zu. Deshalb sind in diesem Lehrbuch eigens zwei Abschnitte (§6.4 und §7.9) gewidmet, in denen numerische Verfahren detailliert besprochen werden.

➤ 1.4.2 Ungleichungen

Äquivalente Umformungen einer Ungleichung sind:

- Addition (bzw. Subtraktion) eines beliebigen Terms auf beiden Seiten der Ungleichung.
- Multiplikation (bzw. Division) beider Seiten mit einer positiven Zahl $K > 0$.
- Multiplikation (bzw. Division) beider Seiten mit einer negativen Zahl $K < 0$; dabei ändert sich das Ungleichheitszeichen von $(< \text{ zu } >)$, $(\leq \text{ zu } \geq)$, $(> \text{ zu } <)$, $(\geq \text{ zu } \leq)$.

Beispiel 1.17 (Mit [MAPLE-Worksheet](#)). Gesucht sind die Lösungen der Betragsgleichung

$$|2x + 2| > 3.$$

Wie bei den Betragsgleichungen müssen beim Lösen von Ungleichungen die Beträge erst durch eine Fallunterscheidung aufgelöst werden.

1. Fall: $2x + 2 \geq 0$, d.h. $x \geq -1$:

In den Betragszeichen stehen für $x \geq -1$ nicht negative Zahlen. Unter dieser Voraussetzung können die Betragszeichen durch eine einfache Klammer ersetzt werden. In diesem Beispiel folgt dann

$$2x + 2 > 3.$$

Das Auflösen nach x liefert

$$x > \frac{1}{2}.$$

$x > \frac{1}{2}$ erfüllt die Bedingung $x \geq -1$, unter der wir die Rechnung durchgeführt haben.

$$\Rightarrow \mathbb{L}_1 = \left(\frac{1}{2}, \infty\right).$$

2. Fall: $2x + 2 < 0$, d.h. $x < -1$:

In den Betragszeichen stehen für $x < -1$ nur negative Zahlen. Unter dieser Voraussetzung werden die Betragszeichen durch $-()$ ersetzt. In diesem Fall folgt

$$-(2x + 2) > 3.$$

Durch Multiplikation mit (-1) folgt

$$2x + 2 < -3$$

und das Auflösen nach x liefert

$$x < -\frac{5}{2}.$$

$x = -\frac{5}{2}$ erfüllt die Bedingung $x < -1$, unter der wir das Ergebnis berechnet haben.

$$\Rightarrow \mathbb{L}_2 = \left(-\infty, -\frac{5}{2}\right).$$

Die Lösungsmenge besteht damit aus zwei Teilintervallen, dem offenen Intervall $(-\infty, -\frac{5}{2})$ vereinigt mit dem offenen Intervall $(\frac{1}{2}, \infty)$:

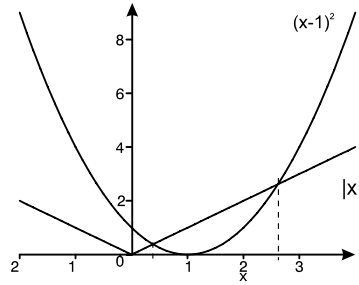
$$\mathbb{L} = \mathbb{L}_1 \cup \mathbb{L}_2 = \left(-\infty, -\frac{5}{2}\right) \cup \left(\frac{1}{2}, \infty\right).$$

□

Beispiel 1.18 (Mit MAPLE-Worksheet). Gesucht sind die Lösungen der Betragsungleichung

$$(x - 2)^2 \leq |x|.$$

Um sich einen Überblick über die beiden Funktionen zu verschaffen, zeichnen wir die linke und die rechte Seite der Ungleichung:



Man erkennt, dass es zwei Schnittpunkte der Graphen gibt, die es zu bestimmen gilt. Die Lösungsmenge besteht dann aus dem abgeschlossenen Intervall, in dem die Quadratfunktion $(x - 1)^2$ kleiner bzw. gleich der Betragsfunktion $|x|$ ist.

Bestimmung der Schnittpunkte: Aus dem Graphen entnimmt man, dass die Schnittpunkte sich im positiven x -Bereich befinden. In diesem Bereich darf man bei $|x|$ die Betragszeichen durch eine einfache Klammer ersetzen, so dass wir die Gleichung

$$x = (x - 1)^2 \quad \Rightarrow \quad x^2 - 2x + 1 = x \quad \Rightarrow \quad x^2 - 3x + 1 = 0$$

zu lösen haben. Die p/q -Formel liefert

$$x_{1/2} = \frac{3}{2} \pm \sqrt{\frac{9}{4} - 1} = \frac{3}{2} \pm \sqrt{\frac{5}{4}}.$$

Damit sind $x_1 = 0.38$ und $x_2 = 2.62$ die Schnittpunkte der Kurven. Die Lösungsmenge besteht aus dem beidseitig abgeschlossenen Intervall

$$\mathbb{L} = \left[\frac{3}{2} - \frac{1}{2}\sqrt{5}, \frac{3}{2} + \frac{1}{2}\sqrt{5} \right] = [0.38, 2.62]. \quad \square$$



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich die MAPLE-Prozedur [visual.solve](#), um Gleichungen und Ungleichungen zu lösen sowie die Lösung graphisch darzustellen.

1.5 Lineare Gleichungssysteme

Lineare Gleichungssysteme (LGS) spielen in Theorie und Anwendungen eine sehr wichtige Rolle. In diesem Abschnitt führen wir eine Methode ein, mit der beliebige LGS gelöst werden können: den *Gauß-Algorithmus*. Auf allgemeine Zusammenhänge und Aussagen über LGS sei auf das Kapitel 3, Matrizen und Determinanten, verwiesen.

1.5.1 Einführung

Anwendungsbeispiel 1.19 (Die Beschreibung elektrischer Netzwerke bei Gleichströmen).

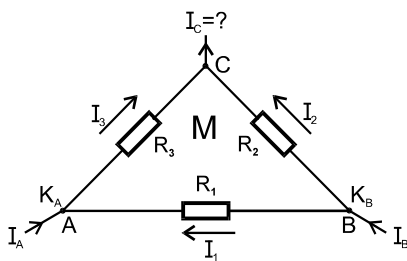


Abb. 1.4. Elektrisches Netzwerk

Gegeben sei das nebenstehende *elektrische Netzwerk* mit den gegebenen Widerständen $R_1 = 1\Omega$, $R_2 = 5\Omega$, $R_3 = 3\Omega$. Diesem Netzwerk werden zwei Gleichströme $I_A = 1A$ und $I_B = 2A$ zugeführt. Gesucht sind die Einzelströme I_1, I_2, I_3 .

Um die Modellgleichungen zu erhalten, wenden wir die **Kirchhoffschen Gesetze** an: Der *Knotensatz* besagt, dass die

Summe der in einem Knoten zu- und abfließenden Ströme gleich Null ist. Der *Maschensatz* besagt, dass in einer Masche die Summe aller Spannungen Null ergibt.

Bei unserem Beispiel gilt für den Knoten K_A , dass I_3 zu- und I_A, I_1 abfließen

$$(K_A) : \quad I_3 = I_A + I_1;$$

für den Knoten K_B , dass I_B zu- und I_1, I_2 abfließen

$$(K_B) : \quad I_B = I_1 + I_2.$$

Für die Masche mit angegebenen Stromrichtungen gilt, dass der Spannungsabfall über R_2 gleich der Summe der Spannungsabfällen über R_1 und R_3 ist:

$$(M) : \quad R_1 I_1 + R_3 I_3 = R_2 I_2.$$

Dies ergibt ein *System* von 3 Gleichungen für die Einzelströme I_1, I_2, I_3 . □

Im Folgenden werden wir die gegebenen Werte in die Gleichungen einsetzen und eine Methode einführen, um das System systematisch zu lösen:

	I_1	I_2	I_3	$r.S.$
$G_1 : \quad 1I_1 - 5I_2 + 3I_3 = 0$	1	-5	3	0
$G_2 : \quad -1I_1 \quad \quad + 1I_3 = 1$	-1	0	1	1
$G_3 : \quad 1I_1 + 1I_2 \quad \quad = 2$	1	1	0	2

Dieses System wird gelöst, indem die Variable I_1 aus Gleichungen G_2 und G_3 eliminiert wird. Dazu bildet man die Summe aus Gleichung G_1 und G_2 bzw. die Differenz aus Gleichung G_1 und G_3 :

	I_1	I_2	I_3	$r.S.$
$G'_1 = G_1 : \quad 1I_1 - 5I_2 + 3I_3 = 0$	1	-5	3	0
$G'_2 = G_1 + G_2 : \quad -5I_2 + 4I_3 = 1$	0	-5	4	1
$G'_3 = G_1 - G_3 : \quad -6I_2 + 3I_3 = -2$	0	-6	3	-2

Anschließend nehmen wir Gleichung G'_2 und G'_3 und eliminieren die Variable I_2 aus G'_3 . Dazu addieren wir das 6-fache von Gleichung G'_2 zum (-5) -fachen von Gleichung G'_3 :

$$\begin{array}{rcl}
 -30I_2 & + & 24I_3 = 6 \\
 30I_2 & - & 15I_3 = 10 \\
 \hline
 & & 9I_3 = 16
 \end{array}$$

Damit erhalten wir schließlich

	I_1	I_2	I_3	$r.S.$
$G''_1 = G_1 : \quad 1I_1 - 5I_2 + 3I_3 = 0$	1	-5	3	0
$G''_2 = G'_2 : \quad -5I_2 + 4I_3 = 1$	0	-5	4	1
$G''_3 = 6G'_2 - 5G'_3 : \quad 9I_3 = 16$	0	0	9	16

Aus Gleichung G''_3 folgt

$$9I_3 = 16 \Rightarrow I_3 = \frac{16}{9}.$$

Eingesetzt in Gleichung G''_2 folgt

$$-5I_2 + 4 \cdot \frac{16}{9} = 1 \Rightarrow I_2 = \frac{11}{9}.$$

Beide Ergebnisse in Gleichung G''_1 eingesetzt liefert

$$I_1 - 5 \cdot \frac{11}{9} + 3 \cdot \frac{16}{9} = 0 \Rightarrow I_1 = \frac{7}{9}.$$

Damit sind die Teilströme I_1, I_2, I_3 berechnet. □

In der letzten Spalte wurde jeweils auf die Angabe der Variablen verzichtet und nur der Koeffizient der Variablen bzw. die Konstanten auf der rechten Seite der Gleichung aufgelistet. Hierbei steht an erster Stelle immer der Koeffizient von I_1 , an zweiter Stelle der Koeffizient von I_2 und an dritter Stelle der Koeffizient von I_3 . Im Prinzip reicht diese Kurzversion des Gleichungssystems aus, um es zu lösen. Die Vorgehensweise, die wir zur Lösung dieses speziellen Gleichungssystems gewählt haben, ist verallgemeinerbar ($\rightarrow$ Gauß-Algorithmus), wenn die gesuchten Größen nur linear ($\rightarrow$ LGS) vorkommen.

► 1.5.2 Begriffsbildung und Notation

Ein *linearer Zusammenhang* zwischen zwei Größen x und y liegt dann vor, wenn x proportional zu y ($x \sim y$) ist, d.h. $ax + by = \text{const.}$ Allgemeiner bezeichnet man eine Gleichung der Form

$$ax_1 + bx_2 + cx_3 = d$$

als *lineare Gleichung* in x_1, x_2, x_3 , da jede der Variablen x_1, x_2 und x_3 nur in linearer Form, also zur Potenz 1 auftritt. Jedes 3-Tupel von reellen Zahlen $(x_1, x_2, x_3) \in \mathbb{R}^3 = \mathbb{R} \times \mathbb{R} \times \mathbb{R}$, das die Gleichung erfüllt, heißt Lösung.

Beispiele 1.20:

- ① $x_1 - x_2 + x_3 = 0$ ist eine lineare Gleichung, da die Variablen x_1, x_2, x_3 proportional in der Gleichung enthalten sind. Diese Gleichung hat z.B. $(0, 1, 1)$, $(1, 1, 0)$, $(1, 2, 1)$ als Lösungen.
- ② Die Gleichung $x^2 + 2x - y = 0$ ist **keine** lineare Gleichung, da die Variable x quadratisch vorkommt.
- ③ $x_1 - x_2 + x_3 = 0$ und $2x_1 + 3x_2 - x_3 = 0$ bilden ein System aus linearen Gleichungen, ein lineares Gleichungssystem. □

Definition: Ein System von m linearen Gleichungen in den n Unbekannten $x_1, x_2, \dots, x_n$

$$\begin{array}{ccccccc} a_{11}x_1 & + & a_{12}x_2 & + & \cdots & + & a_{1n}x_n & = & b_1 \\ a_{21}x_1 & + & a_{22}x_2 & + & \cdots & + & a_{2n}x_n & = & b_2 \\ & & \vdots & & \vdots & & \vdots & & \vdots \\ a_{m1}x_1 & + & a_{m2}x_2 & + & \cdots & + & a_{mn}x_n & = & b_m \end{array}$$

nennt man ein **lineares Gleichungssystem (LGS)**. Die reellen Zahlen a_{ij} heißen die **Koeffizienten** und b_i die **Konstanten der rechten Seite des LGS**.

Abkürzend für das LGS schreiben wir die Koeffizienten und die rechte Seite in das folgende Schema

$$\left(\begin{array}{cccc|c} a_{11} & a_{12} & a_{13} & \cdots & a_{1n} & b_1 \\ a_{21} & a_{22} & a_{23} & \cdots & a_{2n} & b_2 \\ \vdots & \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & a_{m3} & \cdots & a_{mn} & b_m \end{array} \right)$$

Man nennt dieses Schema die erweiterte Koeffizientenmatrix bzw. kurz **Matrix**. Die durchgezogene Linie soll daran erinnern, dass die Koeffizienten links und die Konstanten rechts vom Gleichheitszeichen stehen.

Ein LGS, bei dem alle Konstanten b_i der rechten Seite gleich Null sind, heißt **homogenes LGS**. Ist mindestens eine Konstante b_i ungleich Null, so heißt es ein **inhomogenes LGS**.

Jede Zeile der Matrix steht für eine Gleichung; jede Spalte ist der entsprechenden Unbekannten zugeordnet. Die Lösung besteht aus allen n -Tupeln $(x_1, x_2, \dots, x_n)$, die sämtliche m Gleichungen erfüllen.

Wie wir beim einleitenden Beispiel gesehen haben, werden beim sukzessiven Lösen des LGS nur jeweils die Koeffizienten und die Konstanten verändert, nicht aber die Variablen. Daher verzichtet man beim Lösen von LGS ganz auf die Variablen und führt alle Rechenschritte in der Matrizenschreibweise durch.

➤ 1.5.3 Das Lösen von linearen Gleichungssystemen

Umformungen, welche die Lösungsmenge eines Systems nicht ändern, nennt man **Äquivalenzumformungen**. Folgende Umformungen sind Äquivalenzumformungen eines linearen Gleichungssystems:

- (1) Die Reihenfolge der Gleichungen kann vertauscht werden.
- (2) Eine Gleichung kann mit einer reellen Zahl $\lambda \neq 0$ multipliziert werden.
- (3) Zu einer Gleichung kann eine andere Gleichung des Systems addiert werden.

Wendet man diese 3 Regeln systematisch - wie im Folgenden beschrieben wird - an, ist die Lösungsmenge jedes LGS bestimmbar. Wie im Einleitungsbeispiel gezeigt, wird in jedem Rechenschritt eine Variable aus dem System eliminiert und dadurch um eine Gleichung reduziert, bis zum Schluss nur noch eine Gleichung für eine Variable übrig bleibt. Das auf Gauß (1777-1855) zurückgehende Verfahren heißt das **Gaußsche Eliminationsverfahren** oder der **Gauß-Algorithmus**. Wir beschränken uns bei der Beschreibung der Einfachheit halber auf quadratische Systeme mit n Gleichungen für n Unbekannte. Der Gauß-Algorithmus ist aber auf beliebige $(n \times m)$ -Systeme übertragbar.

Gauß-Algorithmus

- (1) Man wählt sich eine Gleichung mit einem Koeffizienten von x_1 ungleich Null als erste Gleichung.
- (2) Man eliminiert die Variable x_1 aus den restlichen $(n - 1)$ Gleichungen. Dazu wird die 1. Zeile mit $-\frac{a_{21}}{a_{11}}$ multipliziert und zur zweiten Gleichung addiert. Ebenso verfährt man mit den übrigen Zeilen: Man addiert das $-\frac{a_{j1}}{a_{11}}$ -fache der 1. Zeile zur j -ten Zeile. Man erhält so $(n - 1)$ Gleichungen mit den $(n - 1)$ Unbekannten $x_2, x_3, \dots, x_n$.
- (3) Schritt (2) wird auf das reduzierte System angewendet, indem die Unbekannte x_2 aus Zeilen 3 bis n eliminiert wird. Nach insgesamt $(n - 1)$ Schritten bleibt nur noch eine einzige Gleichung mit der Unbekannten x_n übrig.
- (4) Die eliminierten Gleichungen bilden ein *gestaffeltes System* von Zeilen, aus denen sich die Unbekannten in der Reihenfolge $x_n, x_{n-1}, \dots, x_2, x_1$ berechnen lassen.

⚠ Im obigen Algorithmus wird angenommen, dass keiner der Koeffizienten a_{ii} gleich Null ist; ansonsten müssen die Zeilen vertauscht werden. Sind alle verbleibenden Koeffizienten von der zu eliminierenden Variablen x_i gleich Null, so kann dieser Schritt übergangen werden, da das LGS schon die gewünschte Form hat. Bei der numerischen Ausführung des Algorithmus entstehen Rechenungenauigkeiten jedoch bereits dann, wenn diese Koeffizienten sehr klein sind. Um solche Fehler möglichst klein zu halten, ist es günstig, die Zeilen in jedem Schritt so zu vertauschen, dass die Zeile mit dem betragsgrößten Koeffizienten a'_{ii} als oberste Gleichung gewählt wird. Man nennt dies *Pivotisierung*.

Beispiele 1.21:

① **Ein System mit genau einer Lösung:** Gesucht ist die Lösungsmenge des LGS

$$\begin{aligned} 2x_1 + x_2 - x_3 &= 3 \\ 3x_1 + 5x_2 - 4x_3 &= 1 \\ 4x_1 - 3x_2 + 2x_3 &= 2 \end{aligned}$$

In Matrizenschreibweise lautet dieses LGS

$$\begin{aligned} G_1 : & \quad \left(\begin{array}{ccc|c} 2 & 1 & -1 & 3 \end{array} \right) \\ G_2 : & \quad \left(\begin{array}{ccc|c} 3 & 5 & -4 & 1 \end{array} \right) \\ G_3 : & \quad \left(\begin{array}{ccc|c} 4 & -3 & 2 & 2 \end{array} \right) \end{aligned}$$

Zur Lösung wenden wir den Gauß-Algorithmus an. Dazu schreiben wir die erste Zeile ab; multiplizieren G_1 mit (-3) und addieren das Ergebnis zur 2-fachen zweiten Zeile hinzu. Außerdem multiplizieren wir die erste Zeile mit (-2) und addieren das Ergebnis zur dritten Zeile:

$$\begin{aligned} G'_1 : & \quad \left(\begin{array}{ccc|c} 2 & 1 & -1 & 3 \end{array} \right) & (G_1) \\ G'_2 : & \quad \left(\begin{array}{ccc|c} 0 & 7 & -5 & -7 \end{array} \right) & (2G_2 - 3G_1) \\ G'_3 : & \quad \left(\begin{array}{ccc|c} 0 & -5 & 4 & -4 \end{array} \right) & (G_3 - 2G_1) \end{aligned}$$

Jetzt lassen wir die beiden ersten Gleichungen unverändert und formen die letzte Gleichung so um, dass der Koeffizient von x_2 gleich Null wird.

$$\begin{aligned} G''_1 : & \quad \left(\begin{array}{ccc|c} 2 & 1 & -1 & 3 \end{array} \right) & (G'_1) \\ G''_2 : & \quad \left(\begin{array}{ccc|c} 0 & 7 & -5 & -7 \end{array} \right) & (G'_2) \\ G''_3 : & \quad \left(\begin{array}{ccc|c} 0 & 0 & 3 & -63 \end{array} \right) & (7G'_3 + 5G'_2) \end{aligned}$$

Aus dem äquivalenten System $('')$ lassen sich nun die Lösungen leicht berechnen. Die letzte Gleichung liefert

$$3x_3 = -63 \Rightarrow x_3 = -21.$$

Eingesetzt in G''_2 : $7x_2 - 5 \cdot (-21) = -7 \Rightarrow x_2 = -16.$

Beides eingesetzt in G''_1 : $2x_1 + (-16) - (-21) = 3 \Rightarrow x_1 = -1.$

Somit hat das System **genau eine Lösung** $(-1; -16; -21)$ und die Lösungsmenge lautet

$$\mathbb{L} = \left\{ (x_1, x_2, x_3) \in \mathbb{R}^3 : \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -1 \\ -16 \\ -21 \end{pmatrix} \right\}.$$

Man nennt das System $('')$ ein System mit *oberer Dreiecksmatrix*, da die Eintragungen unterhalb der *Hauptdiagonalen* (a''_{11} , a''_{22} , a''_{33}) gleich Null sind. Hat das System obere Dreiecksform ist das Eliminationsverfahren beendet. Durch *Rückwärtsauflösen* lassen sich dann die Unbekannten x_1 , x_2 , x_3 bestimmen.

② **Die Lösung enthält eine Variable:** Um das System

$$\begin{array}{rcl} x_1 - 3x_2 + 2x_3 & = & 4 \\ -2x_1 + x_2 + 3x_3 & = & 2 \\ 2x_1 - 16x_2 + 18x_3 & = & 28 \end{array}$$

zu lösen, formen wir die Koeffizientenmatrix in zwei Schritten so um, dass sie *Dreiecksform* erhält

$$\begin{array}{lcl} G_1 & \left(\begin{array}{ccc|c} 1 & -3 & 2 & 4 \end{array} \right) & \\ G_2 & \left(\begin{array}{ccc|c} -2 & 1 & 3 & 2 \end{array} \right) & \\ G_3 & \left(\begin{array}{ccc|c} 2 & -16 & 18 & 28 \end{array} \right) & \\ G'_1 & \left(\begin{array}{ccc|c} 1 & -3 & 2 & 4 \end{array} \right) & \\ G'_2 & \left(\begin{array}{ccc|c} 0 & -5 & 7 & 10 \end{array} \right) & (G_2 + 2G_1) \\ G'_3 & \left(\begin{array}{ccc|c} 0 & -15 & 21 & 30 \end{array} \right) & (G_3 + G_2) \\ G''_1 & \left(\begin{array}{ccc|c} 1 & -3 & 2 & 4 \end{array} \right) & \\ G''_2 & \left(\begin{array}{ccc|c} 0 & -5 & 7 & 10 \end{array} \right) & \\ G''_3 & \left(\begin{array}{ccc|c} 0 & 0 & 0 & 0 \end{array} \right) & (G'_3 - 3G'_2) \end{array}$$

Aus der letzten Zeile folgt $\boxed{0 \cdot x_3 = 0}$, welches für beliebiges x_3 erfüllt ist. Daher setzen wir $x_3 = \lambda$ (beliebig). In G''_2 eingesetzt, folgt

$$-5x_2 + 7\lambda = 10 \Rightarrow x_2 = -2 + \frac{7}{5}\lambda.$$

Beides in G''_1 eingesetzt, liefert

$$x_1 = 4 + 3(-2 + \frac{7}{5}\lambda) - 2\lambda = -2 + \frac{11}{5}\lambda.$$

Um eine einfachere Schreibweise zu erhalten, setzen wir $\lambda = 5k$, so dass insgesamt die Lösungsmenge lautet

$$\mathbb{L} = \left\{ (x_1, x_2, x_3) \in \mathbb{R}^3 : \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -2 \\ -2 \\ 0 \end{pmatrix} + k \begin{pmatrix} 11 \\ 7 \\ 5 \end{pmatrix} \text{ und } k \in \mathbb{R} \right\}.$$

③ **Das System hat keine Lösung:** Wir betrachten das System aus (2), indem wir die letzte Gleichung abändern: Die Konstante 28 wird durch 27 ersetzt. Durch elementare Umformungen erhält man

$$\left(\begin{array}{ccc|c} 1 & -3 & 2 & 4 \\ 0 & -5 & 7 & 10 \\ 0 & 0 & 0 & -1 \end{array} \right).$$

Aus der letzten Zeile folgt $\boxed{0 \cdot x_3 = -1}$. Diese Gleichung ist nicht erfüllbar, weil die linke Seite immer Null ergibt. Daher ist $\mathbb{L} = \{\}$.

④ **Homogenes LGS:** Nach Beispiel ② können wir sofort die Lösungsmenge des homogenen LGS

$$\begin{aligned}x_1 - 3x_2 + 2x_3 &= 0 \\ -2x_1 + x_2 + 3x_3 &= 0 \\ 2x_1 - 16x_2 + 18x_3 &= 0\end{aligned}$$

angeben, denn die elementaren Zeilenumformungen liefern

$$\left(\begin{array}{ccc|c} 1 & -3 & 2 & 0 \\ 0 & -5 & 7 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

Durch Rückwärtsauflösen erhalten wir aus Zeile 3 $\boxed{0 \cdot x_3 = 0}$. Daher ist x_3 beliebig. Wir setzen $x_3 = 5k$. In Zeile 2 eingesetzt, folgt

$$-5x_2 + 7 \cdot 5k = 0 \Rightarrow x_2 = 7k$$

und beides in Zeile 1 eingesetzt:

$$x_1 = +3 \cdot 7k - 2 \cdot 5k = 11k.$$

Daher ist

$$\mathbb{L} = \left\{ (x_1, x_2, x_3) : \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = k \begin{pmatrix} 11 \\ 7 \\ 5 \end{pmatrix} \text{ und } k \in \mathbb{R} \right\}. \quad \square$$

Die Beispiele ① - ④ legen folgende allgemeingültige Schlussfolgerung nahe, die wir im Kapitel über Matrizen und Determinanten genauer untersuchen:

Lösungsverhalten von linearen Gleichungssystemen

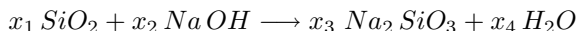
- (1) Ein **inhomogenes** LGS besitzt entweder genau eine Lösung oder unendlich viele Lösungen oder überhaupt keine Lösung.
- (2) Ein **homogenes** LGS besitzt entweder genau eine Lösung, nämlich die triviale Null-Lösung $x = \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix}$, oder unendlich viele Lösungen.
- (3) Falls das inhomogene LGS lösbar ist, setzt sich die Lösung zusammen aus allen homogenen Lösungen plus einer Lösung des inhomogenen Systems:

$$\mathbb{L}_i = \mathbb{L}_h + x_s,$$

wenn $\mathbb{L}_i$ = Lösungsmenge des inhomogenen LGS, $\mathbb{L}_h$ = Lösungsmenge des zugehörigen homogenen LGS und x_s eine spezielle Lösung des inhomogenen Systems ist.

Anwendungsbeispiel 1.22 (Chemische Reaktion).

Aus Quarz (SiO_2) und Natronlauge ($NaOH$) entsteht Natriumsilikat (Na_2SiO_3) und Wasser (H_2O):



Gesucht sind die Anteile der Stoffe x_1, x_2, x_3, x_4 , für welche die Reaktion abläuft. Da nur ganzzahlige Vielfache in Frage kommen, sind natürliche Zahlen x_1, x_2, x_3, x_4 zu bestimmen, so dass jedes der chemischen Elemente Si, O, Na, H auf beiden Seiten der Reaktionsgleichung gleich oft auftritt. Dies führt zu dem folgenden homogenen linearen Gleichungssystem:

$$Si: \quad x_1 = x_3$$

$$Na: \quad x_2 = 2x_3$$

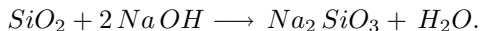
$$O: \quad 2x_1 + x_2 = 3x_3 + x_4$$

$$H: \quad x_2 = 2x_4.$$

In Matrizenform lautet das LGS

$$\left(\begin{array}{cccc|c} 1 & 0 & -1 & 0 & 0 \\ 0 & 1 & -2 & 0 & 0 \\ 2 & 1 & -3 & -1 & 0 \\ 0 & 1 & 0 & -2 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} 1 & 0 & -1 & 0 & 0 \\ 0 & 1 & -2 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right)$$

Aus der letzten Zeile folgt $0 \cdot x_4 = 0$. Daher ist x_4 beliebig. Wir wählen $x_4 = k$. In Zeile 3 eingesetzt, folgt $x_3 = k$. Beide Ergebnisse in Zeile 2 bzw. Zeile 1 eingesetzt liefert $x_2 = 2k$ und $x_1 = k$. Die Lösung mit den kleinsten Anteilen der Substanzen lautet daher (für $k = 1$)



□

Anwendungsbeispiel 1.23 (Mischen von Legierungen).

Edelstahl ist eine Legierung aus Eisen, Chrom und Nickel. Beispielsweise besteht V2A-Stahl aus 74% Eisen, 18% Chrom und 8% Nickel. In der unten stehenden Tabelle sind vorhandene Legierungen (I - IV) angegeben, mit denen 1000 kg V2A-Stahl gemischt werden soll.

	I	II	III	IV
Eisen	70%	72%	80%	85%
Chrom	22%	20%	10%	12%
Nickel	8%	8%	10%	3%

Sind x_1, x_2, x_3, x_4 die Anteile der Legierungen I - IV in Einheiten kg, so gilt für die Summe aller Mischungsanteile in kg

$$x_1 + x_2 + x_3 + x_4 = 1000.$$

Für die Einzelbestandteile Eisen, Chrom und Nickel gelten die Erhaltungsgleichungen

$$0.7 x_1 + 0.72 x_2 + 0.8 x_3 + 0.85 x_4 = 740$$

$$0.22 x_1 + 0.2 x_2 + 0.1 x_3 + 0.12 x_4 = 180$$

$$0.08 x_1 + 0.08 x_2 + 0.1 x_3 + 0.03 x_4 = 80.$$

Man beachte, dass bei 1000 kg Legierung mit 74% Eisen das Eisengewicht 740 kg beträgt, entsprechendes gilt für Chrom und Nickel. Diese vier Gleichungen liefern ein inhomogenes lineares Gleichungssystem

$$\left(\begin{array}{cccc|c} 1 & 1 & 1 & 1 & 1000 \\ 70 & 72 & 80 & 85 & 74000 \\ 22 & 20 & 10 & 12 & 18000 \\ 8 & 8 & 10 & 3 & 8000 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} 1 & 1 & 1 & 1 & 1000 \\ 0 & 2 & 10 & 15 & 4000 \\ 0 & 0 & -2 & 5 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right).$$

Wir wählen $x_4 = k$ beliebig. In Zeile 3 eingesetzt, folgt $x_3 = \frac{5}{2}k$. Beide Ergebnisse in Zeile 2 bzw. Zeile 1 eingesetzt liefert $x_2 = 2000 - 20k$ und $x_1 = -1000 + 16.5k$. Damit die Lösung realisierbar ist, müssen alle Anteile $x_1, x_2, x_3, x_4 \geq 0$ sein. Wegen der Bedingung für x_2 folgt $100 \geq k$ und wegen der Bedingung für x_1 folgt $k \geq 60.6$. D.h. für $100 \geq k \geq 60.6$ ist das Problem physikalisch lösbar. $\square$

MAPLE-Worksheets zu Kapitel 1



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 1 mit MAPLE zur Verfügung. Sie können direkt durch Anklicken aus der pdf-Version des Buches gestartet werden.

- Mengen mit MAPLE
- Natürliche Zahlen mit MAPLE
- Pythagoräische Zahlen mit MAPLE
- Reelle Zahlen mit MAPLE
- Gleichungen und Ungleichungen mit MAPLE
- Visualisierung von Gleichungen mit MAPLE
- Lineare Gleichungssysteme mit MAPLE
- Zusammenstellung der MAPLE-Befehle
- MAPLE-Lösungen zu den Aufgaben

1.6 Aufgaben zu Zahlen, Gleichungen und Gleichungssystemen

- 1.1 Stellen Sie die folgenden Mengen durch Aufzählen ihrer Elemente dar:
 a) $\{x : x \text{ ist Primzahl und } x < 20\}$
 b) $\{x : x \text{ ist reell und } x^2 + 1 = 0\}$
- 1.2 Gegeben sind die Mengen $A = \{x \in \mathbb{R} : 0 < x < 2\}$ und $B = \{x \in \mathbb{R} : 1 \leq x \leq 3\}$. Bestimmen Sie graphisch sowie rechnerisch
 (i) $A \cap B$, (ii) $A \cup B$, (iii) $A \times B$, (iv) $A \setminus B$.
- 1.3 Bilden Sie die Vereinigung, Durchschnitt und beide Differenzmengen aus den folgenden Mengen
 a) $M_1 = \{2, 4, 6, \dots\}$, $M_2 = \{3, 6, 9, \dots\}$
 b) $M_1 = \{x : x^2 + x - 2 = 0\}$, $M_2 = \{x : x^2 - 3x + 2 = 0\}$
- 1.4 Zeigen Sie mit Hilfe von Venn-Diagrammen, dass für drei Mengen M_1 , M_2 und M_3 gilt
 a) $M_1 \cap (M_2 \cup M_3) = (M_1 \cap M_2) \cup (M_1 \cap M_3)$
 b) $M_1 \cup (M_2 \cap M_3) = (M_1 \cup M_2) \cap (M_1 \cup M_3)$
- 1.5 Zeigen Sie durch vollständige Induktion, dass für alle $n \in \mathbb{N}$ gilt
 a) $1^2 + 2^2 + 3^2 + \dots + n^2 = \sum_{k=1}^n k^2 = \frac{n(n+1)(2n+1)}{6}$
 b) $2^0 + 2^1 + 2^2 + \dots + 2^n = \sum_{k=0}^n 2^k = 2^{n+1} - 1$
 c) $\frac{1}{1 \cdot 2} + \frac{1}{2 \cdot 3} + \frac{1}{3 \cdot 4} + \dots + \frac{1}{n(n+1)} = \frac{n}{n+1}$
- 1.6 Zeigen Sie durch vollständige Induktion
 a) $2^n \leq n!$ für jedes $n \geq 4$
 b) $2n + 1 \leq 2^n$ für jedes $n \geq 3$
 c) $n^2 \leq 2^n$ für jedes $n \neq 3$
- 1.7 Berechnen Sie
 a) $\binom{n}{0}, \binom{n}{n}, \binom{3}{1}, \binom{3}{2}, \binom{4}{0}, \binom{4}{1}, \binom{4}{2}, \binom{4}{3}, \binom{4}{4}$
 b) 102^4
- 1.8 Rechnen Sie nach, dass

$$\sum_{k=1}^n a_{k-1} = \sum_{k=0}^{n-1} a_k \quad ; \quad \sum_{k=0}^{n-1} a_{k+1} = \sum_{k=1}^n a_k .$$
- 1.9 Man zeige $\binom{n}{k} \frac{1}{n^k} \leq \frac{1}{k!}$ für jedes $n \in \mathbb{N}$. (Nachrechnen!)
- 1.10 Man entwickle die folgenden Binome
 a) $(x+4)^5$ b) $(1-5y)^4$ c) $(a^2-2b)^3$.
- 1.11 Bestimmen Sie mit MAPLE den Summenwert von $\sum_{k=71}^{125} (k^2 + 1)$.

- 1.12 Bestimmen Sie mit MAPLE eine Formel für die Summenwerte von
- $1^3 + 2^3 + 3^3 + \dots + n^3$
 - $1^4 + 2^4 + 3^4 + \dots + n^4$
 - $\sum_{i=1}^n \frac{1}{i(i+1)}$; $\sum_{i=1}^n \frac{1}{i(i+1)(i+2)}$; $\sum_{i=0}^n x^i$.
- 1.13 Vereinfachen Sie die folgenden Ausdrücke soweit möglich
- $\frac{18x^{a+4}}{2y^{5a+7}} : \frac{4x^{7-3a}}{9y^{8+5a}}$
 - $(a^{n+1}b^{x-1} + a^n b^x + a^{n-1}b^{x+1}) : (a^{n-2}b^{x-1})$
- 1.14 Vereinfachen Sie formal die Wurzelterme soweit möglich
- $\frac{x(2r^2 - 4x^2)}{\sqrt{r^2 - x^2}} - 8x\sqrt{r^2 - x^2}$
 - $2\sqrt{(x-k)^2 + x^2} - \frac{(2x-k)^2}{\sqrt{2x^2 - 2kx + k^2}}$
 - $\sqrt{6x^2 - 6}\sqrt{\frac{3x-3}{2x+2}}$
- 1.15 Berechnen Sie
- $\sqrt[3]{a^6 b^{12}}$
 - $\sqrt[4]{a^2 \sqrt[3]{a^2}}$
 - $\sqrt[3]{\sqrt{a^6 b^8}}$
 - $\sqrt[3]{a^3 \sqrt{a^2 \sqrt[5]{a^8 \sqrt[4]{a^3}}}}$
 - $\frac{\sqrt[6]{a^5 \sqrt[3]{a^2}}}{\sqrt[3]{a^2 \sqrt[6]{a^4}}} : \frac{\sqrt{a^3 \sqrt[9]{a^7}}}{\sqrt[9]{a^7 \sqrt{a}}}$
- 1.16 Berechnen Sie
- $\text{ld } 2^4, \log \sqrt{10}, \ln e^3$
 - $\ln(\sqrt{e})^3, \ln \sqrt{\frac{1}{\sqrt[3]{e^2}}}, \ln \sqrt{e^{3(\ln e^2 + \ln e^6)}}$
 - $\log \sqrt[n+1]{a^n \sqrt[m]{b^{-1}}}$
- 1.17 Zeigen Sie, dass die beiden Mengen zusammen mit den Rechenoperationen $+$ und $\cdot$ die Körperaxiome erfüllen.
- $(\{a + b\sqrt{2} \text{ mit } a, b \in \mathbb{Q}\}, +, \cdot)$ mit den Rechenoperationen in $\mathbb{R}$.
 - $(F_2, +, \cdot)$ bzgl. den in Beispiel 1.9 ④ angegebenen Verknüpfungstabellen.
- 1.18 Geben Sie die reellen Lösungen der folgenden quadratischen Gleichungen an:
- $4x^2 + 8x - 60 = 0$
 - $x^2 - 4x + 13 = 0$
 - $-1 = -9(x-2)^2$
 - $5x^2 + 20x + 20 = 0$
 - $(x-1)(x+3) = -4$
- 1.19 Man bestimme den Parameter c so, dass die Gleichung $2x^2 + 4x = c$ genau eine reelle Lösung besitzt.
- 1.20 Welche reellen Lösungen besitzen die Gleichungen?
- $-2x^3 + 8x^2 = 8x$
 - $t^4 - 13t^2 + 36 = 0$
 - $\frac{1}{2}(3x^2 - 6)(x^2 - 25)(x + 3) = 0$
- 1.21 Lösen Sie mit MAPLE die folgenden Wurzelgleichungen:
- $\sqrt{-3+2x} = 2$
 - $\sqrt{x^2+4} = x-2$
 - $\sqrt{x-1} = \sqrt{x+1}$
 - $\sqrt{2x^2-1} + x = 0$
- 1.22 Welche reellen Lösungen besitzen die Betragsgleichungen?
- $|x^2 - x| = 24$
 - $|2x + 4| = -(x^2 - x - 6)$

1.23 Bestimmen Sie mit MAPLE die reellen Lösungsmengen der folgenden Ungleichungen:

a) $2x - 8 > |x|$ b) $x^2 + x + 1 \geq 0$ c) $|x| \leq x - 2$ d) $|x - 4| > x^2$

1.24 Lösen Sie die folgenden Gleichungssysteme:

a)
$$\begin{array}{rrrrcl} 4x_1 & + & 2x_2 & + & 4x_3 & = & 10 \\ x_1 & + & x_2 & + & x_3 & = & 3 \\ 2x_1 & + & 3x_2 & + & 3x_3 & = & 8 \end{array}$$

b)
$$\begin{array}{rrrrcl} 2x_1 & + & x_2 & + & x_3 & = & 7 \\ 2x_1 & + & 2x_2 & + & x_3 & = & 10 \\ 3x_1 & & & + & x_3 & = & 5 \end{array}$$

c)
$$\begin{array}{rrrrcl} 2x_1 & + & x_2 & + & x_3 & = & 7 \\ 2x_1 & + & x_2 & + & x_3 & = & 0 \\ 3x_1 & & & + & x_3 & = & 5 \end{array}$$

1.25 Man bestimme die Lösungsmenge der folgenden Systeme:

a)
$$\begin{array}{rrrrcl} x_1 & - & 3x_2 & + & x_3 & = & -3 \\ -3x_1 & + & x_2 & + & x_3 & = & 5 \end{array}$$

b)
$$\begin{array}{rrrrcl} x_1 & + & x_2 & + & x_3 & = & 6 \\ x_1 & + & 2x_2 & + & x_3 & = & 7 \\ 2x_1 & + & x_2 & + & 2x_3 & = & 11 \end{array}$$

c)
$$\begin{array}{rrrrcl} x_1 & + & x_2 & + & x_3 & = & 7 \\ x_1 & + & 2x_2 & + & x_3 & = & 7 \\ 2x_1 & + & x_2 & + & 2x_3 & = & 11 \end{array}$$

1.26 Bestimmen Sie die Lösungsmenge der linearen Gleichungssysteme:

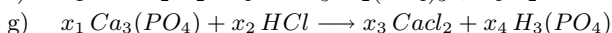
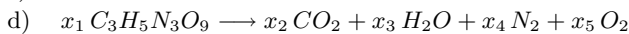
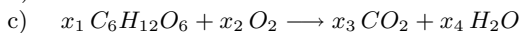
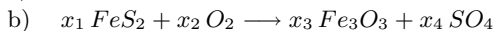
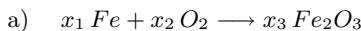
a)
$$2x_1 + 3x_2 + 4x_3 = 4$$

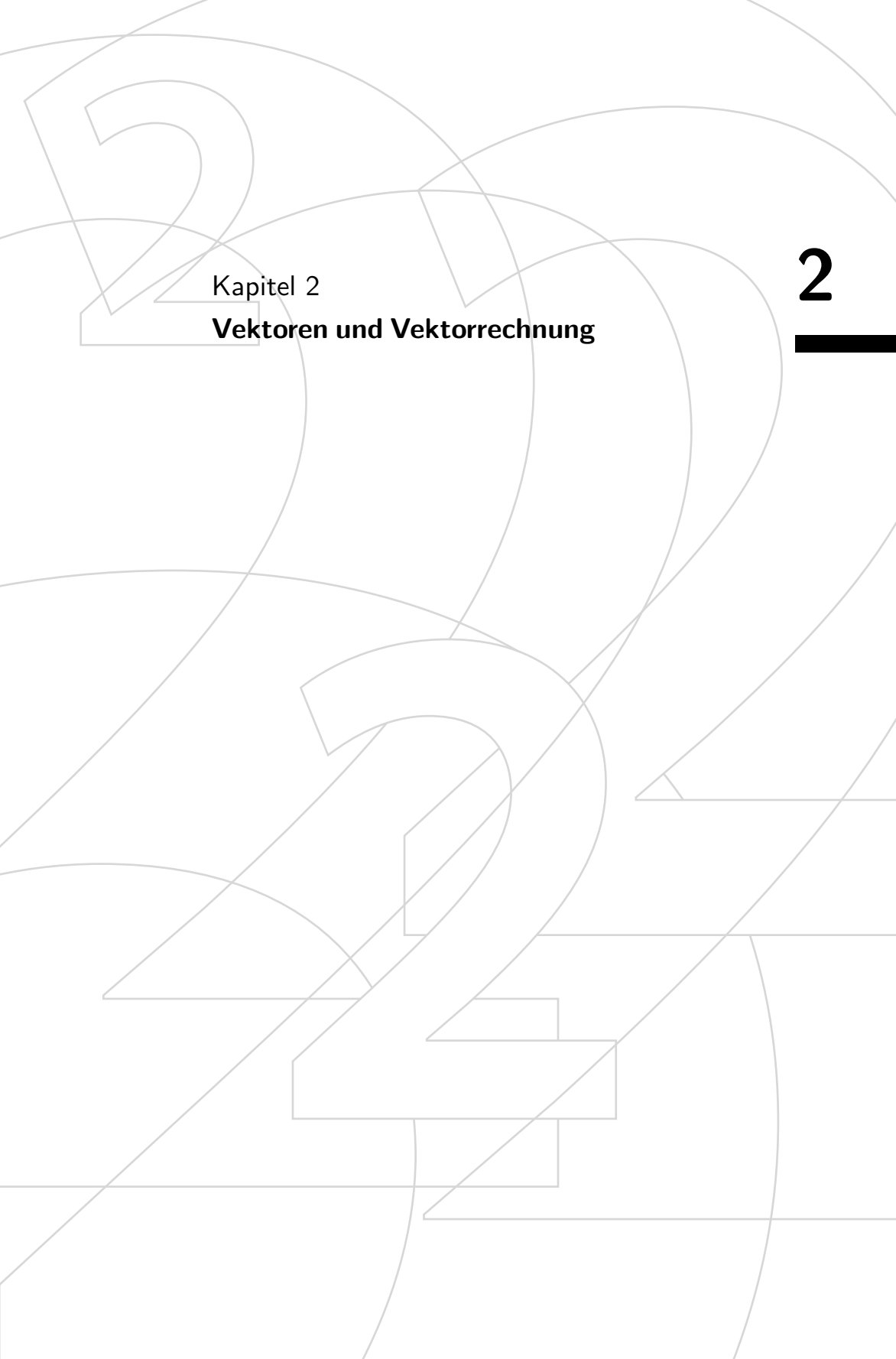
b)
$$\begin{array}{rrrrcl} x_1 & - & x_2 & + & x_3 & = & 1 \\ -3x_1 & + & 3x_2 & - & 3x_3 & = & -3 \\ 5x_1 & - & 5x_2 & + & 5x_3 & = & 5 \end{array}$$

c)
$$\begin{array}{rrrrcl} x_1 & - & x_2 & + & x_3 & = & 1 \\ -3x_1 & + & 3x_2 & - & 3x_3 & = & -1 \\ 5x_1 & - & 5x_2 & + & 5x_3 & = & 5 \end{array}$$

1.27 Welche Aussagen gelten für die entsprechenden homogenen Systeme?

1.28 Die Variablen $x_1, x_2, \dots$ in den folgenden chemischen Reaktionen sollen für möglichst kleine natürliche Zahlen stehen:





Kapitel 2

Vektoren und Vektorrechnung

2

2	Vektoren und Vektorrechnung	39
2.1	Vektoren im $\mathbb{R}^2$	42
2.1.1	Multiplikation eines Vektors mit einem Skalar	42
2.1.2	Addition zweier Vektoren	43
2.1.3	Die Länge (der Betrag) eines Vektors	43
2.1.4	Das Skalarprodukt zweier Vektoren	44
2.1.5	Geometrische Anwendung	48
2.2	Vektoren im $\mathbb{R}^3$	50
2.2.1	Rechenregeln für Vektoren	50
2.2.2	Projektion eines Vektors	53
2.2.3	Das Vektorprodukt (Kreuzprodukt) zweier Vektoren	55
2.2.4	Das Spatprodukt von drei Vektoren	59
2.3	Geraden und Ebenen im $\mathbb{R}^3$	61
2.3.1	Vektorielle Darstellung von Geraden	61
2.3.2	Lage zweier Geraden zueinander	62
2.3.3	Abstandsberechnung zu Geraden	64
2.3.4	Vektorielle Darstellung von Ebenen	67
2.3.5	Lage zweier Ebenen zueinander	70
2.3.6	Abstandsberechnung zu Ebenen	72
2.3.7	Berechnung des Schnittes einer Geraden mit einer Ebene	73
2.4	Vektorräume	76
2.4.1	Vektorrechnung im $\mathbb{R}^n$	76
2.4.2	Vektorräume	78
2.4.3	Linearkombination und Erzeugnis	81
2.4.4	Lineare Abhängigkeit und Unabhängigkeit	84
2.4.5	Basis und Dimension	87
2.5	Aufgaben zur Vektorrechnung	92
Zusätzliche Abschnitte auf der CD-Rom		
2.6	MAPLE: Vektorrechnung	cd
2.7	Punkte, Geraden und Ebenen mit MAPLE	cd

2 Vektoren und Vektorrechnung

Vektoren sind ein unentbehrliches Hilfsmittel bei der Beschreibung physikalischer Größen. Während die Temperatur eines Körpers, die Dichte eines homogenen Mediums, der Ohmsche Widerstand eines elektrischen Elementes durch eine reelle Zahl (zusammen mit einer Einheit) charakterisiert werden, ist dies z.B. bei den folgenden physikalischen Größen nicht möglich:

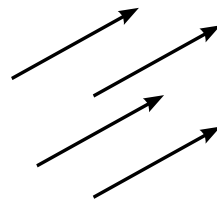
Die Geschwindigkeit eines Massenpunktes im Raum ist festgelegt durch die Größe der Geschwindigkeit und deren Richtung. Die Kraft, die an einem Massenpunkt angreift, wird beschrieben durch die Länge der Kraft und der Richtung, unter welcher die Kraft angreift. Elektrische Feldstärke, Drehmoment usw. sind andere physikalische Größen, die durch Maßzahl (Länge) und Richtung festgelegt werden.

Immer wenn in Physik und Technik Größen auftreten, die sich nicht nur durch die Angabe einer Zahl zusammen mit einer Einheit versehen beschreiben lassen, spricht man von Vektoren.

Stellen wir uns einen Kraftvektor vor, so definieren wir verallgemeinernd:

Definition: Ein Vektor $\vec{a}$ ist eine Klasse von gerichteten Strecken (Pfeilen), die in Richtung und Länge übereinstimmen.

Zwei gerichtete Strecken $\overrightarrow{AB}$ (Anfangspunkt A , Endpunkt B) und $\overrightarrow{CD}$ (Anfangspunkt C , Endpunkt D) stellen genau dann denselben Vektor dar, wenn sie gleich gerichtet und gleichlang sind. Man spricht daher oftmals von *Richtungsvektoren*. Durch Parallelverschiebung entstehende Vektoren sind also gleich. Ein Vektor ist eindeutig durch seinen Anfangspunkt und Endpunkt festgelegt.



Historisch gesehen ist die Vektorrechnung eine recht junge Disziplin verglichen z.B. mit der Differenzialrechnung. Die Begründung der modernen Vektorrechnung geht auf Hermann Großmann (1809 - 1877; 1844) zurück. Der Formalismus der Vektoren und der Vektorrechnung entstand also wesentlich später als z.B. die komplexen Zahlen.

Im Folgenden werden wir zur Beschreibung der Vektoren und der Rechenoperationen mit Vektoren von einem rechtwinkligen (*kartesischen*) *Koordinatensystem* ausgehen.

2.1 Vektoren im $\mathbb{R}^2$

Der *zweidimensionale Raum* $\mathbb{R}^2$ wird durch zwei senkrecht aufeinander stehenden Koordinatenachsen (**kartesisches Koordinatensystem**) festgelegt. In einem solchen Koordinatensystem ist ein Vektor $\vec{a}$ vom Punkt $P_1 = (x_1, y_1)$ zum Punkt $P_2 = (x_2, y_2)$ durch seine **Komponenten** festgelegt, den Projektionen auf die Koordinatenachsen:

$$\vec{a} := \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \end{pmatrix}.$$

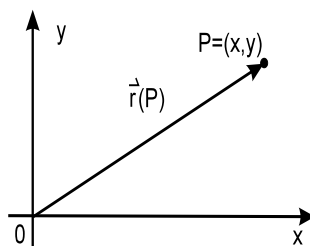
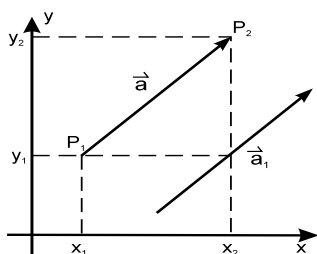


Abb. 2.1. Richtungsvektoren (links)- und Ortsvektor (rechts)

Dabei kommt es nicht auf die spezielle Lage im $\mathbb{R}^2$ an. $\vec{a}$ und $\vec{a}_1$ repräsentieren den gleichen Vektor. Wir sprechen von einem **Richtungsvektor**, wenn der Angriffspunkt keine Rolle spielt. Im Gegensatz zu Richtungsvektoren spricht man von **Ortsvektoren**, wenn der Vektor vom Ursprung O zum Punkt P führt:

$$\vec{r}(P) := \begin{pmatrix} x \\ y \end{pmatrix}.$$

➤ 2.1.1 Multiplikation eines Vektors mit einem Skalar

$$\lambda \vec{a} = \lambda \begin{pmatrix} a_x \\ a_y \end{pmatrix} := \begin{pmatrix} \lambda a_x \\ \lambda a_y \end{pmatrix}; \quad \lambda \in \mathbb{R}.$$

Das Produkt eines Skalars $\lambda \in \mathbb{R}$ mit einem Vektor $\vec{a}$ ist wieder ein Vektor. Die Multiplikation geschieht **komponentenweise**. Geometrisch entspricht dies der *Streckung* des Vektors $\vec{a}$ um den Faktor λ . Für $\lambda = -1$ hat $-\vec{a}$ die gleiche Länge aber umgekehrte Richtung wie $\vec{a}$. Manchmal ist es bequemer den Zahlenfaktor rechts vom Vektor zu schreiben. Wir setzen daher $\vec{a}\lambda := \lambda\vec{a}$.

2.1.2 Addition zweier Vektoren

$$\vec{a} + \vec{b} = \begin{pmatrix} a_x \\ a_y \end{pmatrix} + \begin{pmatrix} b_x \\ b_y \end{pmatrix} := \begin{pmatrix} a_x + b_x \\ a_y + b_y \end{pmatrix}.$$

Die Summe zweier Vektoren ist ein Vektor. Die Addition erfolgt komponentenweise. Entsprechend ist die Subtraktion erklärt. Geometrisch entspricht die Addition zweier Vektoren der *Kräfte-Addition* über das *Kräfte-Parallelogramm*.

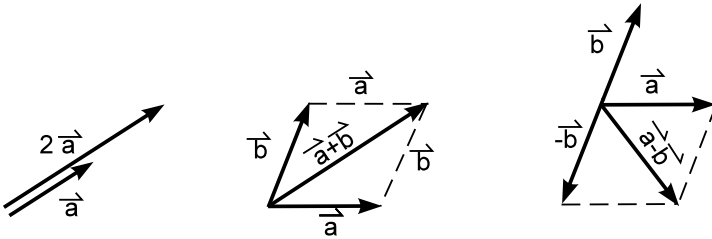


Abb. 2.2. Addition und Subtraktion über das Kräfte-Parallelogramm

2.1.3 Die Länge (der Betrag) eines Vektors

Die *Länge* (= *Betrag*) eines Vektors $\vec{a}$ ergibt sich nach dem Satz von Pythagoras:

$$|\vec{a}| = \sqrt{a_x^2 + a_y^2}.$$

Für den Betrag eines Vektors schreibt man auch kurz $a := |\vec{a}|$.

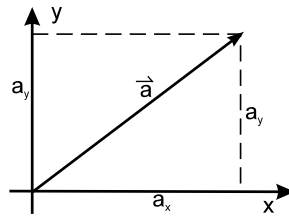


Abb. 2.3. Betrag des Vektors

Beispiele 2.1:

- ① Gegeben sind die Vektoren $\vec{a} = \begin{pmatrix} 5 \\ 3 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 4 \\ 1 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} -3 \\ -2 \end{pmatrix}$. Dann ist

$$\vec{d} = -\vec{a} + 3\vec{b} + 2\vec{c} = \begin{pmatrix} -5 \\ -3 \end{pmatrix} + \begin{pmatrix} 12 \\ 3 \end{pmatrix} + \begin{pmatrix} -6 \\ -4 \end{pmatrix} = \begin{pmatrix} 1 \\ -4 \end{pmatrix}.$$

- ② An einem Massenpunkt wirken die Kräfte

$$\vec{F}_1 = \begin{pmatrix} 2N \\ 1N \end{pmatrix}, \vec{F}_2 = \begin{pmatrix} -1N \\ 5N \end{pmatrix}, \vec{F}_3 = \begin{pmatrix} -4N \\ 2N \end{pmatrix}.$$

Gesucht ist der Betrag der resultierenden Kraft F_R :

$$\vec{F}_R = \vec{F}_1 + \vec{F}_2 + \vec{F}_3$$

$$= \begin{pmatrix} 2N \\ 1N \end{pmatrix} + \begin{pmatrix} -1N \\ 5N \end{pmatrix} + \begin{pmatrix} -4N \\ 2N \end{pmatrix} = \begin{pmatrix} -3N \\ 8N \end{pmatrix}$$

$$F_R = |\vec{F}_R| = \sqrt{9N^2 + 64N^2} = \sqrt{73}N.$$

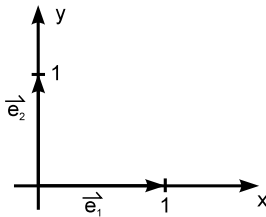


Abb. 2.4. Einheitsvektoren

③ Ein Vektor der Länge 1 heißt **Einheitsvektor**. Spezielle Einheitsvektoren sind die *Koordinaten-Einheitsvektoren*

$$\vec{e}_1 := \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad \text{und} \quad \vec{e}_2 := \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

Diese Vektoren haben die Richtung der entsprechenden Koordinatenachsen und die Länge 1. Mit den Einheitsvektoren lässt sich jeder Vektor $\vec{a}$ schreiben als

$$\vec{a} = a_x \vec{e}_1 + a_y \vec{e}_2$$

(Linearkombination von $\vec{e}_1$ und $\vec{e}_2$).

- ④ Gegeben ist ein Vektor $\vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix}$. Gesucht ist der **Einheitsvektor in Richtung** $\vec{a}$. Wegen $|\vec{a}| = \sqrt{a_x^2 + a_y^2}$ ist

$$\vec{e}_a = \frac{1}{|\vec{a}|} \vec{a} = \frac{1}{\sqrt{a_x^2 + a_y^2}} \begin{pmatrix} a_x \\ a_y \end{pmatrix}$$

der gesuchte Einheitsvektor, denn $|\vec{e}_a| = 1$ und die Richtung von $\vec{e}_a$ und $\vec{a}$ stimmen überein. $\square$

➤ 2.1.4 Das Skalarprodukt zweier Vektoren

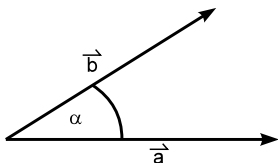


Abb. 2.5. Winkel zw. $\vec{a}$ und $\vec{b}$

Definition: Unter dem **Skalarprodukt** (Punktprodukt) zweier Vektoren $\vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} b_x \\ b_y \end{pmatrix}$ versteht man die reelle Zahl

$$\vec{a} \cdot \vec{b} := |\vec{a}| \cdot |\vec{b}| \cdot \cos \alpha, \quad (1)$$

wenn α der zwischen 0° und 360° gelegene Winkel zwischen den Vektoren $\vec{a}$ und $\vec{b}$ ist. Übliche Bezeichnungen für das Skalarprodukt sind auch $\langle \vec{a}, \vec{b} \rangle$ oder $(\vec{a}, \vec{b})$.

Für das Skalarprodukt gelten aufgrund der Definition die **Rechenregeln**

(S ₁)	$\vec{a} \cdot \vec{b} = \vec{b} \cdot \vec{a}$	Symmetriegesetz
(S ₂)	$\lambda \cdot (\vec{a} \cdot \vec{b}) = (\lambda \cdot \vec{a}) \cdot \vec{b} = \vec{a} \cdot (\lambda \cdot \vec{b})$	Assoziativgesetz
(S ₃)	$\vec{a} \cdot (\vec{b} + \vec{c}) = (\vec{a} \cdot \vec{b}) + (\vec{a} \cdot \vec{c})$	Distributivgesetz

(S₁) und (S₂) sind offensichtlich, (S₃) ist nicht ganz evident. Auf Beweise sei jedoch verzichtet. Wir verwenden stattdessen die Regeln, um eine sehr einfache Darstellung des Skalarprodukts zu erhalten: Aufgrund der Definition des Skalarprodukts ist

$$\vec{e}_1 \cdot \vec{e}_1 = 1, \quad \vec{e}_1 \cdot \vec{e}_2 = 0, \quad \vec{e}_2 \cdot \vec{e}_2 = 1.$$

Daher gilt für zwei Vektoren

$$\vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix} = a_x \vec{e}_1 + a_y \vec{e}_2 \quad \text{und} \quad \vec{b} = \begin{pmatrix} b_x \\ b_y \end{pmatrix} = b_x \vec{e}_1 + b_y \vec{e}_2 :$$

$$\begin{aligned} \vec{a} \cdot \vec{b} &= (a_x \vec{e}_1 + a_y \vec{e}_2) \cdot (b_x \vec{e}_1 + b_y \vec{e}_2) \\ &= a_x b_x \vec{e}_1 \cdot \vec{e}_1 + a_x b_y \vec{e}_1 \cdot \vec{e}_2 + a_y b_x \vec{e}_2 \cdot \vec{e}_1 + a_y b_y \vec{e}_2 \cdot \vec{e}_2 \end{aligned}$$

$$\Rightarrow \quad \vec{a} \cdot \vec{b} = a_x b_x + a_y b_y. \quad (2)$$

Das Skalarprodukt zweier Vektoren lässt sich also einfach angeben ohne den Winkel α zwischen den Vektoren $\vec{a}$ und $\vec{b}$ berechnen zu müssen, indem **die Summe der Produkte der ersten Komponenten und der zweiten Komponenten gebildet wird.**

Anwendungsbeispiel 2.2 (Arbeit bei konstanter Kraft). Physikalisch entspricht das Skalarprodukt der Arbeit, die unter Einwirkung einer konstanten Kraft geleistet wird. Sind Kraft und Bewegungsrichtung gleichgerichtet, dann ist die Arbeit $W = F \cdot s$. Kann sich ein Massenpunkt aber nur entlang einer Richtung $\vec{e}_a$ bewegen, die nicht mit der Richtung der Kraft übereinstimmt, dann ist die geleistete Arbeit die Komponente der Kraft $|\vec{F}_a|$ in Richtung $\vec{e}_a$ multipliziert mit dem zurückgelegten Weg $|\vec{a}| = |a \vec{e}_a|$:

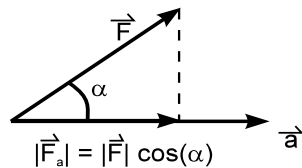


Abb. 2.6. Kraft in Richtung $\vec{a}$

$$W = |\vec{F}_a| \cdot |\vec{a}| = |\vec{F}| \cdot \cos \alpha \cdot |\vec{a}|$$

$$W = \vec{F} \cdot \vec{a}.$$

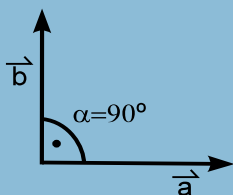
□

Anwendungsbeispiel 2.3 (Berechnung des von zwei Vektoren eingeschlossenen Winkels).

Aus Gleichung (1) und (2) folgt eine wichtige Formel zur Berechnung des von zwei Vektoren eingeschlossenen Winkels:

$$\cos \alpha = \frac{\vec{a} \cdot \vec{b}}{|\vec{a}| \cdot |\vec{b}|} = \frac{a_x b_x + a_y b_y}{\sqrt{a_x^2 + a_y^2} \sqrt{b_x^2 + b_y^2}}.$$

Wendet man anschließend die Umkehrfunktion des Kosinus an (siehe 4.7 Arkusfunktionen), erhält man den von den Vektoren eingeschlossenen Winkel α zwischen 0 und 180° . $\square$



Wichtiger Spezialfall: Stehen $\vec{a}$ und $\vec{b}$ senkrecht aufeinander, so ist $\alpha = 90^\circ$ und $\cos \alpha = 0$. Daher gilt

$$\vec{a} \cdot \vec{b} = 0 \Leftrightarrow \vec{a} \perp \vec{b}.$$

⚠ Achtung: Im Gegensatz zum Produkt von zwei reellen Zahlen ist das Skalarprodukt nicht nur dann Null, wenn mindestens einer der beiden Faktoren der Nullvektor ist, sondern auch, wenn die beiden Vektoren aufeinander senkrecht stehen.

Beispiele 2.4:

① Man bestimme das Skalarprodukt von $\vec{a} = \begin{pmatrix} 4 \\ 2 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} -1 \\ 3 \end{pmatrix}$.

$$\vec{a} \cdot \vec{b} = \begin{pmatrix} 4 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 3 \end{pmatrix} = 4 \cdot (-1) + 2 \cdot 3 = 2.$$

② Die Vektoren $\vec{a} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} -2 \\ 1 \end{pmatrix}$ sind *orthogonal*, d.h. sie stehen senkrecht aufeinander.

Wir zeigen, dass das Skalarprodukt Null ergibt:

$$\vec{a} \cdot \vec{b} = \begin{pmatrix} 1 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} -2 \\ 1 \end{pmatrix} = 1 \cdot (-2) + 2 \cdot 1 = 0.$$

- ③ Der Betrag eines Vektors kann aus dem Skalarprodukt berechnet werden:

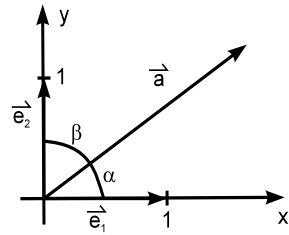
Für $\vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix}$ gilt $\vec{a} \cdot \vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix} \cdot \begin{pmatrix} a_x \\ a_y \end{pmatrix} = a_x^2 + a_y^2 = |\vec{a}|^2$

$$\Rightarrow |\vec{a}| = a = \sqrt{\vec{a} \cdot \vec{a}}.$$

- ④ Gegeben ist der Vektor $\vec{a} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$. Gesucht sind die Winkel α und β , die $\vec{a}$ mit den Koordinatenachsen einschließt. Die Winkel erhalten wir aus dem Skalarprodukt von $\vec{a}$ mit $\vec{e}_1$ bzw. mit $\vec{e}_2$:

$$\cos \alpha = \frac{\vec{a} \cdot \vec{e}_1}{|\vec{a}| \cdot |\vec{e}_1|} = \frac{a_x}{|\vec{a}|} = \frac{2}{\sqrt{5}} \Rightarrow \alpha = 26,6^\circ.$$

$$\cos \beta = \frac{\vec{a} \cdot \vec{e}_2}{|\vec{a}| \cdot |\vec{e}_2|} = \frac{a_y}{|\vec{a}|} = \frac{1}{\sqrt{5}} \Rightarrow \beta = 63,4^\circ.$$



- ⑤ Man bestimme zum Vektor $\vec{a} = \begin{pmatrix} a_x \\ a_y \end{pmatrix}$ einen senkrecht dazu stehenden Vektor $\vec{n}$ mit Länge 1:

Einen zu $\vec{a}$ senkrechten Vektor erhält man, indem man die x -Komponente mit der y -Komponente des Vektors $\vec{a}$ vertauscht und anschließend bei einer Komponenten das Vorzeichen wechselt. $\vec{N} = \begin{pmatrix} -a_y \\ a_x \end{pmatrix}$ steht daher senkrecht auf $\vec{a}$, denn

$$\vec{a} \cdot \vec{N} = \begin{pmatrix} a_x \\ a_y \end{pmatrix} \cdot \begin{pmatrix} -a_y \\ a_x \end{pmatrix} = 0.$$

Dividiert man nun noch durch den Betrag von $\vec{N}$, so ist der zugehörige **Normalen-Einheitsvektor** gegeben durch

$$\vec{n} := \vec{e}_n = \frac{1}{N} \vec{N} = \frac{1}{\sqrt{a_x^2 + a_y^2}} \begin{pmatrix} -a_y \\ a_x \end{pmatrix}.$$

□

► 2.1.5 Geometrische Anwendung

Durch den Ortsvektor entspricht jeder Punkt $P = (x, y)$ im $\mathbb{R}^2$ genau einem Vektor $\vec{r}(P) = \begin{pmatrix} x \\ y \end{pmatrix}$. Eine Gerade g durch zwei Punkte $P_1 = (x_1, y_1)$ und $P_2 = (x_2, y_2)$ lässt sich demnach darstellen als Menge aller Punkte P , für die gilt

$$g: \quad \vec{r}(P) = \vec{r}(P_1) + \lambda (\vec{r}(P_2) - \vec{r}(P_1)) = \vec{r}(P_1) + \lambda \cdot \vec{a}.$$

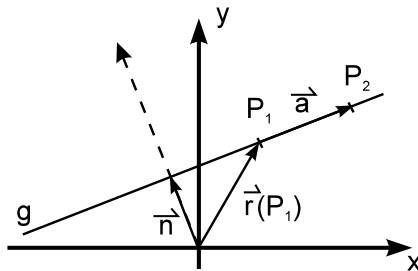


Abb. 2.7. Punkt-Richtungs-Darstellung einer Geraden

Dies ist die **Punkt-Richtungs-Darstellung** einer Geraden, definiert durch den Ortsvektor $\vec{r}(P_1)$ und dem Richtungsvektor $\vec{a} := \vec{r}(P_2) - \vec{r}(P_1)$.

Eine alternative Darstellung der Geradengleichung folgt, wenn wir die Punkt-Richtungs-Darstellung mit dem zu g senkrecht stehenden *Normalen-Einheitsvektor* $\vec{n}$ (siehe Beispiel 2.4 ⑤) skalarmultiplizieren.

$$\begin{pmatrix} x \\ y \end{pmatrix} \cdot \vec{n} = \left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + \lambda \cdot \vec{a} \right) \cdot \vec{n} = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix} \cdot \vec{n} + \lambda \underbrace{\vec{a} \cdot \vec{n}}_{=0}$$

$$\Rightarrow \quad \begin{pmatrix} x \\ y \end{pmatrix} \cdot \begin{pmatrix} n_1 \\ n_2 \end{pmatrix} = d.$$

Dies ist die **Hesse-Normalform** einer Geraden im $\mathbb{R}^2$ und

$$d = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix} \cdot \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$$

ist der kürzeste Abstand der Geraden vom Nullpunkt.

Beispiel 2.5. Gegeben sind zwei Punkte $P_1 = (1, 1)$ und $P_2 = (4, 2)$. Gesucht ist die Punkt-Richtungs-Darstellung sowie die Hesse-Normalform der Geraden g durch die Punkte P_1 und P_2 . Wie groß ist der kleinste Abstand vom Ursprung?

(i) Punkt-Richtungs-Darstellung:

$$g: \vec{r} = \begin{pmatrix} x \\ y \end{pmatrix} = \vec{r}(P_1) + \lambda(\vec{r}(P_2) - \vec{r}(P_1)) = \begin{pmatrix} 1 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 3 \\ 1 \end{pmatrix}.$$

(ii) Hesse-Normalform:

Der Vektor $\vec{N} = \begin{pmatrix} -1 \\ 3 \end{pmatrix}$ steht senkrecht zu $\vec{a} = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$. Wegen $|\vec{N}| = \sqrt{1+9} = \sqrt{10}$ ist $\vec{n} := \frac{1}{N}\vec{N} = \frac{1}{\sqrt{10}} \begin{pmatrix} -1 \\ 3 \end{pmatrix}$ der Normalen-Einheitsvektor zu g .

$$\Rightarrow d = \begin{pmatrix} x \\ y \end{pmatrix} \frac{1}{\sqrt{10}} \begin{pmatrix} -1 \\ 3 \end{pmatrix} \quad \text{ist die Hesse-Normalform.}$$

(iii) Der Minimalabstand der Geraden zum Ursprung erhält man, indem man den Punkt $P_1 = (1, 1)$ in die Hesse-Normalform einsetzt:

$$d = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \frac{1}{\sqrt{10}} \begin{pmatrix} -1 \\ 3 \end{pmatrix} = \frac{2}{\sqrt{10}} = \frac{1}{5}\sqrt{10}.$$

(iv) Berechnen wir noch das Skalarprodukt auf der linken Seite der Hesse-Normalform

$$\frac{1}{\sqrt{10}}(-x + 3y) = \frac{2}{\sqrt{10}}$$

und lösen nach y auf, erhalten wir die übliche Darstellung der Geradengleichung in der Ebene

$$y = \frac{1}{3}x + \frac{2}{3}. \quad \square$$



Visualisierung mit MAPLE: Auf der CD-Rom befinden sich MAPLE-Prozeduren, welche sowohl die [Darstellung von Vektoren im \$\mathbb{R}^2\$](#) ermöglichen als auch die Visualisierung der in 2.1.1 bis 2.1.4 beschriebenen Vektoroperationen.

2.2 Vektoren im $\mathbb{R}^3$

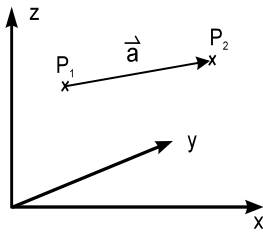


Abb. 2.8. Richtungsvektor

Analog zum Vorgehen im zweidimensionalen Raum $\mathbb{R}^2$ führt man Vektoren im $\mathbb{R}^3$ ein, indem ein Vektor $\vec{a}$ in einem rechtwinkligen Koordinatensystem vom Punkt $P_1 = (x_1, y_1, z_1)$ zum Punkt $P_2 = (x_2, y_2, z_2)$ festgelegt ist:

$$\vec{a} := \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \\ z_2 - z_1 \end{pmatrix}.$$

$\vec{a}$ heißt dann wieder *Richtungsvektor*. Ein *Ortsvektor* $\vec{r}(P)$ stellt einen Vektor vom Ursprung O zum Punkt $P = (x, y, z)$ dar:

$$\vec{r}(P) = \begin{pmatrix} x \\ y \\ z \end{pmatrix}.$$

2.2.1 Rechenregeln für Vektoren

Die Multiplikation eines Vektors $\vec{a}$ mit einem Skalar λ und die Addition zweier Vektoren erfolgen komponentenweise:

$$\lambda \cdot \vec{a} = \lambda \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} := \begin{pmatrix} \lambda a_x \\ \lambda a_y \\ \lambda a_z \end{pmatrix} \quad (\text{Skalare Multiplikation})$$

$$\vec{a} + \vec{b} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} + \begin{pmatrix} b_x \\ b_y \\ b_z \end{pmatrix} := \begin{pmatrix} a_x + b_x \\ a_y + b_y \\ a_z + b_z \end{pmatrix} \quad (\text{Addition})$$

Die *Länge* (bzw. der *Betrag*) eines Vektors $\vec{a}$ ist gegeben durch

$$a := |\vec{a}| = \sqrt{a_x^2 + a_y^2 + a_z^2} \quad (\text{Betrag})$$

und entspricht der Diagonalen eines Quaders mit Kantenlängen a_x, a_y, a_z . Jeder Vektor $\vec{e}$ mit $|\vec{e}| = 1$ heißt *Einheitsvektor*. Die *Koordinaten-Einheitsvektoren* lauten

$$\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Jeder Vektor $\vec{a}$ lässt sich schreiben als *Linearkombination* dieser Einheitsvektoren

$$\vec{a} = a_x \vec{e}_1 + a_y \vec{e}_2 + a_z \vec{e}_3.$$

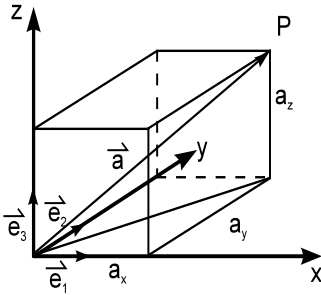


Abb. 2.9. $\vec{a}$ als Linearkombination der Einheitsvektoren

Beispiele 2.6:

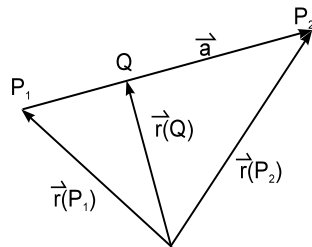
① Der Ortsvektor zum Punkt $P = (5, 1, -3)$ lautet $\vec{r}(P) = \begin{pmatrix} 5 \\ 1 \\ -3 \end{pmatrix}$. Er hat die Länge $|\vec{r}(P)| = \sqrt{5^2 + 1^2 + (-3)^2} = \sqrt{35}$.

② Der Richtungsvektor von $P_1 = (3, 4, 7)$ nach $P_2 = (7, 3, 1)$ ist

$$\vec{a} = \overrightarrow{P_1 P_2} = \vec{r}(P_2) - \vec{r}(P_1) = \begin{pmatrix} 7 \\ 3 \\ 1 \end{pmatrix} - \begin{pmatrix} 3 \\ 4 \\ 7 \end{pmatrix} = \begin{pmatrix} 4 \\ -1 \\ -6 \end{pmatrix}.$$

③ Gesucht sind die Koordinaten des Punktes Q , welcher die Strecke von $P_1 = (3, 4, 7)$ zum Punkt $P_2 = (7, 3, 1)$ im Verhältnis 1:2 schneidet.

$$\begin{aligned} \vec{r}(Q) &= \vec{r}(P_1) + \frac{1}{3}(\vec{r}(P_2) - \vec{r}(P_1)) \\ &= \vec{r}(P_1) + \frac{1}{3}\vec{a} = \begin{pmatrix} 3 \\ 4 \\ 7 \end{pmatrix} + \frac{1}{3}\begin{pmatrix} 4 \\ -1 \\ -6 \end{pmatrix} \\ &= \frac{1}{3}\begin{pmatrix} 13 \\ 11 \\ 15 \end{pmatrix}. \end{aligned}$$



Das **Skalarprodukt** ist im $\mathbb{R}^3$ definiert durch

$$\vec{a} \cdot \vec{b} := |\vec{a}| \cdot |\vec{b}| \cdot \cos \alpha, \quad \alpha = \angle(\vec{a}, \vec{b}),$$

wenn α der von den Vektoren $\vec{a}$ und $\vec{b}$ eingeschlossene Winkel ist. Für die Darstellung des Skalarprodukts berechnet man mit den gleichen Regeln wie im $\mathbb{R}^2$ ($(S_1) - (S_3)$), dass

$$\vec{a} \cdot \vec{b} = a_x b_x + a_y b_y + a_z b_z.$$

Folglich gilt wieder für den von zwei Vektoren eingeschlossenen Winkel α :

$$\cos \alpha = \frac{\vec{a} \cdot \vec{b}}{|\vec{a}| \cdot |\vec{b}|} = \frac{a_x b_x + a_y b_y + a_z b_z}{\sqrt{a_x^2 + a_y^2 + a_z^2} \sqrt{b_x^2 + b_y^2 + b_z^2}}. \quad (3)$$

Folgerung: Zwei Vektoren $\vec{a}$ und $\vec{b}$ stehen **senkrecht** aufeinander, wenn das Skalarprodukt verschwindet:

$$\vec{a} \perp \vec{b} \Leftrightarrow \vec{a} \cdot \vec{b} = 0$$

Beispiele 2.7:

- ① **Orthonormalsystem:** $\vec{e}_1, \vec{e}_2, \vec{e}_3$ bilden ein *Orthonormalsystem* von $\mathbb{R}^3$, d.h. sie stehen senkrecht aufeinander und haben die Länge 1:

$$\vec{e}_1 \cdot \vec{e}_1 = \vec{e}_2 \cdot \vec{e}_2 = \vec{e}_3 \cdot \vec{e}_3 = 1$$

$$\vec{e}_1 \cdot \vec{e}_2 = \vec{e}_2 \cdot \vec{e}_3 = \vec{e}_3 \cdot \vec{e}_1 = 0.$$

- ② Die Vektoren $\vec{a}_1 = \frac{1}{\sqrt{3}} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$, $\vec{a}_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}$, $\vec{a}_3 = \frac{1}{\sqrt{6}} \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix}$ bilden ebenfalls ein Orthonormalsystem.

- ③ **Richtungskosinus:** Durch das Skalarprodukt lassen sich auf einfache Weise die Winkel berechnen, die ein Vektor $\vec{a} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix}$ mit den Koordinatenachsen einschließt.

$$\vec{a} \cdot \vec{e}_1 = a_x; |\vec{e}_1| = 1 \Rightarrow \cos \alpha = \frac{a_x}{a}$$

$$\vec{a} \cdot \vec{e}_2 = a_y; |\vec{e}_2| = 1 \Rightarrow \cos \beta = \frac{a_y}{a}$$

$$\vec{a} \cdot \vec{e}_3 = a_z; |\vec{e}_3| = 1 \Rightarrow \cos \gamma = \frac{a_z}{a}.$$

Die Winkel α, β, γ heißen *Richtungskosinus* von $\vec{a}$. Es gilt

$$\cos^2 \alpha + \cos^2 \beta + \cos^2 \gamma = \frac{a_x^2}{a^2} + \frac{a_y^2}{a^2} + \frac{a_z^2}{a^2} = \frac{a^2}{a^2} = 1 \quad (4)$$

und

$$a_x = |\vec{a}| \cos \alpha, \quad a_y = |\vec{a}| \cos \beta, \quad a_z = |\vec{a}| \cos \gamma.$$

Durch Gleichung (4) sind für ein Vektor $\vec{a}$ die 3 Winkel zu den Koordinatenachsen nicht beliebig wählbar. Nur 2 Winkel sind frei; der dritte bestimmt sich aus (4).

- ④ **Zahlenbeispiel:** Gegeben sind die Vektoren $\vec{a} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} -4 \\ 3 \\ -2 \end{pmatrix}$.

Gesucht ist der Winkel α zwischen $\vec{a}$ und $\vec{b}$:

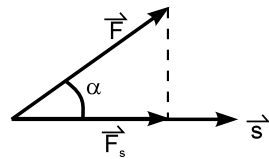
$$\left. \begin{aligned} \vec{a} \cdot \vec{b} &= \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} -4 \\ 3 \\ -2 \end{pmatrix} = -4 + 6 - 6 = -4 \\ |\vec{a}| &= \sqrt{14}, \quad |\vec{b}| = \sqrt{29} \end{aligned} \right\} \Rightarrow \cos \alpha = \frac{-4}{\sqrt{14} \cdot \sqrt{29}}$$

$$\Rightarrow \alpha = 101,45^\circ.$$

□

2.2.2 Projektion eines Vektors

Wir betrachten die folgende physikalische Problemstellung: Ein Massenpunkt ist in eine Schiene eingespannt und kann nur entlang der Richtung $\vec{s}$ bewegt werden. Auf diesen Massenpunkt wirkt eine Kraft $\vec{F}$. Gesucht ist die Kraft $\vec{F}_s$ in Richtung $\vec{s}$:



Der Betrag von $\vec{F}_s$ ist mit Gleichung (3) gegeben durch

$$|\vec{F}_s| = |\vec{F}| \cdot \cos \alpha = |\vec{F}| \cdot \frac{\vec{F} \cdot \vec{s}}{|\vec{F}| \cdot |\vec{s}|}$$

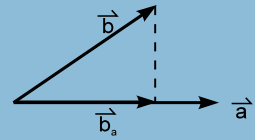
und die Richtung durch $\vec{e}_s = \frac{\vec{s}}{|\vec{s}|}$. Also ist

$$\vec{F}_s = |\vec{F}_s| \cdot \vec{e}_s = \frac{\vec{F} \cdot \vec{s}}{|\vec{s}|} \cdot \frac{\vec{s}}{|\vec{s}|} = \frac{\vec{F} \cdot \vec{s}}{|\vec{s}|^2} \vec{s}.$$

Man nennt $\vec{F}_s$ die *Projektion von $\vec{F}$ in Richtung $\vec{s}$* . Man verallgemeinert diese Konstruktion für zwei beliebige Vektoren $\vec{a}$ und $\vec{b}$:

Die **Projektion von $\vec{b}$ in Richtung $\vec{a}$** ist gegeben durch den Vektor

$$\vec{b}_a = \frac{\vec{a} \cdot \vec{b}}{|\vec{a}|^2} \cdot \vec{a}$$



⚠ Man beachte, dass $\vec{a} \cdot \vec{b}$ das Skalarprodukt bedeutet und daher $\frac{\vec{a} \cdot \vec{b}}{a^2}$ eine reelle Zahl darstellt. Das zweite Produktzeichen ist die Multiplikation des Vektors $\vec{a}$ mit dieser reellen Zahl. Beide „ $\cdot$ “-Zeichen dürfen **nicht** vertauscht werden!

Anwendungsbeispiel 2.8 (Kraft entlang einer vorgegebenen Richtung).

Gegeben ist die Kraft $\vec{F} = \begin{pmatrix} 2 \\ 2 \\ -7 \end{pmatrix}$, die auf eine Masse wirkt, die sich nur entlang der Richtung $\vec{s} = \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix}$ bewegen kann. Gesucht sind die Beschleunigungskraft, die Stärke der Kraft sowie die verrichtete Arbeit.

Die Beschleunigungskraft ergibt sich durch die Projektion der Kraft $\vec{F}$ in Richtung der Bewegung. Zur Bestimmung dieser Kraft berechnen wir zuerst das Skalarprodukt von $\vec{F}$ mit $\vec{s}$:

$$\vec{F} \cdot \vec{s} = \begin{pmatrix} 2 \\ 2 \\ -7 \end{pmatrix} \cdot \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix} = 3, \quad |\vec{s}|^2 = 3.$$

Die Beschleunigungskraft ergibt sich nun aus der Projektionsformel

$$\vec{F}_s = \frac{3}{3} \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix} = \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix} \quad \text{und} \quad |\vec{F}_s| = \sqrt{3}.$$

Die verrichtete Arbeit W ergibt sich nach Beispiel 2.2 durch die Formel

$$W := \vec{F} \cdot \vec{s}.$$

Es ist daher $W = 3$.

□

2.2.3 Das Vektorprodukt (Kreuzprodukt) zweier Vektoren

Im $\mathbb{R}^3$ definiert man für zwei Vektoren das *Vektorprodukt* (auch *Kreuzprodukt* genannt), dessen Ergebnis wieder ein Vektor ist:

Definition: Unter dem **Vektorprodukt (Kreuzprodukt)**

$$\vec{c} = \vec{a} \times \vec{b}$$

zweier Vektoren $\vec{a}$ und $\vec{b}$ versteht man den Vektor $\vec{c}$ mit den folgenden Eigenschaften:

- (1) $\vec{c}$ ist sowohl zu $\vec{a}$ als auch zu $\vec{b}$ senkrecht:
 $\vec{c} \cdot \vec{a} = 0, \vec{c} \cdot \vec{b} = 0$.
- (2) Der Betrag von $\vec{c}$ ist gleich dem Produkt aus den Beträgen der Vektoren $\vec{a}$ und $\vec{b}$ und dem Sinus des eingeschlossenen Winkels α :
 $|\vec{c}| = |\vec{a}| \cdot |\vec{b}| \cdot \sin \alpha$, wenn α der Winkel, den die Vektoren $\vec{a}$ und $\vec{b}$ miteinander einschließen.
- (3) Die Vektoren $\vec{a}, \vec{b}, \vec{c}$ bilden ein Rechtssystem.

Bemerkungen:

- (1) Im Gegensatz zum Skalarprodukt ist das Vektorprodukt ein Vektor.
- (2) Statt $\vec{a} \times \vec{b}$ wird auch oftmals das Symbol $\left[\vec{a}, \vec{b} \right]$ verwendet.
- (3) ⚠ **Achtung:** Das Vektorprodukt ist nur in $\mathbb{R}^3$ definiert!
- (4) ⚠ **Achtung:** Das Vektorprodukt ist **nicht** kommutativ!

Geometrische Deutung: Da $\vec{c} \perp \vec{a}$ und $\vec{c} \perp \vec{b}$ steht, kommt als Richtung des Vektors $\vec{c}$ nur die in Abb. 2.10 gestrichelte Linie in Frage. Da $\vec{a}, \vec{b}, \vec{c}$ in dieser Reihenfolge ein Rechtssystem bilden, bleibt nur der nach oben weisende Teil. Der Flächeninhalt des von $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms ist Grundseite mal Höhe, also

$$\begin{aligned} A &= |\vec{a}| \cdot h = |\vec{a}| \cdot |\vec{b}| \cdot \sin \alpha \\ &= \left| \vec{a} \times \vec{b} \right|. \end{aligned}$$

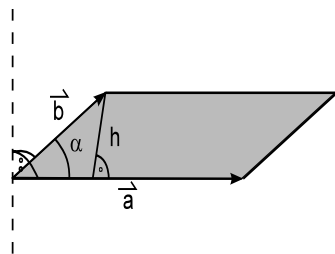


Abb. 2.10. Kreuzprodukt $\vec{a}$ mit $\vec{b}$

Der Betrag des Vektorproduktes entspricht dem **Flächeninhalt** des von den Vektoren $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms.

Beispiele 2.9:

- ① Die Vektorprodukte der Einheitsvektoren lassen sich aufgrund der Definition sofort berechnen:

$$\vec{e}_1 \times \vec{e}_1 = 0, \vec{e}_2 \times \vec{e}_2 = 0, \vec{e}_3 \times \vec{e}_3 = 0;$$

$$\vec{e}_1 \times \vec{e}_2 = \vec{e}_3, \vec{e}_2 \times \vec{e}_3 = \vec{e}_1, \vec{e}_3 \times \vec{e}_1 = \vec{e}_2.$$

- ② Kriterium für **kollineare Vektoren**:

Verschwundet das Kreuzprodukt von $\vec{a} \neq 0$ und $\vec{b} \neq 0$, so ist entweder $\vec{a} \uparrow \vec{b}$ ($\vec{a}$ parallel zu $\vec{b}$) oder $\vec{a} \downarrow \vec{b}$ ($\vec{a}$ antiparallel zu $\vec{b}$). $\square$

Wir geben für das Vektorprodukt die wesentlichen **Rechenregeln** an:

(V ₁)	$\vec{a} \times (\vec{b} + \vec{c})$	$= \vec{a} \times \vec{b} + \vec{a} \times \vec{c}$	Distributivgesetze
	$(\vec{a} + \vec{b}) \times \vec{c}$	$= \vec{a} \times \vec{c} + \vec{b} \times \vec{c}$	
(V ₂)	$\vec{a} \times \vec{b}$	$= -\vec{b} \times \vec{a}$	Anti-Symmetriengesetz
(V ₃)	$\lambda \cdot (\vec{a} \times \vec{b})$	$= (\lambda \vec{a}) \times \vec{b}$	Multiplikation mit Skalar λ
		$= \vec{a} \times (\lambda \vec{b})$	

③ **Formel für das Vektorprodukt**

Mit Hilfe der Rechengesetze erhalten wir eine für die Praxis brauchbare Darstellung des Vektorproduktes über die Komponenten der Vektoren, denn es gilt

$$\begin{aligned}
 \vec{a} \times \vec{b} &= (a_x \vec{e}_1 + a_y \vec{e}_2 + a_z \vec{e}_3) \times (b_x \vec{e}_1 + b_y \vec{e}_2 + b_z \vec{e}_3) \\
 &= a_x b_x \underbrace{(\vec{e}_1 \times \vec{e}_1)}_{=0} + a_x b_y \underbrace{(\vec{e}_1 \times \vec{e}_2)}_{\vec{e}_3} + a_x b_z \underbrace{(\vec{e}_1 \times \vec{e}_3)}_{-\vec{e}_2} + \\
 &\quad a_y b_x \underbrace{(\vec{e}_2 \times \vec{e}_1)}_{-\vec{e}_3} + a_y b_y \underbrace{(\vec{e}_2 \times \vec{e}_2)}_{=0} + a_y b_z \underbrace{(\vec{e}_2 \times \vec{e}_3)}_{\vec{e}_1} + \\
 &\quad a_z b_x \underbrace{(\vec{e}_3 \times \vec{e}_1)}_{\vec{e}_2} + a_z b_y \underbrace{(\vec{e}_3 \times \vec{e}_2)}_{-\vec{e}_1} + a_z b_z \underbrace{(\vec{e}_3 \times \vec{e}_3)}_{=0} \\
 &= (a_y b_z - a_z b_y) \vec{e}_1 + (a_z b_x - a_x b_z) \vec{e}_2 + (a_x b_y - a_y b_x) \vec{e}_3.
 \end{aligned}$$

Somit ist

$$\vec{a} \times \vec{b} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix} \times \begin{pmatrix} b_x \\ b_y \\ b_z \end{pmatrix} = \begin{pmatrix} a_y b_z - a_z b_y \\ a_z b_x - a_x b_z \\ a_x b_y - a_y b_x \end{pmatrix}.$$

Merkregel: Formal lässt sich das Vektorprodukt in der Form einer *dreireihigen Determinante* ($\rightarrow$ Kapitel 3.2) darstellen, wenn man nach der ersten Spalte entwickelt:

$$\vec{a} \times \vec{b} = \begin{vmatrix} \vec{e}_1 & a_x & b_x \\ \vec{e}_2 & a_y & b_y \\ \vec{e}_3 & a_z & b_z \end{vmatrix} := \begin{vmatrix} a_y & b_y \\ a_z & b_z \end{vmatrix} \vec{e}_1 - \begin{vmatrix} a_x & b_x \\ a_z & b_z \end{vmatrix} \vec{e}_2 + \begin{vmatrix} a_x & b_x \\ a_y & b_y \end{vmatrix} \vec{e}_3.$$

Der Wert einer *zweireihigen Determinante* ist definiert durch die Differenz von Haupt- und Nebendiagonal-Produkten.

$$\begin{vmatrix} a & b \\ c & d \end{vmatrix} := a \cdot d - b \cdot c.$$

$$\Rightarrow \vec{a} \times \vec{b} = (a_y b_z - b_y a_z) \vec{e}_1 - (a_x b_z - a_z b_x) \vec{e}_2 + (a_x b_y - a_y b_x) \vec{e}_3.$$

Beispiel 2.10. Gesucht ist ein Vektor $\vec{c}$, der auf den beiden Vektoren

$$\vec{a} = \begin{pmatrix} 2 \\ 1 \\ -2 \end{pmatrix} \text{ und } \vec{b} = \begin{pmatrix} -1 \\ 2 \\ -1 \end{pmatrix}$$

senkrecht steht und der *Normalen-Einheitsvektor* $\vec{e}_c$, der zusätzlich die Länge 1 besitzt:

$$\begin{aligned} c = \vec{a} \times \vec{b} &= \begin{vmatrix} \vec{e}_1 & 2 & -1 \\ \vec{e}_2 & 1 & 2 \\ \vec{e}_3 & -2 & -1 \end{vmatrix} = \begin{vmatrix} 1 & 2 \\ -2 & -1 \end{vmatrix} \vec{e}_1 - \begin{vmatrix} 2 & -1 \\ -2 & -1 \end{vmatrix} \vec{e}_2 + \begin{vmatrix} 2 & -1 \\ 1 & 2 \end{vmatrix} \vec{e}_3 \\ &= \begin{pmatrix} 3 \\ 4 \\ 5 \end{pmatrix}. \end{aligned}$$

Der zugehörige Normalen-Einheitsvektor ist $\vec{e}_c = \frac{1}{|\vec{c}|} \vec{c} = \frac{1}{\sqrt{50}} \vec{c}$ □

Beispiel 2.11. Gegeben sind die Vektoren $\vec{a} = \begin{pmatrix} 1 \\ -4 \\ 1 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix}$.

Gesucht ist das Kreuzprodukt dieser beiden Vektoren sowie der Flächeninhalt des von beiden Vektoren aufgespannten Parallelogramms.

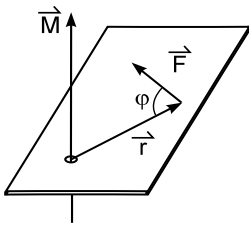
$$\begin{aligned} \text{(i) } \vec{a} \times \vec{b} &= \begin{vmatrix} \vec{e}_1 & 1 & 2 \\ \vec{e}_2 & -4 & 0 \\ \vec{e}_3 & 1 & 2 \end{vmatrix} = \begin{vmatrix} -4 & 0 \\ 1 & 2 \end{vmatrix} \vec{e}_1 - \begin{vmatrix} 1 & 2 \\ 1 & 2 \end{vmatrix} \vec{e}_2 + \begin{vmatrix} 1 & 2 \\ -4 & 0 \end{vmatrix} \vec{e}_3 \\ &= -8 \vec{e}_1 - 0 \vec{e}_2 + 8 \vec{e}_3 = \begin{pmatrix} -8 \\ 0 \\ 8 \end{pmatrix}. \end{aligned}$$

(ii) Der Flächeninhalt des von $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms ist

$$A = |\vec{a} \times \vec{b}| = \left| \begin{pmatrix} -8 \\ 0 \\ 8 \end{pmatrix} \right| = \sqrt{64 + 64} = \sqrt{128} = 8\sqrt{2}. \quad \square$$

Anwendungsbeispiel 2.12 (Auftreten des Kreuzproduktes in der Physik).

In der Physik tritt das Kreuzprodukt z.B. in den folgenden wichtigen Fällen auf:



① **Drehmoment:** Ein Körper sei um einen festen Punkt O drehbar und im Punkte P dieses Körpers greife eine Kraft $\vec{F}$ an. Dann ist die Größe M das *Drehmoment* von $\vec{F}$ bezüglich O

$$M = |\vec{r}| \cdot |\vec{F}| \cdot \sin \varphi$$

(Kraft mal Hebelarm). Der Drehmomentvektor steht senkrecht zu der durch $\vec{r}$ und $\vec{F}$ gebildeten Ebene und kann als Richtung der Drehachse aufgefasst werden:

$$\vec{M} = \vec{r} \times \vec{F}.$$

② **Drehimpuls:** Sei O ein fester Bezugspunkt. Eine Masse m befinde sich in einem bestimmten Augenblick in P und besitze die Geschwindigkeit $\vec{v}$. Dann lautet der momentane *Drehimpuls* $\vec{L}$ des Massenpunktes bzgl. O

$$\vec{L} = m \vec{r} \times \vec{v},$$

wenn $\vec{r} = \overrightarrow{OP} = \vec{r}(P)$ der Ortsvektor zum Punkt P .

③ **Lorentz-Kraft:** Bewegt sich ein geladenes Teilchen (Ladung q) mit der Geschwindigkeit $\vec{v}$ durch ein Magnetfeld, so erfährt es eine zu $\vec{B}$ und $\vec{v}$ senkrechte *Lorentz-Kraft*:

$$\vec{F}_L = q \vec{v} \times \vec{B}$$

□

2.2.4 Das Spatprodukt von drei Vektoren

In der Mechanik kommt das Produkt $(\vec{a} \times \vec{b}) \cdot \vec{c}$ vor. Der Klammerausdruck ist ein Vektor, der skalarmultipliziert mit $\vec{c}$ wird. Das Ergebnis ist also eine reelle Zahl.

Definition: Unter dem **Spatprodukt** $[\vec{a}, \vec{b}, \vec{c}]$ von drei Vektoren versteht man die reelle Zahl

$$[\vec{a}, \vec{b}, \vec{c}] := (\vec{a} \times \vec{b}) \cdot \vec{c}.$$

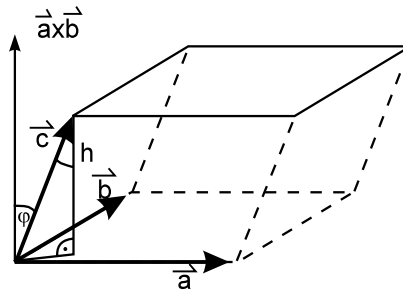


Abb. 2.11. Spatprodukt von drei Vektoren

Für das Spatprodukt gelten die **Rechenregeln**

$$(1) \quad [\lambda \vec{a}, \vec{b}, \vec{c}] = \lambda [\vec{a}, \vec{b}, \vec{c}] \quad \text{usw.}$$

$$(2) \quad [\vec{a} + \vec{b}, \vec{b}, \vec{c}] = [\vec{a}, \vec{b}, \vec{c}] \quad \text{usw.}$$

$$(3) \quad [\vec{e}_1, \vec{e}_2, \vec{e}_3] = 1$$

Bilden die Vektoren $\vec{a}, \vec{b}, \vec{c}$ ein *Rechtssystem* (Linkssystem), so ist das Spatprodukt *positiv* (negativ).

Geometrische Interpretation: Das Volumen des von den Vektoren $\vec{a}, \vec{b}$ und $\vec{c}$ aufgespannten Spates (Parallelotops) ist gegeben durch Grundfläche G mal Höhe h . Die Grundfläche ist nach Definition des Kreuzproduktes $G = |\vec{a} \times \vec{b}|$ und die Höhe $h = |\vec{c}| \cos \varphi$.

$$\Rightarrow V = |\vec{a} \times \vec{b}| \cdot |\vec{c}| \cdot \cos \varphi = |(\vec{a} \times \vec{b}) \cdot \vec{c}|.$$

Das **Volumen des Spates** ist gleich dem Betrag des Spatproduktes. Der Wert des Spatproduktes erhält man durch Ausrechnen

$$\left[\vec{a}, \vec{b}, \vec{c} \right] = a_x b_y c_z + a_y b_z c_x + a_z b_x c_y - a_z b_y c_x - a_y b_x c_z - a_x b_z c_y.$$

Die Rechnung kann man auch auffassen als das Ergebnis der Entwicklung der Determinante

$$\left[\vec{a}, \vec{b}, \vec{c} \right] = \begin{vmatrix} a_x & b_x & c_x \\ a_y & b_y & c_y \\ a_z & b_z & c_z \end{vmatrix}.$$

Beispiel 2.13. Gesucht ist das Spatprodukt der drei folgenden Vektoren:

$$\vec{a} = \begin{pmatrix} 4 \\ 2 \\ 1 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 3 \\ -1 \\ 2 \end{pmatrix}, \quad \vec{c} = \begin{pmatrix} 0 \\ 2 \\ -5 \end{pmatrix}.$$

Zur Bestimmung des Spatproduktes der drei Vektoren berechnen wir die dreireihigen Determinante durch Entwicklung nach der ersten Spalte:

$$\begin{vmatrix} 4 & 3 & 0 \\ 2 & -1 & 2 \\ 1 & 2 & -5 \end{vmatrix} = 4 \cdot \begin{vmatrix} -1 & 2 \\ 2 & -5 \end{vmatrix} - 2 \cdot \begin{vmatrix} 3 & 0 \\ 2 & -5 \end{vmatrix} + 1 \cdot \begin{vmatrix} 3 & 0 \\ -1 & 2 \end{vmatrix} \\ = 4((-1)(-5) - 2 \cdot 2) - 2(3 \cdot (-5) - 2 \cdot 0) + (3 \cdot 2 - (-1) \cdot 0) = 40. \quad \square$$

Aus der Interpretation des Spatproduktes als das Volumen des von $\vec{a}$, $\vec{b}$ und $\vec{c}$ aufgespannten Spates ergibt sich folgende wichtige Folgerung

Folgerung: Das Spatprodukt ist Null, wenn die drei Vektoren in einer Ebene liegen:

$$\left[\vec{a}, \vec{b}, \vec{c} \right] = 0 \Leftrightarrow \vec{a}, \vec{b}, \vec{c} \text{ liegen in einer Ebene.}$$

Es gilt bei der Berechnung des Spatproduktes die folgende **Regel**

$$(\vec{a} \times \vec{b}) \cdot \vec{c} = (\vec{b} \times \vec{c}) \cdot \vec{a} = (\vec{c} \times \vec{a}) \cdot \vec{b}.$$

Das Vorzeichen des Spatproduktes ändert sich, wenn man von der zyklischen Reihenfolge abc , bca oder cab abweicht.



Visualisierung mit MAPLE: Auf der CD-Rom befinden sich MAPLE-Prozeduren, welche sowohl die **Darstellung von Vektoren im $\mathbb{R}^3$** ermöglichen als auch die Visualisierung der Vektoroperationen.

2.3 Geraden und Ebenen im $\mathbb{R}^3$

In diesem Abschnitt werden einige Anwendungen der Vektoren und der Vektorrechnung angegeben: Die Beschreibung von Geraden und Ebenen im $\mathbb{R}^3$ sowie Abstandsberechnungen und Lage von Punkten, Geraden und Ebenen zueinander.

2.3.1 Vektorielle Darstellung von Geraden

Eine Gerade g ist eindeutig durch die Angabe zweier verschiedener Punkte $P_1 = (x_1, y_1, z_1)$ und $P_2 = (x_2, y_2, z_2)$ festgelegt. Denn durch $\vec{a} := \overrightarrow{P_1 P_2}$ ist der Richtungsvektor der Geraden festgelegt und jeder Punkt $P = (x, y, z)$ der Geraden lässt sich nach Abb. 2.12 darstellen als

$$g : \vec{r}(P) = \vec{r}(P_1) + \lambda \vec{a}, \quad \lambda \in \mathbb{R},$$

(Punkt-Richtungsform einer Geraden).

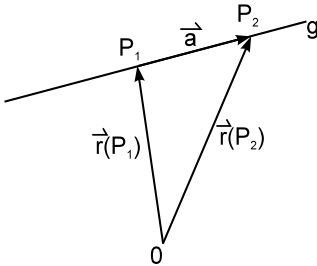


Abb. 2.12. Beschreibung einer Geraden

Ersetzt man den Vektor $\vec{a} := \overrightarrow{P_1 P_2}$ durch $\vec{r}(P_2) - \vec{r}(P_1)$, so erhält man:

$$g : \vec{r}(P) = \vec{r}(P_1) + \lambda (\vec{r}(P_2) - \vec{r}(P_1)), \quad \lambda \in \mathbb{R},$$

(Zweipunkteform einer Geraden).

Ein Punkt Q liegt auf einer Geraden g , falls die entsprechende Vektorgleichung

$$\vec{r}(Q) = \vec{r}(P_1) + \lambda \vec{a}$$

eine Lösung für λ besitzt.

Beispiel 2.14. Gegeben sind die Punkte $P_1 = (2, 0, 4)$ und $P_2 = (2, 2, 2)$. Liegt der Punkt $Q = (2, -2, 6)$ auf der Geraden g durch die Punkte P_1 und P_2 ?

Die Geradengleichung für g lautet mit dem Richtungsvektor

$$\vec{a} = \overrightarrow{P_1 P_2} = \vec{r}(P_2) - \vec{r}(P_1) = \begin{pmatrix} 2 \\ 2 \\ 2 \end{pmatrix} - \begin{pmatrix} 2 \\ 0 \\ 4 \end{pmatrix} = \begin{pmatrix} 0 \\ 2 \\ -2 \end{pmatrix}$$

$$g: \vec{r}(P) = \vec{r}(P_1) + \lambda \vec{a} = \begin{pmatrix} 2 \\ 0 \\ 4 \end{pmatrix} + \lambda \begin{pmatrix} 0 \\ 2 \\ -2 \end{pmatrix}.$$

Der Punkt Q liegt auf der Geraden g , wenn die Vektorgleichung

$$\vec{r}(Q) = \vec{r}(P_1) + \lambda \vec{a}$$

eine Lösung besitzt. Dies ist dann der Fall, wenn die Gleichung $\begin{pmatrix} 2 \\ -2 \\ 6 \end{pmatrix} =$

$\begin{pmatrix} 2 \\ 0 \\ 4 \end{pmatrix} + \lambda \begin{pmatrix} 0 \\ 2 \\ -2 \end{pmatrix}$ lösbar ist. Bringen wir den Vektor $\vec{r}(P_1)$ auf die linke Seite gilt $\lambda \begin{pmatrix} 0 \\ 2 \\ -2 \end{pmatrix} = \begin{pmatrix} 0 \\ -2 \\ 2 \end{pmatrix}$. Für $\lambda = -1$ ist diese Gleichung erfüllt und der Punkt Q liegt daher auf g . $\square$

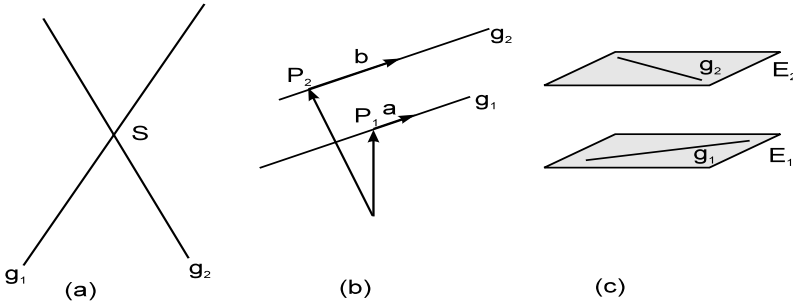
2.3.2 Lage zweier Geraden zueinander

Zwei Geraden

$$g_1: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} \quad \text{und} \quad g_2: \vec{x} = \vec{r}(P_2) + \mu \vec{b}$$

können im $\mathbb{R}^3$ vier verschiedene Lagen zueinander besitzen:

- (1) g_1 und g_2 **schneiden sich** in genau einem Punkt S (Abb. 2.13 (a)).
- (2) g_1 und g_2 **fallen zusammen**. Dies ist dann der Fall, wenn $\vec{a} \parallel \vec{b}$ und $\overrightarrow{P_1 P_2} \parallel \vec{a}$.
- (3) g_1 und g_2 sind **parallel, fallen aber nicht zusammen** (Abb. 2.13 (b)). Dies ist dann der Fall, wenn $\vec{a} \parallel \vec{b}$ und $\overrightarrow{P_1 P_2} \nparallel \vec{a}$.
- (4) g_1 und g_2 sind **windschief**: Sie verlaufen weder parallel noch schneiden sie sich in einem Punkt (Abb. 2.13 (c)).

Abb. 2.13. Lage zweier Geraden g_1 und g_2

Um die Lage zweier Geraden rechnerisch zu bestimmen, genügt es die Vektorgleichung $\vec{x}_{g_1} = \vec{x}_{g_2}$ zu lösen:

$$\begin{aligned} \Rightarrow \vec{r}(P_1) + \lambda \vec{a} &= \vec{r}(P_2) + \mu \vec{b} \\ \Rightarrow \lambda \vec{a} - \mu \vec{b} &= \vec{r}(P_2) - \vec{r}(P_1) = \overrightarrow{P_1 P_2}. \end{aligned}$$

Dies ist ein LGS für die Unbekannten λ und μ , denn für

$$\vec{a} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix}, \vec{b} = \begin{pmatrix} b_x \\ b_y \\ b_z \end{pmatrix}, \vec{r}(P_1) = \begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix}, \vec{r}(P_2) = \begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix}$$

lautet das LGS komponentenweise

$$\begin{aligned} \lambda a_x - \mu b_x &= x_2 - x_1 \\ \lambda a_y - \mu b_y &= y_2 - y_1 \\ \lambda a_z - \mu b_z &= z_2 - z_1 \end{aligned}$$

bzw. in Matrizen Schreibweise

$$\left(\begin{array}{cc|c} a_x & -b_x & x_2 - x_1 \\ a_y & -b_y & y_2 - y_1 \\ a_z & -b_z & z_2 - z_1 \end{array} \right).$$

Es gilt dann

- (1) Besitzt das lineare Gleichungssystem für λ und μ genau eine Lösung, dann schneiden sich g_1 und g_2 genau in einem Punkt.
- (2) Besitzt das lineare Gleichungssystem für λ und μ unendlich viele Lösungen, dann fallen g_1 und g_2 zusammen.
- (3) Besitzt das lineare Gleichungssystem für λ und μ keine Lösung, dann sind g_1 und g_2 windschief **oder** sie sind parallel aber nicht zusammenfallend.

Beispiel 2.15. Gegeben ist die Gerade g_1 definiert durch den Richtungsvektor $\vec{a} = \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix}$ und den Punkt $P_1 = (3, 2, 1)$ sowie die Gerade g_2 durch die Punkte $P_2 = (4, 0, -1)$, $P_3 = (-2, -1, -1)$. Man bestimme die Lage der beiden Geraden zueinander. Es ist

$$g_1 : \vec{x} = \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix} \quad g_2 : \vec{x} = \begin{pmatrix} 4 \\ 0 \\ -1 \end{pmatrix} + \mu \begin{pmatrix} -6 \\ -1 \\ 0 \end{pmatrix}.$$

Da die Richtungsvektoren von g_1 und g_2 nicht parallel sind, können beide Geraden sich entweder nur schneiden oder sie sind windschief. Wir setzen die Vektorgleichung $\vec{x}_{g_1} = \vec{x}_{g_2}$ an:

$$\begin{aligned} \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix} &= \begin{pmatrix} 4 \\ 0 \\ -1 \end{pmatrix} + \mu \begin{pmatrix} -6 \\ -1 \\ 0 \end{pmatrix} \\ \Rightarrow \lambda \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix} + \mu \begin{pmatrix} 6 \\ 1 \\ 0 \end{pmatrix} &= \begin{pmatrix} 4 \\ 0 \\ -1 \end{pmatrix} - \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ -2 \\ -2 \end{pmatrix}. \end{aligned}$$

In Matrizenschreibweise lautet dieses lineare Gleichungssystem

$$\left(\begin{array}{cc|c} 1 & 6 & 1 \\ 2 & 1 & -2 \\ -1 & 0 & -2 \end{array} \right) \hookrightarrow \left(\begin{array}{cc|c} 1 & 6 & 1 \\ 0 & -11 & -4 \\ 0 & 6 & -1 \end{array} \right).$$

Aus der letzten Zeile folgt $6\mu = -1 \Rightarrow \mu = -\frac{1}{6}$, und aus der vorletzten Zeile folgt $-11\mu = -4 \Rightarrow \mu = \frac{4}{11}$. Dies ist ein Widerspruch! Also lässt sich die Vektorgleichung **nicht** lösen und es gibt keinen Schnittpunkt von g_1 mit g_2 . $\Rightarrow g_1$ und g_2 sind windschief. $\square$

2.3.3 Abstandsberechnung zu Geraden

Der **Abstand** eines Punktes Q von einer Geraden

$$g : \vec{x} = \vec{r}(P_1) + \lambda \vec{a}$$

ist gegeben durch die Höhe d des Parallelogramms, welches durch die Vektoren $\vec{a}$ und $\overrightarrow{P_1 Q}$ aufgespannt wird (siehe Abb. 2.14). Die Parallelogrammfläche A ist nach Definition des Kreuzproduktes $A = \left| \vec{a} \times \overrightarrow{P_1 Q} \right| = |\vec{a}| \cdot d$. Nach d aufgelöst folgt:

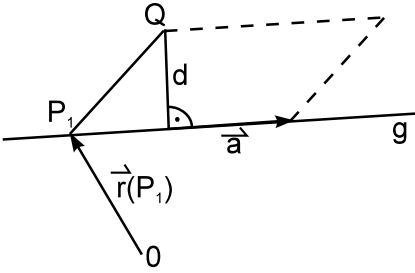


Abb. 2.14. Abstand des Punktes Q zur Geraden g

Der Abstand eines Punktes Q von einer Geraden

$$g : \vec{x} = \vec{r}(P_1) + \lambda \vec{a}$$

ist gegeben durch

$$d = \frac{|\vec{a} \times \overrightarrow{P_1 Q}|}{|\vec{a}|}.$$

Für $d = 0$ liegt der Punkt Q auf der Geraden! Aus aus obiger Formel ergibt sich auch der **Abstand zweier paralleler Geraden**

$$g_1 : \vec{x} = \vec{r}(P_1) + \lambda \vec{a} \quad \text{und} \quad g_2 : \vec{x} = \vec{r}(P_2) + \mu \vec{b},$$

indem man einen beliebigen Punkt auf der Geraden g_2 wählt, z.B. den Punkt P_2 , und den Abstand dieses Punktes zur Geraden g_1 bestimmt. Für $d = 0$ sind die Geraden zusammenfallend!

Um den Abstand zweier windschiefer Geraden

$$g_1 : \vec{x} = \vec{r}(P_1) + \lambda \vec{a} \quad \text{und} \quad g_2 : \vec{x} = \vec{r}(P_2) + \mu \vec{b}$$

zu berechnen, bestimmen wir den Vektor $\vec{L} = \vec{a} \times \vec{b}$.

$\vec{L}$ steht senkrecht auf $\vec{a}$ und auf $\vec{b}$. Für $\vec{L} = 0$ sind $\vec{a}$ und $\vec{b}$ parallel, für $\vec{L} \neq 0$ gehen wir zum Einheitsvektor

$$\vec{l} = \frac{1}{|\vec{L}|} \vec{L}$$

über. Der Abstand von g_1 und g_2 ist gegeben durch die Projektion von $\overrightarrow{P_1 P_2}$ auf $\vec{l}$, also $d = \vec{l} \cdot \overrightarrow{P_1 P_2}$:

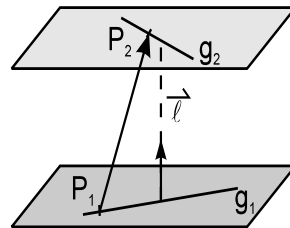


Abb. 2.15. Windschiefe Geraden

$$d = \frac{|\vec{a} \times \vec{b}|}{|\vec{a} \times \vec{b}|} \cdot \overrightarrow{P_1 P_2}$$

Abstand zweier windschiefer Geraden

$$g_1 : \vec{x} = \vec{r}(P_1) + \lambda \vec{a} \quad \text{und}$$

$$g_2 : \vec{x} = \vec{r}(P_2) + \mu \vec{b}.$$

Ist der Abstand $d = 0$, so schneiden sich die Geraden und der Schnittwinkel ergibt sich durch den Winkel, den die beiden Richtungsvektoren $\vec{a}$ und $\vec{b}$ miteinander einschließen:

$$\cos \varphi = \frac{|\vec{a} \cdot \vec{b}|}{|\vec{a}| |\vec{b}|}$$

Schnittwinkel zweier sich schneidender Geraden.

Beispiel 2.16. Gesucht ist der Abstand der beiden parallelen Geraden

$$g_1 : \vec{x} = \begin{pmatrix} 1 \\ 1 \\ 4 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \quad \text{und} \quad g_2 : \vec{x} = \begin{pmatrix} 3 \\ 0 \\ 1 \end{pmatrix} + \mu \begin{pmatrix} 2 \\ 2 \\ 2 \end{pmatrix}.$$

$$\text{Wegen } \vec{a} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \text{ ist } |\vec{a}| = \sqrt{3} \text{ und } \vec{a} \times \overrightarrow{P_1 P_2} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \times \begin{pmatrix} 2 \\ -1 \\ -3 \end{pmatrix} = \begin{pmatrix} 2 \\ 5 \\ -3 \end{pmatrix}.$$

$$\text{Also ist } d = \frac{|\vec{a} \times \overrightarrow{P_1 P_2}|}{|\vec{a}|} = \frac{\sqrt{38}}{\sqrt{3}} = 3,559. \quad \square$$

Beispiel 2.17. Gesucht ist der Schnittpunkt S und der Schnittwinkel φ der Geraden

$$g_1 : \vec{x} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} \quad \text{und} \quad g_2 : \vec{x} = \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}.$$

Um den Schnittpunkt S zu bestimmen, setzen wir $\vec{x}_{g_1} = \vec{x}_{g_2}$:

$$\begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}$$

$$\hookrightarrow \lambda \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 1 \\ -2 \end{pmatrix} = \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix} - \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}.$$

In der Matrixschreibweise folgt weiter

$$\left(\begin{array}{cc|c} 2 & -1 & 1 \\ 1 & 1 & -1 \\ 1 & -2 & 2 \end{array} \right) \hookrightarrow \left(\begin{array}{cc|c} 2 & -1 & 1 \\ 0 & -3 & 3 \\ 0 & 3 & -3 \end{array} \right).$$

Sowohl aus der letzten als auch vorletzten Zeile folgt $\mu = -1$, d.h. das LGS ist lösbar. Aus der ersten Zeile berechnet man

$$2\lambda - (-1) = 1 \quad \Rightarrow \quad \lambda = 0.$$

Somit ist $\vec{r}(S) = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + 0 \cdot \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ und die Koordinaten von S sind $S = (1, 1, 0)$. Der Schnittwinkel ist

$$\varphi = \arccos \left(\frac{\vec{a} \cdot \vec{b}}{|\vec{a}| |\vec{b}|} \right) = \arccos \frac{3}{\sqrt{6}\sqrt{6}} = 60^\circ,$$

denn $\vec{a} \cdot \vec{b} = \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} = 2 - 1 + 2 = 3$ und $|\vec{a}| = \sqrt{6} = |\vec{b}|$. $\square$

► 2.3.4 Vektorielle Darstellung von Ebenen

Eine Ebene E ist eindeutig bestimmt durch die Angabe dreier Punkte $P_1 = (x_1, y_1, z_1)$, $P_2 = (x_2, y_2, z_2)$ und $P_3 = (x_3, y_3, z_3)$, die nicht auf einer Geraden liegen. Dann legen diese 3 Punkte zwei Richtungsvektoren $\vec{a} = \overrightarrow{P_1 P_2}$ und $\vec{b} = \overrightarrow{P_1 P_3}$ und einen Ortsvektor $\vec{r}(P_1)$ fest. Jeder Punkt $P = (x, y, z)$ der Ebene entspricht einem Vektor $\vec{x} = \vec{r}(P)$ mit

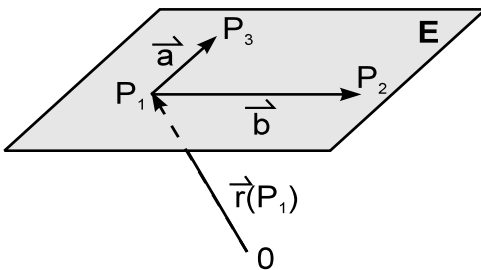


Abb. 2.16. Vektorielle Darstellung einer Ebene E

$$E: \vec{r}(P) = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b} \quad (\lambda, \mu \in \mathbb{R})$$

(Punkt-Richtungsform der Ebene).

Wenn man die Richtungsvektoren $\vec{a} = \overrightarrow{P_1 P_2}$ und $\vec{b} = \overrightarrow{P_1 P_3}$ durch die jeweiligen Ortsvektoren $\vec{r}(P_2) - \vec{r}(P_1)$ bzw. $\vec{r}(P_3) - \vec{r}(P_1)$ ersetzt, erhält man die Dreipunkteform der Ebene:

$$E : \vec{r}(P) = \vec{r}(P_1) + \lambda(\vec{r}(P_2) - \vec{r}(P_1)) + \mu(\vec{r}(P_3) - \vec{r}(P_1)) \quad (\lambda, \mu \in \mathbb{R})$$

(Dreipunkteform der Ebene).

Beispiel 2.18. Gegeben sind die Punkte $P_1 = (5, 2, 1)$, $P_2 = (4, 0, -4)$ und $P_3 = (1, 1, 1)$. Dann lautet die Ebenengleichung durch die 3 Punkte

$$\begin{aligned} E : \vec{r}(P) &= \begin{pmatrix} 5 \\ 2 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 4-5 \\ 0-2 \\ -4-1 \end{pmatrix} + \mu \begin{pmatrix} 1-5 \\ 1-2 \\ 1-1 \end{pmatrix} \\ &= \begin{pmatrix} 5 \\ 2 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} -1 \\ -2 \\ -5 \end{pmatrix} + \mu \begin{pmatrix} -4 \\ -1 \\ 0 \end{pmatrix}. \end{aligned} \quad \square$$

⊗ Hesse-Normalform

Eine alternative Darstellung der Ebenengleichung erhält man, wenn die Punkt-Richtungsform der Ebene mit dem zu E senkrecht stehenden Normalenvektor skalarmultipliziert wird. $\vec{N} := \vec{a} \times \vec{b}$ steht senkrecht auf $\vec{a}$ und auf $\vec{b}$ und der Normalen-Einheitsvektor lautet

$$\vec{n} := \frac{1}{|\vec{N}|} \vec{N} = \frac{1}{|\vec{a} \times \vec{b}|} \vec{a} \times \vec{b}.$$

Dieser hat die Eigenschaft $|\vec{n}| = 1$, $\vec{a} \cdot \vec{n} = 0$ und $\vec{b} \cdot \vec{n} = 0$. Somit folgt für jeden Punkt P auf der Ebene E

$$\vec{r}(P) \cdot \vec{n} = \vec{r}(P_1) \cdot \vec{n} \quad \text{oder} \quad (\vec{r}(P) - \vec{r}(P_1)) \cdot \vec{n} = 0.$$

Setzt man $\vec{r}(P) = \begin{pmatrix} x \\ y \\ z \end{pmatrix}$ folgt:

$$\begin{pmatrix} x - x_1 \\ y - y_1 \\ z - z_1 \end{pmatrix} \cdot \vec{n} = 0$$

Hesse-Normalform der Ebene, wenn P_1 ein Punkt der Ebene und $\vec{n}$ der Normalen-Einheitsvektor zur Ebene.

Setzt man den Richtungsvektor $\vec{n} = \begin{pmatrix} n_1 \\ n_2 \\ n_3 \end{pmatrix}$ ein, erhält man durch Ausmultiplizieren des Skalarprodukts

$$n_1(x - x_1) + n_2(y - y_1) + n_3(z - z_1) = 0.$$

Folglich bilden alle Lösungen einer linearen Gleichung

$$Ax + By + Cz = D$$

eine Ebene im $\mathbb{R}^3$.

Beispiel 2.19. Gesucht sind die Darstellungsformen der Ebene E , die durch den Punkt $P = (2, -5, 3)$ geht und senkrecht zum Normalenvektor $\vec{N} = \begin{pmatrix} 4 \\ 2 \\ 5 \end{pmatrix}$ steht: Aus der Hesse-Normalform folgt

$$\begin{aligned} \vec{N} \cdot (\vec{r}(P) - \vec{r}(P_1)) &= \begin{pmatrix} 4 \\ 2 \\ 5 \end{pmatrix} \cdot \begin{pmatrix} x - 2 \\ y + 5 \\ z - 3 \end{pmatrix} = 4(x - 2) + 2(y + 5) + 5(z - 3) = 0 \\ \Rightarrow \quad 4x + 2y + 5z &= 13. \quad (*) \end{aligned}$$

Der Übergang zur Punkt-Richtungsform der Ebene erhält man, indem bei diesem linearen Gleichungssystem $z = \mu$ (beliebig) und $y = \lambda$ (beliebig) gewählt werden. Aus der Gleichung (*) folgt dann

$$4x + 2\lambda + 5\mu = 13 \quad \Rightarrow \quad x = \frac{13}{4} - \frac{1}{2}\lambda - \frac{5}{4}\mu.$$

Somit ist die Lösung der Gleichung (*) die Punkt-Richtungsform der Ebene

$$\vec{x} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} \frac{13}{4} \\ 0 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -\frac{1}{2} \\ 1 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -\frac{5}{4} \\ 0 \\ 1 \end{pmatrix}.$$

Der Normalenvektor $\vec{N} = \vec{a} \times \vec{b} = \begin{pmatrix} -\frac{1}{2} \\ 1 \\ 0 \end{pmatrix} \times \begin{pmatrix} -\frac{5}{4} \\ 0 \\ 1 \end{pmatrix} = \frac{1}{4} \begin{pmatrix} 4 \\ 2 \\ 5 \end{pmatrix}$ wird normiert durch $\vec{n} = \frac{1}{|\vec{N}|} \vec{N}$. Somit ist eine andere Hesse-Normalform gegeben durch

$$\frac{1}{12\sqrt{5}} \begin{pmatrix} x - \frac{13}{4} \\ y - 0 \\ z - 0 \end{pmatrix} \cdot \begin{pmatrix} 4 \\ 2 \\ 5 \end{pmatrix} = 0. \quad \square$$

2.3.5 Lage zweier Ebenen zueinander

Um die **Lage zweier Ebenen**

$$E_1: \vec{x} = \vec{r}(P_1) + \lambda \vec{a}_1 + \mu \vec{b}_1 \quad \text{und} \quad E_2: \vec{x} = \vec{r}(P_2) + \tau \vec{a}_2 + \sigma \vec{b}_2$$

zu bestimmen, genügt es die Vektorgleichung $\vec{x}_{E_1} = \vec{x}_{E_2}$ zu lösen:

$$\vec{r}(P_1) + \lambda \vec{a}_1 + \mu \vec{b}_1 = \vec{r}(P_2) + \tau \vec{a}_2 + \sigma \vec{b}_2.$$

- (1) Ist das Gleichungssystem nicht lösbar, dann sind die Ebenen E_1 und E_2 **parallel** und **fallen nicht zusammen** ($E_1 \parallel E_2, E_1 \neq E_2$) (siehe Abb. 2.17 (a)).
- (2) Ist das Gleichungssystem mit einem Parameter lösbar, dann **schneiden sich die Ebenen** E_1 und E_2 in einer Geraden ($E_1 \nparallel E_2, E_1 \cap E_2 = g$) (siehe Abb. 2.17 (b)).
- (3) Ist das Gleichungssystem mit zwei Parameter lösbar, dann **fallen sie zusammen** ($E_1 = E_2$).

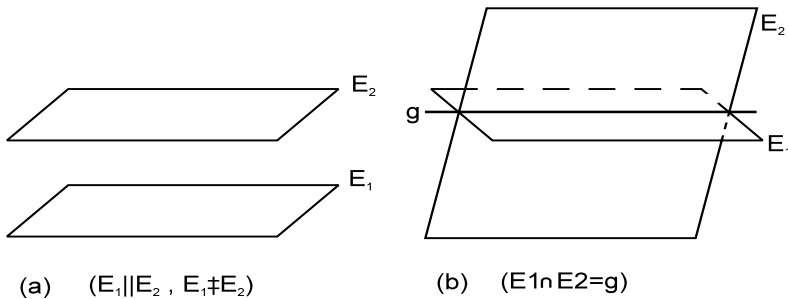


Abb. 2.17. Lage zweier Ebenen E_1 und E_2 zueinander

Beispiel 2.20 (Musterbeispiel).

Man bestimme die Lage der Ebenen E_1 und E_2 zueinander, wenn

$$E_1: \vec{x} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -1 \\ 2 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix};$$

$$E_2: \vec{x} = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \tau \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \sigma \begin{pmatrix} 0 \\ 4 \\ 0 \end{pmatrix}.$$

Mit $\vec{x}_{E_1} = \vec{x}_{E_2}$ folgt

$$\begin{aligned} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -1 \\ 2 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} &= \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \tau \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \sigma \begin{pmatrix} 0 \\ 4 \\ 0 \end{pmatrix} \\ &\hookrightarrow \lambda \begin{pmatrix} -1 \\ 2 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} + \tau \begin{pmatrix} -1 \\ -2 \\ -1 \end{pmatrix} + \sigma \begin{pmatrix} 0 \\ -4 \\ 0 \end{pmatrix} = \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}. \end{aligned}$$

In der Matrixdarstellung lautet das lineare Gleichungssystem für $\lambda, \mu, \tau, \sigma$

$$\begin{aligned} \left(\begin{array}{cccc|c} -1 & -1 & -1 & 0 & -1 \\ 2 & 0 & -2 & -4 & 1 \\ 0 & 1 & -1 & 0 & 0 \end{array} \right) &\hookrightarrow \left(\begin{array}{cccc|c} -1 & -1 & -1 & 0 & -1 \\ 0 & -2 & -4 & -4 & -1 \\ 0 & 1 & -1 & 0 & 0 \end{array} \right) \\ &\hookrightarrow \left(\begin{array}{cccc|c} -1 & -1 & -1 & 0 & -1 \\ 0 & -2 & -4 & -4 & -1 \\ 0 & 0 & -6 & -4 & -1 \end{array} \right). \end{aligned}$$

Es ist demnach $\sigma = t$ (beliebig) und $-6\tau - 4t = -1 \Rightarrow \tau = \frac{1}{6} - \frac{2}{3}t$. Die Lösung des LGS besitzt einen freien Parameter t ; also schneiden sich die Ebenen E_1 und E_2 in einer Geraden g . Die Darstellung der Geradengleichung erhalten wir, indem wir $\sigma = t$ und $\tau = \frac{1}{6} - \frac{2}{3}t$ in die Definitionsgleichung für E_2 einsetzen:

$$\begin{aligned} g: \vec{x} &= \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \left(\frac{1}{6} - \frac{2}{3}t\right) \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + t \begin{pmatrix} 0 \\ 4 \\ 0 \end{pmatrix} \\ &= \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \frac{1}{6} \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + t \left(-\frac{2}{3} \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \begin{pmatrix} 0 \\ 4 \\ 0 \end{pmatrix} \right) \\ g: \vec{x} &= \begin{pmatrix} \frac{1}{6} \\ \frac{4}{3} \\ \frac{1}{6} \end{pmatrix} + t \begin{pmatrix} -\frac{2}{3} \\ \frac{8}{3} \\ -\frac{2}{3} \end{pmatrix}. \end{aligned}$$

□

⑤ Um die Lage einer Ebene

$$E: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b}$$

und einer Geraden

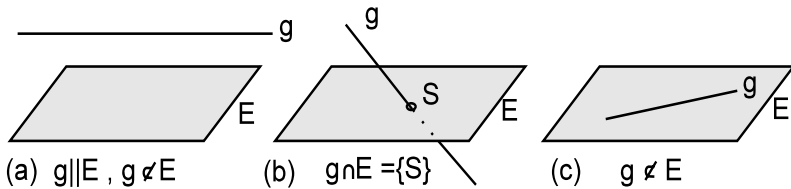
$$g: \vec{x} = \vec{r}(P_2) + \tau \vec{c}$$

zu bestimmen, lösen wir die Vektorgleichung $\vec{x}_E = \vec{x}_g$:

$$\vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b} = \vec{r}(P_2) + \tau \vec{c}.$$

Dies ist ein LGS für die Unbekannten λ, μ, τ . Es gilt

- (1) Ist das Gleichungssystem nicht lösbar, dann ist g **parallel** zu E aber nicht in E enthalten ($g \parallel E, g \not\subset E$) (siehe Abb. 2.18 (a)).
- (2) Ist das Gleichungssystem eindeutig lösbar, dann **schneiden** sich die Ebene E und die Gerade g in einem **Punkt** S ($g \cap E = \{S\}$) (siehe Abb. 2.18 (b)).
- (3) Ist das Gleichungssystem mit einem Parameter lösbar, dann liegt die **Gerade** g **in der Ebene** E ($g \subset E$) (siehe Abb. 2.18 (c)).

Abb. 2.18. Lage einer Geraden g zu einer Ebene E

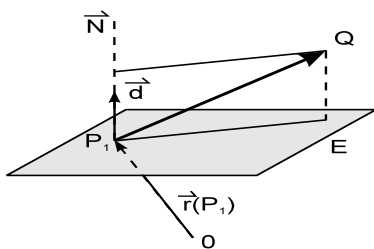
2.3.6 Abstandsberechnung zu Ebenen

Der **Abstand eines Punktes** Q **von einer Ebene** $E: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b}$ ist gegeben durch die Projektion des Vektors $\overrightarrow{P_1 Q}$ auf die Normale $\vec{N}$ der Ebene: $\vec{d} = \frac{\vec{N} \cdot \overrightarrow{P_1 Q}}{N^2} \cdot \vec{N}$. Also ist der Abstand $d = |\vec{d}| = \frac{|\vec{N} \cdot \overrightarrow{P_1 Q}|}{N^2} \cdot |\vec{N}| = \frac{|\vec{N} \cdot \overrightarrow{P_1 Q}|}{N}$:

$$d = \frac{|\vec{N} \cdot \overrightarrow{P_1 Q}|}{|\vec{N}|}$$

Abstand des Punktes Q von der Ebene

$E: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b}$
mit Normalenvektor $\vec{N} = \vec{a} \times \vec{b}$.

Abb. 2.19. Abstand Punkt zu Ebene E

Ist $d = 0$, so liegt der Punkt Q in der Ebene E .

Der **Abstand einer zu E parallelen Geraden** $g: \vec{x} = \vec{r}(P_2) + \tau \vec{c}$ ergibt sich direkt aus obiger Formel, indem man sich einen beliebigen Punkt auf der Geraden z.B. P_2 wählt und den Abstand dieses Punktes zur Ebene bestimmt:

$$d = \frac{|\vec{N} \cdot \overrightarrow{P_1 P_2}|}{|\vec{N}|} \quad \text{Abstand der Geraden } g: \vec{x} = \vec{r}(P_2) + \tau \vec{c} \\ \text{von der Ebene } E: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b} \\ \text{mit Normalenvektor } \vec{N} = \vec{a} \times \vec{b}.$$

Für $d = 0$ liegt die Gerade in der Ebene E .

Der **Abstand einer zu E parallelen Ebene E_2** : $\vec{x} = \vec{r}(P_2) + \tau \vec{c} + \sigma \vec{d}$ ergibt sich ebenfalls direkt aus obiger Formel, indem man einen Punkt der Ebene E_2 wählt (z.B. P_2) und einsetzt:

$$d = \frac{|\vec{N} \cdot \overrightarrow{P_1 P_2}|}{|\vec{N}|} \quad \text{Abstand der Ebene } E_2: \vec{x} = \vec{r}(P_2) + \tau \vec{c} + \sigma \vec{d} \\ \text{von der Ebene } E: \vec{x} = \vec{r}(P_1) + \lambda \vec{a} + \mu \vec{b} \\ \text{mit Normalenvektor } \vec{N} = \vec{a} \times \vec{b}.$$

Ist $d = 0$, so fallen beide Ebenen zusammen.

Beispiel 2.21. Gesucht ist der Abstand des Punktes $Q = (3, 1, 5)$ von der Ebene $E: \vec{x} = \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} 0 \\ 2 \\ 2 \end{pmatrix}$: Wegen $\vec{N} := \vec{a} \times \vec{b} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} \times \begin{pmatrix} 0 \\ 2 \\ 2 \end{pmatrix} = \begin{pmatrix} 2 \\ -2 \\ 2 \end{pmatrix}$ ist $|\vec{N}| = \sqrt{12}$ und $\overrightarrow{P_1 Q} = \begin{pmatrix} 2 \\ -1 \\ 5 \end{pmatrix}$.

$$\Rightarrow d = \frac{|\vec{N} \cdot \overrightarrow{P_1 Q}|}{|\vec{N}|} = \frac{1}{\sqrt{12}} \left| \begin{pmatrix} 2 \\ -2 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ -1 \\ 5 \end{pmatrix} \right| = \frac{16}{\sqrt{12}} = \frac{8}{3}\sqrt{3}. \quad \square$$

2.3.7 Berechnung des Schnittes einer Geraden mit einer Ebene

Um den Schnittpunkt einer Geraden

$$g: \vec{x} = \vec{r}(P_2) + \lambda \vec{a}$$

mit einer Ebene E zu bestimmen, gehen wir davon aus, dass die Ebene E in der Hesse-Normalform

$$E: \vec{N} \cdot (\vec{x} - \vec{r}(P_1)) = 0$$

gegeben ist, d.h. P_1 ein Punkt auf der Ebene und $\vec{N}$ ein Normalenvektor ist. Wir gehen davon aus, dass die Gerade nicht parallel zur Ebene liegt.

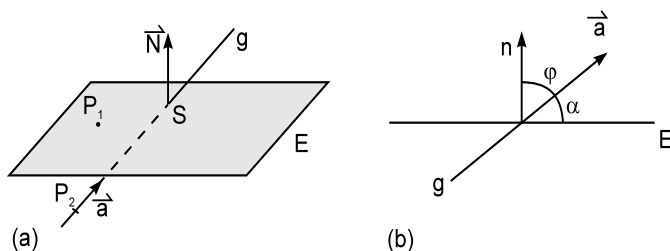


Abb. 2.20. Schnittpunkt und -winkel einer Geraden mit einer Ebene

Der **Schnittpunkt** S hat die Eigenschaft, dass $\vec{x}_g = \vec{x}_E$, d.h. wir setzen $\vec{x}_g$ in die Ebenengleichung ein:

$$\vec{N} \cdot (\vec{r}(P_2) + \lambda \vec{a} - \vec{r}(P_1)) = \vec{N} \cdot (\vec{r}(P_2) - \vec{r}(P_1)) + \lambda \vec{N} \cdot \vec{a} = 0.$$

Da die Gerade nicht parallel zur Ebene liegt, ist $\vec{N} \cdot \vec{a} \neq 0$, so dass mit $\overrightarrow{P_1 P_2} = \vec{r}(P_2) - \vec{r}(P_1)$ folgt

$$\lambda = -\frac{\vec{N} \cdot \overrightarrow{P_1 P_2}}{\vec{N} \cdot \vec{a}}.$$

Setzt man dieses λ in die Geradengleichung ein, folgt für den Schnittpunkt S

$$\vec{r}(S) = \vec{r}(P_2) - \frac{\vec{N} \cdot \overrightarrow{P_1 P_2}}{\vec{N} \cdot \vec{a}} \cdot \vec{a} \quad \text{Ortsvektor zum Schnittpunkt } S.$$

Für den Winkel φ zwischen der Normalen der Ebene und der Geraden gilt

$$\cos \varphi = \frac{\vec{N} \cdot \vec{a}}{|\vec{N}| |\vec{a}|}.$$

φ ist der Ergänzungswinkel zu α :

$$\varphi = 90^\circ \pm \alpha,$$

je nachdem wie die Richtung des Normalenvektors ist. Daher ist

$$\cos \varphi = \cos(90^\circ \pm \alpha) = \mp \sin \alpha \quad \text{und}$$

$$\sin \alpha = \pm \frac{\vec{N} \cdot \vec{a}}{|\vec{N}| |\vec{a}|}$$

Schnittwinkel zwischen der Geraden
 $g: \vec{x} = \vec{r}(P_2) + \lambda \vec{a}$
und der Ebene $E: \vec{N}(\vec{x} - \vec{r}(P_1)) = 0$.

Beispiel 2.22. Gesucht ist der Schnittpunkt S und der Schnittwinkel α der Geraden $g: \vec{x} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix}$ mit $E: \vec{x} = \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} + \lambda \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} 4 \\ 2 \\ -1 \end{pmatrix}$.

Aus den Richtungsvektoren der Ebene $\vec{b}_1 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ und $\vec{b}_2 = \begin{pmatrix} 4 \\ 2 \\ -1 \end{pmatrix}$ erhält man den Normalenvektor $\vec{N} = \vec{b}_1 \times \vec{b}_2 = \begin{pmatrix} -1 \\ 0 \\ -4 \end{pmatrix}$. Damit ist

$$\vec{N} \cdot \overrightarrow{P_1 P_2} = \begin{pmatrix} -1 \\ 0 \\ -4 \end{pmatrix} \cdot \left(\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} - \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} \right) = 1$$

$$\vec{N} \cdot \vec{a} = \begin{pmatrix} -1 \\ 0 \\ -4 \end{pmatrix} \cdot \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix} = -10.$$

Der Schnittpunkt S berechnet sich aus

$$\vec{r}(S) = \vec{r}(P_2) - \frac{\vec{N} \cdot \overrightarrow{P_1 P_2}}{\vec{N} \cdot \vec{a}} \cdot \vec{a} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \frac{1}{10} \begin{pmatrix} 2 \\ 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 1,2 \\ 2 \\ 3,2 \end{pmatrix}.$$

Der Schnittwinkel α folgt aus

$$\sin \alpha = \frac{|\vec{N} \cdot \vec{a}|}{|\vec{N}| \cdot |\vec{a}|} = \frac{-10}{\sqrt{17} \sqrt{8}} \Rightarrow \alpha = -59,04^\circ. \quad \square$$

Bemerkung: Der **Schnittwinkel zweier sich schneidenden Ebenen**

$$E_1: \vec{N}_1(\vec{r}(P) - \vec{r}(P_1)) = 0 \quad \text{und} \quad E_2: \vec{N}_2(\vec{r}(P) - \vec{r}(P_2)) = 0$$

ist der gleiche Schnittwinkel wie der Schnittwinkel ihrer Normalenvektoren.

$$\text{Daher gilt } \cos \varphi = \frac{|\vec{N}_1 \cdot \vec{N}_2|}{|\vec{N}_1| |\vec{N}_2|}.$$



Visualisierung mit MAPLE: Auf der CD-Rom befinden sich MAPLE-Prozeduren zur **Darstellung sowohl von Geraden als auch Ebenen**.

Ebenfalls auf der CD befindet sich die Prozedur **geomnet**. Sie bestimmt die Lage zweier Objekte (Punkte, Geraden oder Ebenen) zueinander.

2.4 Vektorräume

Wir übertragen die Definition von Vektoren aus dem $\mathbb{R}^3$ in den $\mathbb{R}^n$ und erweitern die elementare Rechenoperationen (Addition und S-Multiplikation) auf diese Vektoren. Aus den allgemeinen Eigenschaften kommt man auf den Begriff des Vektorraums, der eine wichtige Bedeutung bei der Beschreibung linearer physikalischer Systeme spielt. Insbesondere bei der formalen Behandlung z.B. von linearen Gleichungssystemen oder der linearen Differenzialgleichungen und -systemen benötigt man diesen Formalismus. Zentral sind bei der Charakterisierung der Vektorräume die Begriffe der Linearkombination, der linearen Unabhängigkeit und der Basis als minimales Erzeugendensystem.

2.4.1 Vektorrechnung im $\mathbb{R}^n$

Wir übertragen den Begriff des Vektors von $\mathbb{R}^3$ in den $\mathbb{R}^n$:

Definition: Die Menge aller n -Tupel reeller Zahlen heißt $\mathbb{R}^n$:

$$\mathbb{R}^n := \left\{ \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} : x_1 \in \mathbb{R}, x_2 \in \mathbb{R}, \dots, x_n \in \mathbb{R} \right\}.$$

Analog dem Koordinatensystem in der Ebene bzw. im Raum wird das Koordinatensystem im $\mathbb{R}^n$ durch n aufeinander senkrecht stehenden Einheitsvektoren

$$\vec{e}_1 := \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \vec{e}_2 := \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}, \dots, \vec{e}_n := \begin{pmatrix} 0 \\ \vdots \\ 1 \\ 0 \end{pmatrix}$$

gebildet. Jeder Vektor $\vec{a} \in \mathbb{R}^n$ lässt sich durch die Angabe seiner Komponenten beschreiben:

$$\vec{a} = a_1 \vec{e}_1 + a_2 \vec{e}_2 + \dots + a_n \vec{e}_n = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix}.$$

Es übertragen sich dann die Begriffe des Betrags, der Gleichheit von Vektoren, der Multiplikation mit einem Skalar, die Addition, das Skalarprodukt, der Orthogonalität usw. auf den $\mathbb{R}^n$.

Addition und S-Multiplikation

Für zwei Vektoren $\vec{a} = \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix}$, $\vec{b} = \begin{pmatrix} b_1 \\ \vdots \\ b_n \end{pmatrix}$ und einem Skalar $\lambda \in \mathbb{R}$ setzt man

$$\vec{a} + \vec{b} := \begin{pmatrix} a_1 + b_1 \\ a_2 + b_2 \\ \vdots \\ a_n + b_n \end{pmatrix} \quad \lambda \cdot \vec{a} := \begin{pmatrix} \lambda a_1 \\ \lambda a_2 \\ \vdots \\ \lambda a_n \end{pmatrix}$$

(Addition) (S-Multiplikation)

Sowohl die Addition als auch die S-Multiplikation werden komponentenweise ausgeführt. Durch die Addition und S-Multiplikation hat man zwei *Operationen*

$$\begin{aligned} + : \mathbb{R}^n \times \mathbb{R}^n &\rightarrow \mathbb{R}^n & \text{mit} & \quad (\vec{a}, \vec{b}) \mapsto \vec{a} + \vec{b} \\ \cdot : \mathbb{R} \times \mathbb{R}^n &\rightarrow \mathbb{R}^n & \text{mit} & \quad (\lambda, \vec{a}) \mapsto \lambda \cdot \vec{a} \end{aligned}$$

festgelegt. Formal unterscheiden sich die Vektoraddition $+$ und die S-Multiplikation $\cdot$ dadurch, dass zum einen zwei Vektoren und zum anderen eine skalare Zahl mit einem Vektor verknüpft werden. Man bezeichnet daher „ $+$ “ als *innere Verknüpfung* und „ $\cdot$ “ als *äußere Verknüpfung*.

Da sowohl die Addition als auch die S-Multiplikation komponentenweise erklärt sind, übertragen sich die folgenden Rechengesetze von den reellen Zahlen auf diese Vektoren.

Es gelten die **Rechengesetze der Addition**

- | | | |
|-------------------|---|-------------------------|
| (A ₁) | $\vec{a} + (\vec{b} + \vec{c}) = (\vec{a} + \vec{b}) + \vec{c}$ | <i>Assoziativgesetz</i> |
| (A ₂) | $\vec{a} + \vec{b} = \vec{b} + \vec{a}$ | <i>Kommutativgesetz</i> |
| (A ₃) | Der Nullvektor hat die Eigenschaft
$\vec{a} + \vec{0} = \vec{a}$ | <i>Nullvektor</i> |
| (A ₄) | Zu jedem Vektor $\vec{a}$ gibt es einen Vektor
$(-\vec{a})$ mit $\vec{a} + (-\vec{a}) = \vec{0}$ | <i>Negativer Vektor</i> |

Es gelten die **Rechengesetze der S-Multiplikation**:

(S_1)	$k \cdot (l \cdot \vec{a}) = (k \cdot l) \cdot \vec{a}$	Assoziativgesetz
(S_2)	$k \cdot (\vec{a} + \vec{b}) = k \vec{a} + k \vec{b}$	Distributivgesetz 1
(S_3)	$(k + l) \cdot \vec{a} = k \vec{a} + l \vec{a}$	Distributivgesetz 2
(S_4)	$1 \cdot \vec{a} = \vec{a}$	Gesetz der Eins

➤ 2.4.2 Vektorräume

Die Gesetzmäßigkeiten bezüglich der Addition und S-Multiplikation gelten nicht nur für n -Tupel, sondern auch für andere Objekte, die keine Veranschaulichung durch Pfeile zulassen (z.B. Funktionen). Um auch solche Objekte zu erfassen, führt man den Begriff des Vektorraums formal für alle Objekte ein, die zwei Verknüpfungen $+$ und $\cdot$ mit den angegebenen Rechengesetzen besitzen.

Definition: Eine Menge $\mathbb{W}$ bildet einen Vektorraum über $\mathbb{R}$, wenn folgende Axiome gelten:

(1) In $\mathbb{W}$ ist eine innere Verknüpfung „+“ erklärt,

$$+ : \mathbb{W} \times \mathbb{W} \rightarrow \mathbb{W} \quad \text{mit} \quad (\vec{a}, \vec{b}) \mapsto \vec{a} + \vec{b} \quad (\text{Addition}),$$

so dass $(\mathbb{W}, +)$ die Gesetze der Addition $(A_1) - (A_4)$ erfüllt.

(2) In $\mathbb{W}$ ist eine äußere Verknüpfung „ $\cdot$ “ erklärt,

$$\cdot : \mathbb{R} \times \mathbb{W} \rightarrow \mathbb{W} \quad \text{mit} \quad (\lambda, \vec{a}) \mapsto \lambda \cdot \vec{a} \quad (\text{S-Multiplikation}),$$

so dass $(\mathbb{W}, \cdot)$ die Gesetze der S-Multiplikation $(S_1) - (S_4)$ erfüllt.

Die Elemente eines Vektorraums bezeichnet man als **Vektoren**, auch wenn der Vektorraum nicht dem Anschauungsraum $\mathbb{R}^3$ entspricht. Hat man als Zahlenmenge nicht $\mathbb{R}$, sondern einen anderen Körper K , so spricht man von einem *Vektorraum über K* .

Beispiele 2.23:

- ① $\mathbb{R}^3$ ist ein Vektorraum bestehend aus allen 3-dimensionalen Pfeilen (den 3-Tupeln von reellen Zahlen).

② $\mathbb{R}^n = \left\{ \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} : x_i \in \mathbb{R} \quad (i = 1, \dots, n) \right\}$ ist ein Vektorraum, dessen Elemente die n -Tupel sind. $\mathbb{R}^n$ heißt auch der *arithmetische Vektorraum*.

③ Die Menge der auf dem Intervall $[a, b]$ definierten, reellwertigen Funktionen

$$F[a, b] := \{f : [a, b] \rightarrow \mathbb{R}\}$$

ist ein Vektorraum, wenn man für die Addition und S-Multiplikation definiert:

$$+ : F[a, b] \times F[a, b] \rightarrow F[a, b] \text{ mit } (f, g) \mapsto f + g \text{ und}$$

$$(f + g)(x) := f(x) + g(x).$$

$$\cdot : \mathbb{R} \times F[a, b] \rightarrow F[a, b] \text{ mit } (\lambda, f) \mapsto \lambda \cdot f \text{ und}$$

$$(\lambda f)(x) := \lambda \cdot f(x).$$

Die Rechengesetze übertragen sich aus dem Reellen. Die konstante Nullfunktion 0 mit $0(x) = 0$ für alle $x \in [a, b]$ bildet den Nullvektor.

④ Die Menge aller Polynomfunktionen vom Grad kleiner gleich n

$$P[n] := \{f : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f(x) = \sum_{i=0}^n a_i x^i, a_i \in \mathbb{R}\}$$

bildet einen Vektorraum, wenn „+“ und „ $\cdot$ “ wie unter ③ erklärt sind. $\square$

Beispiel 2.24. Die Lösungsmenge eines homogenen, linearen Gleichungssystems bildet einen Vektorraum, wenn „+“ und „ $\cdot$ “ wie unter 2.23 ② definiert sind. Als **Zahlenbeispiel** betrachten wir das LGS

$$\begin{aligned} -3x_1 - 5x_2 + 2x_3 &= 0 \\ 4x_1 - x_2 + 3x_3 &= 0 \end{aligned}$$

Die Lösung erhalten wir mit dem Gauß-Algorithmus

$$\left(\begin{array}{ccc|c} -3 & -5 & 2 & 0 \\ 4 & -1 & 3 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -3 & -5 & 2 & 0 \\ 0 & -23 & 17 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

$$\hookrightarrow x_3 = 23\lambda, x_2 = 17\lambda, x_1 = -13\lambda$$

$$\Rightarrow \mathbb{L} = \left\{ \vec{x} \in \mathbb{R}^3 : \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \lambda \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix}; \quad \lambda \in \mathbb{R} \text{ beliebig} \right\}.$$

Wir zeigen nur die *Abgeschlossenheit* bezüglich „+“ und „·“, d.h. dass die Addition zweier Vektoren von $\mathbb{L}$ wieder einen Vektor aus $\mathbb{L}$ ergibt und $r \cdot \vec{x} \in \mathbb{L}$, wenn $\vec{x} \in \mathbb{L}$ ist. Die Gesetze der Addition $(A_1) - (A_4)$ sowie die der S-Multiplikation $(S_1) - (S_4)$ übertragen sich dann von $\mathbb{R}^3$ auf $\mathbb{L}$.

$$\begin{aligned}\vec{x}_1 + \vec{x}_2 &= \lambda_1 \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} + \lambda_2 \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} = (\lambda_1 + \lambda_2) \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} \in \mathbb{L} \\ r \cdot \vec{x}_1 &= r \cdot \left(\lambda_1 \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} \right) = (r \cdot \lambda_1) \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} = \lambda \begin{pmatrix} -13 \\ 17 \\ 23 \end{pmatrix} \in \mathbb{L}.\end{aligned}$$

Wir haben damit nachgerechnet, dass mit zwei Lösungen $\vec{x}_1$ und $\vec{x}_2$ des LGS auch die Summe $\vec{x}_1 + \vec{x}_2$ eine Lösung ist und jedes Vielfache einer Lösung ebenfalls das Gleichungssystem erfüllt. Physikalisch bedeutet diese Eigenschaft, dass das *Superpositionsgesetz* gültig ist. Da die Teilmenge $\mathbb{L} \subset \mathbb{R}^3$ selbst wieder einen Vektorraum darstellt, nennt man $\mathbb{L}$ einen *Untervektorraum* von $\mathbb{R}^3$. Diesen Begriff verwendet man immer, wenn eine Teilmenge eines Vektorraums wieder einen Vektorraum bildet: $\square$

Definition: Sei $(\mathbb{W}, +, \cdot)$ ein Vektorraum. Eine Teilmenge $U \subset \mathbb{W}$ heißt **Untervektorraum**, wenn U bezüglich den linearen Operationen „+“ und „·“ einen Vektorraum bildet.

Beispiele 2.25:

- ① $\mathbb{W} = \mathbb{R}^n, U = \{\text{Lösungen eines homogenen, linearen Gleichungssystems}\}.$
- ② $\mathbb{W} = \{\text{Menge aller reellwertigen Funktionen}\}$
 $U = \{\text{Menge der auf dem Intervall } [a, b] \text{ stetigen, reellwertigen Funktionen}\}.$
- ③ $\mathbb{W} = \{\text{Menge aller Polynomfunktionen vom Grade } \leq n\},$
 $U = \{\text{Menge aller Polynomfunktionen vom Grade } \leq n \text{ und } f(1) = 0\}.$ $\square$

Für eine Teilmenge $U \neq \emptyset$ eines Vektorraums $\mathbb{W}$ muss man nicht mehr die Rechengesetze nachprüfen, um zu zeigen, dass U selbst wieder einen Vektorraum darstellt. Die Rechengesetze übertragen sich von $\mathbb{W}$ auf U , wenn U bezüglich „+“ und „·“ abgeschlossen ist:

Satz: (Untervektorraum-Kriterium). Eine nichtleere Teilmenge $U \subset \mathbb{W}$ ist genau dann ein Untervektorraum, wenn U bezüglich den linearen Operationen „+“ und „·“ abgeschlossen ist.

Folglich sind für eine Teilmenge U eines Vektorraums $\mathbb{W}$ drei Eigenschaften zu prüfen, um zu zeigen, dass U selbst wieder einen Vektorraum darstellt:

$$U \subset \mathbb{W} \text{ ist Vektorraum} \Leftrightarrow \begin{array}{ll} UV1: & \vec{0} \in U. \\ UV2: & \vec{a}, \vec{b} \in U \Rightarrow \vec{a} + \vec{b} \in U. \\ UV3: & \vec{a} \in U, \lambda \in \mathbb{R} \Rightarrow \lambda \cdot \vec{a} \in U. \end{array}$$

Beispiele 2.26:

- ① Die Menge $\{0\}$ ist ein Untervektorraum jedes Vektorraums. Es ist der kleinstmögliche Vektorraum.
- ② $U = \{\text{Menge der reellwertigen Funktionen } f : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f(1) = 0\}$,
 $\mathbb{W} = \{\text{Menge der reellwertigen Funktionen } f : \mathbb{R} \rightarrow \mathbb{R}\}$. $U \subset \mathbb{W}$ ist ein Untervektorraum, denn

UV1: Die Nullabbildung $0 : \mathbb{R} \rightarrow \mathbb{R}$ mit $0(x) = 0$ hat die Eigenschaft $0(1) = 0$. Also ist $\{0\} \in U$.

UV2: Mit $f_1, f_2 \in U$ ist $f_1(1) = f_2(1) = 0$ und damit auch
 $(f_1 + f_2)(1) = f_1(1) + f_2(1) = 0 \Rightarrow f_1 + f_2 \in U$.

UV3: Mit $f \in U$ ist $f(1) = 0$ und damit auch
 $(\lambda f)(1) = \lambda \cdot f(1) = \lambda \cdot 0 = 0 \Rightarrow \lambda f \in U$.

- ③ $U = \{\text{Menge der reellwertigen Funktionen } f : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f(1) = 1\}$,
 $\mathbb{W} = \{\text{Menge der reellwertigen Funktionen } f : \mathbb{R} \rightarrow \mathbb{R}\}$. $U \subset \mathbb{W}$ ist **kein** Untervektorraum!

Denn z.B. aus $f_1, f_2 \in U$, d.h. $f_1(1) = f_2(1) = 1$, folgt $(f_1 + f_2)(1) = f_1(1) + f_2(1) = 2 \Rightarrow f_1 + f_2 \notin U$. Der Nullvektor ist ebenfalls nicht in U enthalten. $\square$

► 2.4.3 Linearkombination und Erzeugnis

Im Folgenden gehen wir davon aus, dass $(\mathbb{W}, +, \cdot)$ ein Vektorraum ist und $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ Vektoren aus diesem Vektorraum sind.

Definition: Ein Vektor $\vec{b}$ der Form

$$\vec{b} = \lambda_1 \vec{a}_1 + \lambda_2 \vec{a}_2 + \dots + \lambda_n \vec{a}_n$$

heißt **Linearkombination** der Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$.

Beispiele 2.27:

① $\begin{pmatrix} 9 \\ 7 \\ 9 \end{pmatrix} = 5 \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + 3 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + 4 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$. Der Vektor $\begin{pmatrix} 9 \\ 7 \\ 9 \end{pmatrix}$ ist also eine Linearkombination der Vektoren $\begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$.

② $\begin{pmatrix} 9 \\ 7 \\ 9 \end{pmatrix} = 9 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + 7 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + 9 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = 9 \vec{e}_1 + 7 \vec{e}_2 + 9 \vec{e}_3$; d.h. der Vektor $\begin{pmatrix} 9 \\ 7 \\ 9 \end{pmatrix}$ ist auch eine Linearkombination der Vektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3$. $\square$

Wie wir aus den $\mathbb{R}^3$ wissen, lässt sich jeder Vektor $\vec{b} \in \mathbb{R}^3$ als Linearkombination von $\vec{e}_1, \vec{e}_2, \vec{e}_3$ darstellen. Man definiert verallgemeinernd

Definition: Die Menge M aller Linearkombinationen der Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ heißt **Erzeugnis** von $\vec{a}_1, \dots, \vec{a}_n$. Man schreibt hierfür

$$\begin{aligned} M &= \{ \text{Menge aller Linearkombinationen von } \vec{a}_1, \dots, \vec{a}_n \} \\ &= [\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n] \\ &= \{ \vec{b} : \vec{b} = \lambda_1 \vec{a}_1 + \lambda_2 \vec{a}_2 + \dots + \lambda_n \vec{a}_n \quad (\lambda_i \in \mathbb{R}) \}. \end{aligned}$$

Beispiel 2.28. Ist $\vec{b} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$ im Erzeugnis von $\vec{a}_1 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix},$

$\vec{a}_3 = \begin{pmatrix} 1 \\ 4 \\ 5 \end{pmatrix}$? Gesucht sind $\lambda_1, \lambda_2, \lambda_3 \in \mathbb{R}$, so dass

$$\vec{b} = \lambda_1 \vec{a}_1 + \lambda_2 \vec{a}_2 + \lambda_3 \vec{a}_3,$$

denn dann ist $\vec{b}$ eine Linearkombination von $\vec{a}_1, \vec{a}_2, \vec{a}_3$.

Ansatz:
$$\begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} = \lambda_1 \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 4 \\ 5 \end{pmatrix}.$$

In der Komponentendarstellung entspricht dies dem inhomogenen, linearen Gleichungssystem

$$\begin{aligned} 0 \cdot \lambda_1 + 1 \cdot \lambda_2 + 1 \cdot \lambda_3 &= 1 \\ 1 \cdot \lambda_1 + 2 \cdot \lambda_2 + 4 \cdot \lambda_3 &= 0 \\ 1 \cdot \lambda_1 + 3 \cdot \lambda_2 + 5 \cdot \lambda_3 &= 1. \end{aligned}$$

Zur Lösung des LGS verwenden wir den Gauß-Algorithmus in Matrizenschreibweise

$$\left(\begin{array}{ccc|c} 0 & 1 & 1 & 1 \\ 1 & 2 & 4 & 0 \\ 1 & 3 & 5 & 1 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 2 & 4 & 0 \\ 0 & 1 & 1 & 1 \\ 0 & 1 & 1 & 1 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 2 & 4 & 0 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

$$\Rightarrow \lambda_3 = t \quad (\text{beliebig}) \quad \hookrightarrow \quad \lambda_2 = 1 - t \quad \hookrightarrow \quad \lambda_1 = -2 - 2t.$$

Setzt man z.B. $t = 1$, so folgt $\lambda_3 = 1$, $\lambda_2 = 0$, $\lambda_1 = -4$ und

$$\vec{b} = -4 \cdot \vec{a}_1 + 0 \cdot \vec{a}_2 + 1 \cdot \vec{a}_3.$$

Damit ist $\vec{b}$ im Erzeugnis von $[\vec{a}_1, \vec{a}_2, \vec{a}_3]$. □

Satz: Die Vektorgleichung $\vec{b} = \lambda_1 \vec{a}_1 + \lambda_2 \vec{a}_2 + \dots + \lambda_n \vec{a}_n$ ist genau dann lösbar, wenn $\vec{b} \in [\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n]$.

Da $M = [\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n]$ die Menge aller Linearkombinationen ist, stellt $M \subset \mathbb{W}$ selbst wieder einen Vektorraum dar, was man mit dem Untervektorraum-Kriterium sofort nachprüfen kann:

Satz: $M = [\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n]$ ist ein Vektorraum.

Im Fall, dass M schon den ganzen Vektorraum $\mathbb{W}$ aufspannt, nennt man M ein *Erzeugendensystem*:

Definition: Eine Teilmenge von Vektoren $\{\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n\} \subset \mathbb{W}$ heißt **Erzeugendensystem von $\mathbb{W}$** , wenn das Erzeugnis von $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ mit $\mathbb{W}$ zusammenfällt: $[\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n] = \mathbb{W}$.

Beispiele 2.29:

- ① $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$ ist ein Erzeugendensystem von $\mathbb{R}^3$, denn jeder Vektor $\vec{x}$ lässt sich als Linearkombination von $\vec{e}_1, \vec{e}_2, \vec{e}_3$ darstellen: $\vec{x} = x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3$.

- ② $\left\{ \vec{e}_1, \vec{e}_2, \vec{e}_3, \vec{d} := \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$ ist ebenfalls ein Erzeugendensystem von $\mathbb{R}^3$, denn jeder Vektor $\vec{x}$ lässt sich als Linearkombination dieser 4 Vektoren darstellen:

$$\vec{x} = x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3 + 0 \cdot \vec{d}$$

$$\text{oder } \vec{x} = (x_1 - 1) \vec{e}_1 + (x_2 - 1) \vec{e}_2 + (x_3 - 1) \vec{e}_3 + 1 \cdot \vec{d}. \quad \square$$

Sowohl $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$ als auch $\{\vec{e}_1, \vec{e}_2, \vec{e}_3, \vec{d}\}$ bilden ein Erzeugendensystem von $\mathbb{R}^3$. Gesucht ist ein Kriterium, um ein kleinstmögliches Erzeugendensystem zu charakterisieren. Dazu benötigt man den Begriff der *linearen Unabhängigkeit*.

2.4.4 Lineare Abhängigkeit und Unabhängigkeit

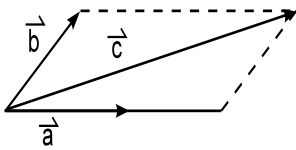


Abb. 2.21. Lineare Abhängigkeit

Beispiel 2.30. Gegeben sind die Vektoren

$$\vec{a} = \begin{pmatrix} 2 \\ -3 \\ 5 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 4 \\ 3 \\ -2 \end{pmatrix}, \quad \vec{c} = \begin{pmatrix} 8 \\ -3 \\ 8 \end{pmatrix}. \quad \text{Dann ist}$$

$$\vec{c} = 2 \cdot \vec{a} + 1 \cdot \vec{b}.$$

$\vec{c}$ lässt sich also als Linearkombination der Vektoren $\vec{a}$ und $\vec{b}$ darstellen. Man nennt $\vec{c}$ daher *linear abhängig* von $\vec{a}$ und $\vec{b}$. Stellt man die Gleichung $\vec{c} = 2 \cdot \vec{a} + 1 \cdot \vec{b}$ um, so ist

$$2 \cdot \vec{a} + 1 \cdot \vec{b} - 1 \cdot \vec{c} = \vec{0}.$$

D.h. der Nullvektor lässt sich darstellen als Linearkombination von $\vec{a}$, $\vec{b}$, $\vec{c}$ mit von Null verschiedenen Koeffizienten. Wir verallgemeinern: $\square$

Definition: Die Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{W}$ heißen **linear abhängig**, wenn in der Gleichung

$$k_1 \vec{a}_1 + k_2 \vec{a}_2 + \dots + k_n \vec{a}_n = \vec{0} \quad (*)$$

mindestens ein $k_i \neq 0$ ist.

Denn dann lässt sich die Gleichung $(*)$ nach $\vec{a}_i$ auflösen:

$$\vec{a}_i = -\frac{1}{k_i} (k_1 \vec{a}_1 + \dots + k_{i-1} \vec{a}_{i-1} + k_{i+1} \vec{a}_{i+1} + \dots + k_n \vec{a}_n),$$

und $\vec{a}_i$ ist durch die restlichen Vektoren darstellbar. Lässt sich die Gleichung $(*)$ niemals nach einem Vektor $\vec{a}_i$ ($i \in \{1, \dots, n\}$) auflösen, dann nennt man die Vektoren *linear unabhängig*. Dies ist genau dann der Fall, wenn $k_1 = k_2 = \dots = k_n = 0$ die einzigen Lösungen sind.

Definition: Die Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{W}$ heißen **linear unabhängig**, wenn gilt:

$$k_1 \vec{a}_1 + k_2 \vec{a}_2 + \dots + k_n \vec{a}_n = \vec{0} \quad \Rightarrow \quad k_1 = 0, k_2 = 0, \dots, k_n = 0.$$

Beispiele 2.31:

- ① Die Vektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3$ sind linear unabhängig: Denn der Ansatz

$$k_1 \vec{e}_1 + k_2 \vec{e}_2 + k_3 \vec{e}_3 = \vec{0}$$

liefert

$$k_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + k_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + k_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

$$\Rightarrow \left(\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{array} \right) \Rightarrow k_3 = 0, k_2 = 0, k_1 = 0.$$

D.h. aus dem Ansatz $k_1 \vec{e}_1 + k_2 \vec{e}_2 + k_3 \vec{e}_3 = \vec{0}$ folgt für die Koeffizienten $k_1 = k_2 = k_3 = 0$. Damit sind nach obiger Definition die Vektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3$ linear **un**abhängig.

- ② Die Vektoren $\vec{a}_1 = \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 3 \\ 0 \\ 1 \end{pmatrix}$ sind linear unabhängig: Aus $k_1 \vec{a}_1 + k_2 \vec{a}_2 + k_3 \vec{a}_3 = \vec{0}$ folgt

$$k_1 \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} + k_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + k_3 \begin{pmatrix} 3 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

In Matrizenschreibweise erhalten wir

$$\left(\begin{array}{ccc|c} 2 & 0 & 3 & 0 \\ -1 & 1 & 0 & 0 \\ 3 & 0 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 2 & 3 & 0 \\ 0 & 3 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 2 & 3 & 0 \\ 0 & 0 & 7 & 0 \end{array} \right).$$

Durch Rückwärtsauflösen ist

$$7 \cdot k_3 = 0 \Rightarrow \boxed{k_3 = 0}; \quad 2 \cdot k_2 = 0 \Rightarrow \boxed{k_2 = 0}; \quad -1 \cdot k_1 = 0 \Rightarrow \boxed{k_1 = 0}.$$

D.h. aus $k_1 \vec{a}_1 + k_2 \vec{a}_2 + k_3 \vec{a}_3 = \vec{0}$ folgt $k_1 = k_2 = k_3 = 0$. Damit sind die Vektoren $\vec{a}_1, \vec{a}_2, \vec{a}_3$ linear **un**abhängig.

- ③ Die Vektoren $\vec{a}_1 = \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 2 \\ 0 \\ 3 \end{pmatrix}$ sind linear abhängig: Denn aus $k_1 \vec{a}_1 + k_2 \vec{a}_2 + k_3 \vec{a}_3 = \vec{0}$, folgt in der Matrizen-

schreibweise

$$\left(\begin{array}{ccc|c} 2 & 0 & 2 & 0 \\ -1 & 1 & 0 & 0 \\ 3 & 0 & 3 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 2 & 2 & 0 \\ 0 & 3 & 3 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right)$$

$$\Rightarrow k_3 = \lambda \text{ (beliebig); } k_2 = -\lambda; k_1 = -\lambda.$$

Z.B. für $\lambda = 1$ ist

$$(-1) \vec{a}_1 + (-1) \vec{a}_2 + 1 \vec{a}_3 = 0 \Rightarrow \vec{a}_1 = -\vec{a}_2 + \vec{a}_3.$$

Das Gleichungssystem ist also **nicht** nur durch $k_1 = k_2 = k_3 = 0$ lösbar, und damit die Vektorgleichung nach dem Vektor $\vec{a}_1$ auflösbar. Die Vektoren $\vec{a}_1, \vec{a}_2, \vec{a}_3$; sind linear **abhängig**.

- ④ Die Vektoren $f_i(x) = x^i$ ($i = 0, 1, 2, \dots, n$) sind linear **unabhängig** im Vektorraum $P[n] := \{\text{Menge aller Polynomfunktionen vom Grade } \leq n\}$. Denn

$$k_0 f_0 + k_1 f_1 + \dots + k_n f_n = \vec{0}$$

bedeutet $k_0 f_0(x) + k_1 f_1(x) + \dots + k_n f_n(x) = 0(x) = 0$ für alle $x \in \mathbb{R}$, d.h.

$$k_0 + k_1 x + \dots + k_n x^n = 0$$

für alle $x \in \mathbb{R}$. Durch Einsetzen von $x = 0$ folgt $k_0 = 0$. Anschließend kann auf der linken Seite ein x ausgeklammert werden. Wieder durch Einsetzen von $x = 0$ folgt $k_1 = 0$. Insgesamt erhält man so $k_0 = k_1 = \dots = k_n = 0$. $\square$

In den Beispielen ① - ④ zeigt sich, dass die Vektorgleichung

$$k_1 \vec{a}_1 + k_2 \vec{a}_2 + \dots + k_n \vec{a}_n = \vec{0}$$

entweder eindeutig lösbar ist; dann ist $k_1 = k_2 = \dots = k_n = 0$ und die Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ sind linear unabhängig. Oder die Vektorgleichung ist nicht eindeutig lösbar, dann sind die Vektoren linear abhängig. Es gilt folgende Charakterisierung von linear unabhängigen Vektoren:

Satz: Für die Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{W}$ sind äquivalent:

- (1) $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ sind linear unabhängig.
- (2) Jeder Vektor $\vec{b} \in [\vec{a}_1, \dots, \vec{a}_n]$ lässt sich eindeutig aus den Vektoren $\vec{a}_1, \dots, \vec{a}_n$ linear kombinieren.

2.4.5 Basis und Dimension

Einer der für die Beschreibung von Vektorräumen wichtigsten Begriffe ist der der *Basis*.

Definition: Eine Menge **linear unabhängiger** Vektoren, die den gesamten Vektorraum erzeugen, nennt man eine **Basis** des Vektorraums.

Also:

$$\begin{array}{|l} \vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{W} \\ \text{ist \textbf{Basis} von } \mathbb{W} \end{array} \Leftrightarrow \begin{array}{|l} (B_1) \vec{a}_1, \dots, \vec{a}_n \text{ sind linear unabhängig.} \\ (B_2) [\vec{a}_1, \dots, \vec{a}_n] = \mathbb{W}. \end{array}$$

Eine Basis ist die kleinste Menge von Vektoren, welche den Vektorraum erzeugt, und sie ist gleichzeitig die größte Menge von linear unabhängigen Vektoren aus $\mathbb{W}$; denn für Basen sind die folgenden Aussagen gleichbedeutend:

Satz: (Charakterisierung von Basen). Für eine Teilmenge $B := \{\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n\} \subset \mathbb{W}$ ist gleichbedeutend:

- (1) B ist eine **Basis** von $\mathbb{W}$.
- (2) B ist eine unverlängerbare, linear unabhängige Teilmenge von $\mathbb{W}$.
- (3) B ist ein unverkürzbares Erzeugendensystem von $\mathbb{W}$.

Beispiele 2.32:

- ① $(\vec{e}_1, \vec{e}_2, \vec{e}_3)$ ist eine Basis von $\mathbb{R}^3$: Denn $\vec{e}_1, \vec{e}_2, \vec{e}_3$ sind linear unabhängig und jedes $\vec{x} \in \mathbb{R}^3$ lässt sich als Linearkombination von $\vec{e}_1, \vec{e}_2, \vec{e}_3$ darstellen:

$$\vec{x} = x_1 \vec{e}_1 + x_2 \vec{e}_2 + x_3 \vec{e}_3.$$

- ② $(\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n)$ ist eine Basis von $\mathbb{R}^n$.

- ③ $(1, x, x^2, x^3, x^4, \dots, x^n)$ ist eine Basis des Vektorraums der Polynomfunktionen vom Grad $\leq n$.

④ $\vec{a}_1 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ ist eine Basis von $\mathbb{R}^3$:

Wir prüfen die Bedingungen (B_1) und (B_2) nach.

(B_1) : Aus $k_1 \vec{a}_1 + k_2 \vec{a}_2 + k_3 \vec{a}_3 = \vec{0}$ folgt

$$\left(\begin{array}{ccc|c} 0 & 1 & 1 & 0 \\ 1 & 2 & 1 & 0 \\ 1 & 3 & 0 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 3 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 1 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 3 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & -2 & 0 \end{array} \right).$$

Dieses LGS hat als eindeutige Lösung $k_1 = k_2 = k_3 = 0$ und damit sind $\vec{a}_1, \vec{a}_2, \vec{a}_3$ linear unabhängig.

(B_2) : Ist $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} \in \mathbb{R}^3$ beliebig, dann müssen λ, μ, τ gefunden werden,

so dass

$$\lambda \vec{a}_1 + \mu \vec{a}_2 + \tau \vec{a}_3 = \vec{b}.$$

In Matrizenschreibweise ist das LGS für die Unbekannten λ, μ und τ :

$$\begin{aligned} \left(\begin{array}{ccc|c} 0 & 1 & 1 & b_1 \\ 1 & 2 & 1 & b_2 \\ 1 & 3 & 0 & b_3 \end{array} \right) &\hookrightarrow \left(\begin{array}{ccc|c} 1 & 3 & 0 & b_3 \\ 0 & 1 & -1 & b_3 - b_2 \\ 0 & 1 & 1 & b_1 \end{array} \right) \\ &\hookrightarrow \left(\begin{array}{ccc|c} 1 & 3 & 0 & b_3 \\ 0 & 1 & 1 & b_1 \\ 0 & 0 & -2 & b_3 - b_2 - b_1 \end{array} \right). \end{aligned}$$

Hieraus erhält man durch Rückwärtsauflösen die gesuchten Größen

$$\tau = -\frac{1}{2}(b_3 - b_2 - b_1); \mu = \frac{1}{2}(b_1 - b_2 + b_3); \lambda = \frac{3}{2}(-b_1 + b_2 - \frac{1}{3}b_3).$$

Damit gibt es zu jedem Vektor $\vec{b} \in \mathbb{R}^3$ Parameter $\lambda, \mu, \tau \in \mathbb{R}$, so dass

$$\lambda \vec{a}_1 + \mu \vec{a}_2 + \tau \vec{a}_3 = \vec{b} \quad \Rightarrow \quad [\vec{a}_1, \vec{a}_2, \vec{a}_3] = \mathbb{W}.$$

Aus (B_1) und (B_2) folgt, dass $(\vec{a}_1, \vec{a}_2, \vec{a}_3)$ eine Basis von $\mathbb{R}^3$ ist.

⑤ Die Vektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3, \vec{a}_4$ mit $\vec{a}_4 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$ bilden **keine** Basis von $\mathbb{R}^3$, da sie linear abhängig sind: $\vec{a}_4 = \vec{e}_1 + \vec{e}_2 + \vec{e}_3$.

⑥ Die Vektoren $\vec{a}_1 = \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix}$, $\vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix}$ bilden **keine** Basis von $\mathbb{R}^3$.

$\vec{a}_1$, $\vec{a}_2$ sind zwar linear unabhängig, aber nicht jeder Vektor $\vec{b} \in \mathbb{R}^3$ lässt sich als Linearkombination von $\vec{a}_1$, $\vec{a}_2$ darstellen. Denn aus

$$\lambda \vec{a}_1 + \mu \vec{a}_2 = \vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} \text{ folgt } \left(\begin{array}{cc|c} 2 & 0 & b_1 \\ 0 & 1 & b_2 \\ 1 & 2 & b_3 \end{array} \right) \\ \hookrightarrow \left(\begin{array}{cc|c} 2 & 0 & b_1 \\ 0 & -4 & -2b_3 + b_1 \\ 0 & 1 & b_2 \end{array} \right).$$

Aus der letzten Zeile des Gleichungssystems folgt $1 \cdot \mu = b_2$ aus der zweit-
letzten Zeile $\mu = \frac{1}{4}(b_1 - 2b_3)$. Damit das LGS lösbar ist, muss für den
Vektor $\vec{b}$ gelten:

$$b_2 = \frac{1}{4}(b_1 - 2b_3).$$

Dies ist aber nicht für alle Vektoren $\vec{b} \in \mathbb{R}^3$ erfüllt und damit ist $(\vec{a}_1, \vec{a}_2)$
kein Erzeugendensystem von $\mathbb{R}^3$. (B_2) ist also nicht erfüllt. $\square$

In einem Vektorraum kann es beliebig viele Basen geben. Hat man jedoch eine
endliche Basis $\vec{a}_1, \dots, \vec{a}_n$ gefunden, so besteht jede andere Basis ebenfalls
aus genau n Vektoren. Die maximale Anzahl linear unabhängiger Vektoren ist
charakteristisch für einen Vektorraum:

Definition: Sei $\mathbb{W}$ ein Vektorraum. Besteht eine Basis aus n Vektoren,
so heißt die Zahl n die **Dimension** des Vektorraums $\mathbb{W}$.

Bezeichnung: $\dim(\mathbb{W}) = n$.

Man beachte, dass zwar jeder Vektorraum eine Basis besitzt; diese muss je-
doch nicht notgedrungen aus endlich vielen Vektoren bestehen. In der Regel
betrachten wir hier nur endlich-dimensionale Vektorräume. Für diese endlich-
dimensionalen Vektorräume gilt

Satz: Sei $\mathbb{W}$ ein n -dimensionaler Vektorraum. Dann gilt für n Vektoren $\vec{a}_1, \dots, \vec{a}_n \in \mathbb{W}$ die Aussage:

$\vec{a}_1, \dots, \vec{a}_n$
sind linear unabhängig.

$\Leftrightarrow$

$(\vec{a}_1, \dots, \vec{a}_n)$ ist eine
Basis von $\mathbb{W}$.

Beispiele 2.33:

- ① $\vec{a}_1 = \begin{pmatrix} 5 \\ 2 \end{pmatrix}$ und $\vec{a}_2 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ bilden eine Basis von $\mathbb{R}^2$, denn $\vec{a}_1, \vec{a}_2$ sind linear unabhängig:

$$k_1 \vec{a}_1 + k_2 \vec{a}_2 = \vec{0}$$

$$\hookrightarrow \left(\begin{array}{cc|c} 5 & 2 & 0 \\ 2 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cc|c} 5 & 2 & 0 \\ 0 & -1 & 0 \end{array} \right) \Rightarrow k_2 = k_1 = 0.$$

- ② $\vec{a}_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 0 \end{pmatrix}$, $\vec{a}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{a}_3 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \end{pmatrix}$, $\vec{a}_4 = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 1 \end{pmatrix}$ bilden eine Basis von $\mathbb{R}^4$, denn $\vec{a}_1, \vec{a}_2, \vec{a}_3, \vec{a}_4$ sind linear unabhängig: Setzen wir die Vektoren in die Vektorgleichung

$$k_1 \vec{a}_1 + k_2 \vec{a}_2 + k_3 \vec{a}_3 + k_4 \vec{a}_4 = \vec{0}$$

ein, erhalten wir das zugehörige LGS, welches wir mit dem Gauß-Algorithmus lösen:

$$\left(\begin{array}{cccc|c} 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 0 & 0 & -1 & 1 & 0 \\ 0 & 0 & 1 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} 1 & 1 & 1 & 1 & 0 \\ 0 & 1 & 0 & 1 & 0 \\ 0 & 0 & -1 & 1 & 0 \\ 0 & 0 & 0 & 2 & 0 \end{array} \right)$$

$$\Rightarrow k_1 = k_2 = k_3 = k_4 = 0.$$

- ③ $P[5] := \{f : \mathbb{R} \rightarrow \mathbb{R} : f(x) = a_0 + a_1x + \dots + a_5x^5\}$ ist ein 6-dimensionaler Vektorraum: $(x^0, x^1, x^2, \dots, x^5)$ sind linear unabhängige Funktionen und jedes $f \in P[5]$ lässt sich als Linearkombination dieser Funktionen darstellen $\Rightarrow (x^0, x^1, \dots, x^5)$ ist eine Basis von $P[5] \Rightarrow \dim P[5] = 6$. $\square$

MAPLE-Worksheets zu Kapitel 2



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 2 mit MAPLE zur Verfügung.

- Darstellung von Vektoren im $\mathbb{R}^2$
- Darstellung von Vektoren im $\mathbb{R}^3$
- Vektorrechnung mit MAPLE
- Graphische Darstellung von Geraden und Ebenen im $\mathbb{R}^3$
- Punkte, Geraden und Ebenen mit MAPLE
- Die Prozedur **geomet**
- Zusammenstellung der MAPLE-Befehle
- MAPLE-Lösungen zu den Aufgaben

2.5 Aufgaben zur Vektorrechnung

2.1 Gegeben sind die Vektoren $\vec{a} = \begin{pmatrix} 2 \\ 3 \\ -1 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 0 \\ -2 \\ 4 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} -5 \\ 3 \\ 1 \end{pmatrix}$.

Man berechne die folgenden Vektoren und ihre Beträge

a) $\vec{s}_1 = 3\vec{a} - 4\vec{b} + \vec{c}$ b) $\vec{s}_2 = -3(5\vec{b} + \vec{c}) + 5(-\vec{a} + 3\vec{b})$
 c) $\vec{s}_3 = 3(\vec{a} - 2\vec{b}) + 5\vec{c}$ d) $\vec{s}_4 = 3(\vec{a} \cdot \vec{b})\vec{c} - 5(\vec{b} \cdot \vec{c})\vec{a}$

2.2 Welche Gegenkraft $\vec{F}$ hebt die vier Einzelkräfte $\vec{F}_1$, $\vec{F}_2$, $\vec{F}_3$, $\vec{F}_4$ in ihrer Gesamtkraft auf? (Krafteinheit 1N.)

$$\vec{F}_1 = \begin{pmatrix} 200 \\ 110 \\ -50 \end{pmatrix}; \vec{F}_2 = \begin{pmatrix} -10 \\ 30 \\ -40 \end{pmatrix}; \vec{F}_3 = \begin{pmatrix} 40 \\ 85 \\ 120 \end{pmatrix}; \vec{F}_4 = -\begin{pmatrix} 30 \\ 50 \\ 40 \end{pmatrix}.$$

2.3 Normieren Sie die folgenden Vektoren:

$$\vec{a} = \begin{pmatrix} 2 \\ 3 \\ 1 \end{pmatrix}, \quad \vec{b} = 3\vec{e}_1 - 5\vec{e}_2 + 2\vec{e}_3, \quad \vec{c} = \begin{pmatrix} -1 \\ 0 \\ -1 \end{pmatrix}.$$

2.4 Wie lautet der Einheitsvektor $\vec{e}$, der die zum Vektor $\vec{a} = -\begin{pmatrix} 4 \\ 3 \\ 0 \end{pmatrix}$ entgegengesetzte Richtung hat?

2.5 Bestimmen Sie die Koordinaten des Punktes Q , der vom Punkte $P = (1, -2, 3)$ in Richtung des Vektors $\vec{a} = \begin{pmatrix} -2 \\ -1 \\ -1 \end{pmatrix}$ 10 Längeneinheiten entfernt ist.

2.6 Bestimmen Sie die Koordinaten der Mitte Q von $\overline{P_1 P_2}$ mit $P_1 = (2, 4, 3)$ und $P_2 = (-1, 3, 2)$.

2.7 Bilden Sie mit den Vektoren $\vec{a} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$; $\vec{b} = \begin{pmatrix} 2 \\ 1 \\ 2 \end{pmatrix}$; $\vec{c} = \begin{pmatrix} -4 \\ 2 \\ -2 \end{pmatrix}$ die

Skalarprodukte:

a) $\vec{a} \cdot \vec{b}$ b) $(\vec{a} - 3\vec{b}) \cdot 4\vec{c}$ c) $(\vec{a} + \vec{b}) \cdot (\vec{a} - \vec{c})$

2.8 Welchen Winkel schließen die Vektoren $\vec{a}$ und $\vec{b}$ ein?

a) $\vec{a} = \begin{pmatrix} 3 \\ 5 \\ 1 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 4 \\ 1 \\ 3 \end{pmatrix}$ b) $\vec{a} = \begin{pmatrix} 2 \\ -1 \\ 2 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} -10 \\ -1 \\ -10 \end{pmatrix}$

2.9 Zeigen Sie, dass die folgenden Vektoren $\vec{e}_1$, $\vec{e}_2$, $\vec{e}_3$ ein orthonormales System bilden; d.h. die Vektoren stehen paarweise aufeinander senkrecht und besitzen die Länge 1:

$$\vec{e}_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \quad \vec{e}_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}, \quad \vec{e}_3 = \begin{pmatrix} 0 \\ -1 \\ 0 \end{pmatrix}$$

2.10 Zeigen Sie: Die drei Vektoren

$$\vec{a} = \begin{pmatrix} 1 \\ 4 \\ -2 \end{pmatrix}, \vec{b} = \begin{pmatrix} -2 \\ 2 \\ 3 \end{pmatrix}, \vec{c} = \begin{pmatrix} -1 \\ 6 \\ 1 \end{pmatrix} \text{ bilden ein rechtwinkliges Dreieck.}$$

2.11 Bestimmen Sie den Betrag und die Winkel, die der Vektor $\vec{a}$ mit den Koordinatenachsen einschließt:

$$\text{a) } \vec{a} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \quad \text{b) } \vec{a} = \begin{pmatrix} 5 \\ -2 \\ 1 \end{pmatrix}$$

2.12 Durch die drei Punkte $A = (-1, 2, 4)$, $B = (5, 0, 0)$ und $C = (3, 4, -2)$ wird ein Dreieck festgelegt. Berechnen Sie die Länge der drei Seiten, die Winkel im Dreieck, sowie den Flächeninhalt.

2.13 Berechnen Sie die Komponente des Vektors $\vec{b}$ in Richtung des Vektors $\vec{a} = \begin{pmatrix} 2 \\ -2 \\ 1 \end{pmatrix}$ für a) $\vec{b} = \begin{pmatrix} 5 \\ 1 \\ 3 \end{pmatrix}$ b) $\vec{b} = \begin{pmatrix} -2 \\ 5 \\ 0 \end{pmatrix}$

2.14 Ein Vektor $\vec{a}$ ist durch den Betrag $|\vec{a}| = 10$ und $\alpha = 30^\circ$, $\beta = 60^\circ$, $90^\circ \leq \gamma \leq 180^\circ$ festgelegt. Wie lauten die Komponenten von $\vec{a}$?

2.15 Man bestimme die Richtungswinkel α, β, γ der Vektoren

$$\text{a) } \vec{a} = \begin{pmatrix} -1 \\ 1 \\ 4 \end{pmatrix} \quad \text{b) } \vec{a} = \begin{pmatrix} 4 \\ 2 \\ -3 \end{pmatrix}$$

2.16 Berechnen Sie für $\vec{a} = \begin{pmatrix} 4 \\ 2 \\ 1 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 2 \\ 3 \\ 3 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} 3 \\ -2 \\ 0 \end{pmatrix}$:

a) $\vec{a} \times \vec{b}$ b) $(\vec{a} - \vec{b}) \times (3\vec{c})$
c) $(-\vec{a} + 2\vec{c}) \times (-\vec{b})$ d) $(2\vec{a}) \times (-\vec{b} - \vec{c})$

2.17 An einem Verteilermast greifen 4 Kräfte an, die in einer Ebene liegen. Ermitteln Sie rechnerisch den Betrag und die Richtung der Resultierenden $\vec{F}_R = \vec{F}_1 + \vec{F}_2 + \vec{F}_3 + \vec{F}_4$, wenn $|\vec{F}_1| = 380 \text{ N}$, $|\vec{F}_2| = 400 \text{ N}$, $|\vec{F}_4| = 440 \text{ N}$, und wenn der Winkel zwischen $\vec{F}_1$ und $\vec{F}_2$ $\alpha = 80^\circ$, der Winkel zwischen $\vec{F}_2$ und $\vec{F}_3$ $\beta = 120^\circ$, der Winkel zwischen $\vec{F}_3$ und $\vec{F}_4$ $\gamma = 70^\circ$ beträgt.

2.18 Gegeben sei ein Körper, der sich nur entlang der Richtung $\vec{a} = \begin{pmatrix} 2 \\ 1 \\ -2 \end{pmatrix}$ bewegen kann. Auf diesen Körper wirkt eine Kraft $\vec{F} = \begin{pmatrix} 20 \\ 20 \\ 10 \end{pmatrix} \text{ N}$.

- a) Wie groß ist der Betrag der Kraft $\vec{F}$?
b) Welche Winkel schließen der Kraftvektor und der Richtungsvektor ein?
c) Welche Kraft wirkt auf den Körper in Richtung $\vec{a}$?

- 2.19 Ein starrer Körper in Form einer Kreisscheibe ist um seine Symmetrieachse drehbar gelagert. Eine im Punkt P angreifende Kraft erzeugt ein Drehmoment

$$\vec{M} = \vec{r} \times \vec{F}. \text{ Seien } \vec{F} = \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} N \text{ und } \vec{r}(P) = \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} m.$$

- a) Welchen Winkel schließen $\vec{r}(P)$ und $\vec{F}$ ein?
 b) Man berechne das Drehmoment $\vec{M}$ und seinen Betrag.
 c) Welche Kraft $\vec{F}_r$ wirkt in Richtung $\vec{r}(P)$?

- 2.20 Gegeben sind die Punkte $A = (1, -1, 2)$, $B = (2, 1, 3)$, $C = (4, 0, 1)$. Unter der Einwirkung der konstanten Kraft $\vec{F} = (1, 1, 1)$ bewegt sich ein Massenpunkt m von A nach B . Wie groß ist die dabei verrichtete Arbeit (Krafteinheit $1 N$, Längeneinheit $1 m$), falls

- a) m sich auf kürzestem Weg von A nach B bewegt?
 b) m sich von A nach B längs der Strecken $\overline{AC}$ und $\overline{CB}$ bewegt?

- 2.21 Überprüfen Sie die Ergebnisse von den Aufgaben 2.1 - 2.20 mit MAPLE.

- 2.22 Wie lautet die Vektorgleichung der Geraden g durch den Punkt P parallel zum Vektor $\vec{a}$? Welche Punkte gehören zu den Parameterwerten $\lambda = 1$, $\lambda = 2$, $\lambda = -5$? $P = (4, 0, 3)$; $\vec{a} = \begin{pmatrix} -1 \\ 0 \\ -1 \end{pmatrix}$.

- 2.23 Man bestimme die Gleichung der Geraden g durch die Punkte $P_1 = (1, 3, -2)$ und $P_2 = (6, 5, 8)$.

- 2.24 Liegen die drei Punkte $P_1 = (3, 0, 4)$, $P_2 = (1, 1, 1)$ und $P_3 = (-1, 2, -2)$ auf einer Geraden?

- 2.25 Man berechne den Abstand des Punktes $Q = (4, 1, 1)$ von der Geraden g , die bestimmt ist durch den Punkt $P_1 = (4, 2, 3)$ und den Richtungsvektor $\vec{a} = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$.

- 2.26 Eine Gerade g verlaufe durch den Punkt $P = (5, 3, 1)$ parallel zu dem Vektor $\vec{a}$ mit den Richtungswinkeln $\alpha = 30^\circ$, $\beta = 90^\circ$ und γ mit $\cos \gamma < 0$. Wie lautet die Gleichung dieser Geraden?

- 2.27 Welche Lage besitzen die folgenden Geradenpaare g_1 , g_2 zueinander? Man bestimme gegebenenfalls Abstand, Schnittpunkt und Schnittwinkel.

- a) g_1 durch $P_1 = (3, 4, 6)$ und $P_2 = (-1, -2, 4)$

$$g_2 \text{ durch } P_3 = (3, 7, -2) \text{ und } P_4 = (5, 15, -6)$$

- b) g_1 durch $\vec{x} = \vec{r}_1 + \lambda \vec{a} = \begin{pmatrix} 5 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -2 \\ 1 \\ 3 \end{pmatrix}$

$$g_2 \text{ durch } \vec{x} = \vec{r}_2 + \lambda \vec{b} = \begin{pmatrix} 1 \\ 1 \\ 5 \end{pmatrix} + \lambda \begin{pmatrix} 6 \\ -3 \\ -9 \end{pmatrix}$$

c) g_1 durch $P_1 = (1, 2, 0)$ mit Richtungsvektor $\vec{a} = \begin{pmatrix} 2 \\ 0 \\ 5 \end{pmatrix}$

g_2 durch $P_2 = (6, 0, 13)$ mit Richtungsvektor $\vec{b} = \begin{pmatrix} 1 \\ -2 \\ 3 \end{pmatrix}$

- 2.28 Zeigen Sie, dass die beiden Geraden g_1 und g_2 windschief sind und berechnen Sie ihren Abstand:

$$g_1 : \vec{x} = \vec{r}_1 + \lambda \vec{a} = \begin{pmatrix} 1 \\ -2 \\ 3 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$$

$$g_2 : \vec{x} = \vec{r}_2 + \lambda \vec{b} = \begin{pmatrix} 3 \\ 3 \\ 3 \end{pmatrix} + \lambda \begin{pmatrix} 0 \\ 2 \\ 1 \end{pmatrix}$$

- 2.29 Wie lautet die Vektorgleichung der Ebene E , die den Punkt $P_1 = (3, 5, 1)$ enthält und parallel zu den Richtungsvektoren $\vec{a} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$ verläuft? Man bestimme den Normalenvektor $\vec{n}$ der Ebene. Welcher Punkt gehört zu dem Parameterpaar $\lambda = 1, \mu = 3$?

- 2.30 Man bestimme die Gleichung der Ebene E durch die Punkte $P_1 = (3, 1, 0)$; $P_2 = (-4, 1, 1)$; $P_3 = (5, 9, 3)$.

- 2.31 Liegen die vier Punkte $P_1 = (1, 1, 1)$; $P_2 = (3, 2, 0)$; $P_3 = (4, -1, 5)$ und $P_4 = (12, -4, 12)$ in einer Ebene?

- 2.32 Eine Ebene verläuft senkrecht zum Vektor $\vec{n} = \begin{pmatrix} 4 \\ 3 \\ 1 \end{pmatrix}$ und enthält den Punkt $A = (5, 8, 10)$. Man bestimme die Vektorgleichung dieser Ebene.

- 2.33 Welche Lage haben die Gerade g und Ebene E zueinander? Man bestimme gegebenenfalls Abstand, Schnittpunkt und Schnittwinkel.

a) g durch $P_1 = (5, 1, 2)$ mit Richtungsvektor $\vec{a} = \begin{pmatrix} 3 \\ 1 \\ 2 \end{pmatrix}$

E durch $P_0 = (2, 1, 8)$ mit Normalenvektor $\vec{n} = \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix}$

b) $g : \vec{x}_g = \vec{r}(P_1) + \lambda \begin{pmatrix} 2 \\ 5 \\ 1 \end{pmatrix} = \begin{pmatrix} 5 \\ 3 \\ 6 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 5 \\ 1 \end{pmatrix}$

$E : \vec{n}(\vec{r}(P) - \vec{r}(P_0)) = \begin{pmatrix} 3 \\ -1 \\ -1 \end{pmatrix} \begin{pmatrix} x-1 \\ y-1 \\ z-1 \end{pmatrix} = 0$

c) g durch $P_1 = (2, 0, 3)$ und $P_2 = (5, 6, 18)$

E durch $P_3 = (1, -2, -2)$, $P_4 = (0, -1, -1)$ und $P_5 = (-1, 0, -1)$

- 2.34 Man zeige die Parallelität der beiden Ebenen und berechne ihren Abstand

$$E_1 \text{ durch } P_1 = (3, 5, 6) \text{ mit Normalenvektor } \vec{n}_1 = \begin{pmatrix} 1 \\ 3 \\ -2 \end{pmatrix}$$

$$E_2 \text{ durch } P_2 = (1, 5, -2) \text{ mit Normalenvektor } \vec{n}_2 = \begin{pmatrix} -3 \\ -9 \\ 6 \end{pmatrix}.$$

2.35 Man bestimme Schnittgerade und Schnittwinkel der beiden Ebenen

$$E_1 : \vec{n}_1 \cdot (\vec{x}_E - \vec{r}(P_1)) = \begin{pmatrix} 3 \\ 1 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} x-2 \\ y-5 \\ z-6 \end{pmatrix} = 0$$

$$E_2 : \vec{n}_2 \cdot (\vec{x}_E - \vec{r}(P_2)) = \begin{pmatrix} 2 \\ 0 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} x-1 \\ y-5 \\ z-1 \end{pmatrix} = 0.$$

2.36 Spannen die Vektoren $\vec{a}_1 = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$, $\vec{a}_2 = \begin{pmatrix} 1 \\ 0 \\ -2 \end{pmatrix}$, $\vec{a}_3 = \begin{pmatrix} 3 \\ 1 \\ 1 \end{pmatrix}$ den $\mathbb{R}^3$ auf?

2.37 Sind die folgenden Vektoren des $\mathbb{R}^4$ linear unabhängig?

$$\vec{a}_1 = \begin{pmatrix} 2 \\ -1 \\ 3 \\ 0 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 2 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 3 \\ 0 \\ 1 \\ 4 \end{pmatrix}, \vec{a}_4 = \begin{pmatrix} 5 \\ -2 \\ 2 \\ 3 \end{pmatrix}.$$

2.38 Im $\mathbb{R}^4$ sind die Vektoren

$$\vec{a}_1 = \begin{pmatrix} 2 \\ 0 \\ 1 \\ 3 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 2 \\ 3 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 1 \\ -1 \\ 0 \\ 0 \end{pmatrix}, \vec{a}_4 = \begin{pmatrix} 0 \\ -2 \\ 1 \\ 0 \end{pmatrix}, \vec{b} = \begin{pmatrix} 0 \\ 5 \\ 2 \\ 6 \end{pmatrix}$$

gegeben. Man stelle $\vec{b}$ als Linearkombination von $\vec{a}_1, \vec{a}_2, \vec{a}_3, \vec{a}_4$ dar.

2.39 Untersuchen Sie die folgenden Vektoren des $\mathbb{R}^5$ auf lineare Abhängigkeit:

$$\vec{a}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_4 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, \vec{a}_5 = \begin{pmatrix} 0 \\ 1 \\ 1 \\ 1 \\ 1 \end{pmatrix}.$$

2.40 Ist der Vektor $\vec{b}$ im Erzeugnis der Vektoren $\vec{a}_1, \vec{a}_2, \vec{a}_3$?

$$\text{a) } \vec{a}_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \vec{b} = \begin{pmatrix} 2 \\ 2 \\ 1 \end{pmatrix}$$

$$\text{b) } \vec{a}_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}, \vec{b} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$$

2.41 Zeigen Sie, dass die Vektoren $\vec{a}_1, \vec{a}_2, \vec{a}_3$ eine Basis des $\mathbb{R}^3$ bilden und stellen Sie $\vec{d}$ als Linearkombination von $\vec{a}_1, \vec{a}_2, \vec{a}_3$ dar:

$$\vec{a}_1 = \begin{pmatrix} 3 \\ 4 \\ 3 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 5 \\ -1 \\ 0 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} -2 \\ 2 \\ -3 \end{pmatrix}, \vec{d} = \begin{pmatrix} 1 \\ -11 \\ -3 \end{pmatrix}$$



Kapitel 3

Matrizen und Determinanten

3

3	Matrizen und Determinanten	97
3.1	Matrizen	99
3.1.1	Einführung, spezielle Matrizen	99
3.1.2	Rechenoperationen für Matrizen	101
3.1.3	Inverse Matrix	105
3.1.4	Lineare Abbildungen	109
3.1.5	Anwendungsbeispiele	110
3.2	Determinanten	114
3.2.1	Einführung	114
3.2.2	Rechenregeln für zweireihige Determinanten	115
3.2.3	n -reihige Determinanten	118
3.2.4	Anwendungen von Determinanten	123
3.3	Lösbarkeit von linearen Gleichungssystemen	124
3.3.1	Lineare Gleichungssysteme, Rang	124
3.3.2	Anwendungen	129
3.4	Aufgaben zu Matrizen und Determinanten	135
 Zusätzliche Abschnitte auf der CD-Rom		
3.5	MAPLE: Matrizen und Determinanten	cd

3 Matrizen und Determinanten

Durch die Konstruktion der Koeffizientenmatrix in Kapitel 1 werden lineare Gleichungssysteme sehr kompakt beschrieben. In diesem Kapitel werden wir den Begriff der Matrix und der den quadratischen Matrizen zugeordneten Determinanten nicht nur als abkürzende Bezeichnungen kennen lernen, sondern mit ihnen Rechenoperationen durchführen, die wir dann beim Lösen von linearen Gleichungssystemen einsetzen. Beim Anwendungsbeispiel der gekoppelten Pendel werden wir aufzeigen, dass man mit der Beschreibung des physikalischen Systems durch die Systemmatrix und der Berechnung der zugehörigen Determinante die Eigenfrequenzen des Systems bestimmt.

3.1

3.1 Matrizen

Grundlegend für dieses Kapitel ist der Begriff der Matrix und das Rechnen mit Matrizen. Es werden die Addition und die Multiplikation von Matrizen definiert sowie das Gauß-Jordan-Verfahren zur Berechnung der inversen Matrix eingeführt.

➤ 3.1.1 Einführung, spezielle Matrizen

Definition: Unter einer $(m \times n)$ -Matrix $A = (a_{ij})_{mn}$ versteht man ein rechteckiges Zahlenschema

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1j} & \cdots & a_{1n} \\ \vdots & \vdots & & \vdots & & \vdots \\ a_{i1} & a_{i2} & \cdots & a_{ij} & \cdots & a_{in} \\ \vdots & \vdots & & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mj} & \cdots & a_{mn} \end{pmatrix} \quad \leftarrow i\text{-te Zeile} = (a_{ij})_{mn}$$

$\uparrow$
j-te Spalte

mit $a_{ij} \in \mathbb{R}$. Man nennt a_{ij} die Matrixelemente, i den Zeilenindex und j den Spaltenindex.

Eine Matrix setzt sich zusammen aus ihren *Spaltenvektoren* (kurz: Spalten) bzw. aus ihren *Zeilenvektoren* (kurz: Zeilen). Eine $(m \times n)$ -Matrix A hat n Spalten und m Zeilen. Der Index i gibt die Nummer der Zeile und der Index j die Nummer der Spalte an.

Beispiele 3.1:

① $A = \begin{pmatrix} 5 & 3 & \boxed{1} \\ 2 & 0 & 1 \end{pmatrix}$ ist eine (2×3) -Matrix. Das Element $a_{13} = 1$.

② $B = \begin{pmatrix} 0 & 1 \\ 0 & \boxed{2} \\ 4 & 2 \end{pmatrix}$ ist eine (3×2) -Matrix. Das Element $b_{22} = 2$. $\square$

③ **Quadratische Matrizen:**

Die Matrix A heißt **quadratisch**, wenn $n = m$. Falls $n = m$ ist, nennt man die *Matrixelemente* $a_{11}, a_{22}, \dots, a_{nn}$ die *Hauptdiagonale* (kurz: *Diagonale*) der Matrix

$$A = \begin{pmatrix} a_{11} & \cdots & \cdots & \cdots & a_{1n} \\ \vdots & \ddots & & & \vdots \\ \vdots & & \ddots & & \vdots \\ \vdots & & & \ddots & \vdots \\ a_{n1} & \cdots & \cdots & \cdots & a_{nn} \end{pmatrix}$$


 Diagonale

④ **Spezielle quadratische Matrizen:**

Eine Matrix D heißt **Diagonalmatrix**, wenn alle Nichtdiagonal-Elemente Null sind. Die **Einheitsmatrix** I_n ist eine Diagonalmatrix, die auf der Diagonalen nur Elemente mit dem Wert 1 enthält.

$$D = \begin{pmatrix} a_{11} & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & a_{nn} \end{pmatrix}, \quad I_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}.$$

Eine Matrix O heißt **obere Dreiecksmatrix**, wenn alle Elemente unterhalb der Diagonalen Null sind; eine Matrix U heißt **untere Dreiecksmatrix**, wenn alle Elemente oberhalb der Diagonalen Null sind.

$$O = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ & \ddots & \vdots \\ 0 & & a_{nn} \end{pmatrix}, \quad U = \begin{pmatrix} a_{11} & & 0 \\ \vdots & \ddots & \\ a_{n1} & \cdots & a_{nn} \end{pmatrix}.$$

Eine quadratische Matrix heißt **symmetrisch**, wenn

$$a_{ij} = a_{ji} \quad \text{für alle } i, j = 1 \dots n.$$

Eine symmetrische Matrix S besteht aus Elementen, die spiegelsymmetrisch zur Hauptdiagonalen angeordnet sind. Die **Nullmatrix** N hat als Elemente nur Nullen.

$$S = \begin{pmatrix} 1 & 2 & 3 \\ 2 & 5 & 4 \\ 3 & 4 & 6 \end{pmatrix}, \quad N = \begin{pmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{pmatrix}.$$

► 3.1.2 Rechenoperationen für Matrizen

➤ (1) Multiplikation einer Matrix mit einem Skalar

Definition: Ist $\alpha \in \mathbb{R}$, $A = (a_{ij})_{mn}$ eine $(m \times n)$ -Matrix, dann ist

$$\alpha \cdot A = \alpha (a_{ij})_{mn} := (\alpha a_{ij})_{mn}.$$

Eine Matrix wird mit einer Zahl multipliziert, indem jedes Element der Matrix mit der Zahl multipliziert wird.

➤ (2) Addition zweier $(m \times n)$ -Matrizen

Definition: Sind $A = (a_{ij})_{mn}$ und $B = (b_{ij})_{mn}$ zwei $(m \times n)$ -Matrizen, dann ist die **Summe** (bzw. **Differenz**) der Matrizen

$$A \pm B = (a_{ij})_{mn} \pm (b_{ij})_{mn} := (a_{ij} \pm b_{ij})_{mn}.$$

Zwei $(m \times n)$ -Matrizen werden addiert bzw. subtrahiert, indem die jeweiligen Matrixelemente addiert bzw. subtrahiert werden. Man beachte, dass sowohl für die Addition als auch Subtraktion beide Matrizen $(m \times n)$ -Matrizen sein müssen.

Beispiele 3.2: $A = \begin{pmatrix} 5 & 0 & 1 \\ 2 & 3 & 7 \end{pmatrix}; B = \begin{pmatrix} 2 & -4 & 3 \\ 1 & 2 & 3 \end{pmatrix}.$

$$\textcircled{1} \quad 2A + 4B = \begin{pmatrix} 10 & 0 & 2 \\ 4 & 6 & 14 \end{pmatrix} + \begin{pmatrix} 8 & -16 & 12 \\ 4 & 8 & 12 \end{pmatrix} = \begin{pmatrix} 18 & -16 & 14 \\ 8 & 14 & 26 \end{pmatrix}.$$

$$\textcircled{2} \quad 4A - 2B = \begin{pmatrix} 20 & 0 & 4 \\ 8 & 12 & 28 \end{pmatrix} - \begin{pmatrix} 4 & -8 & 6 \\ 2 & 4 & 6 \end{pmatrix} = \begin{pmatrix} 16 & 8 & -2 \\ 6 & 8 & 22 \end{pmatrix}.$$

□

⊗ (3) Transponieren einer Matrix

Definition: Vertauscht man Zeilen und Spalten einer $(m \times n)$ -Matrix $A = (a_{ij})_{mn}$, so entsteht die zu A **transponierte Matrix** A^t

$$A^t := (a_{ji})_{nm}.$$

Beispiel 3.3. $A = \begin{pmatrix} 5 & 0 & 1 \\ 2 & 3 & 7 \end{pmatrix} \Rightarrow A^t = \begin{pmatrix} 5 & 2 \\ 0 & 3 \\ 1 & 7 \end{pmatrix}.$ □

Bemerkungen:

- (1) Ist A eine $(m \times n)$ -Matrix, dann ist A^t eine $(n \times m)$ -Matrix.
- (2) Die transponierte Matrix erhält man, indem man die Matrixelemente an der Hauptdiagonalen spiegelt.
- (3) $(A^t)^t = A$.
- (4) Eine quadratische Matrix ist **symmetrisch**, falls $A^t = A$.

⊗ (4) Multiplikation von Matrizen

Um die Matrizenmultiplikation zu motivieren, führen wir die folgenden Vorüberlegungen durch. Das inhomogene LGS

$$\left. \begin{array}{l} 3x_1 + 4x_2 - 3x_3 = 4 \\ 5x_1 - x_2 + 2x_3 = 6 \\ 4x_1 - 2x_2 - 2x_3 = 1 \end{array} \right\} \quad \left(\begin{array}{ccc|c} 3 & 4 & -3 & 4 \\ 5 & -1 & 2 & 6 \\ 4 & -2 & -2 & 1 \end{array} \right)$$

wird mit der Matrizen Schreibweise abgekürzt. Schreiben wir es in der Form

$$\begin{pmatrix} 3 & 4 & -3 \\ 5 & -1 & 2 \\ 4 & -2 & -2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 4 \\ 6 \\ 1 \end{pmatrix}$$

ist $A \vec{x} = \vec{b}$ eine Vektorgleichung mit der Matrix

$$A = \begin{pmatrix} 3 & 4 & -3 \\ 5 & -1 & 2 \\ 4 & -2 & -2 \end{pmatrix} \quad \text{und den Vektoren} \quad \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 4 \\ 6 \\ 1 \end{pmatrix}.$$

Die Ausgangsgleichungen erhält man aus der Matrizen Schreibweise zurück, indem der Spaltenvektor $\vec{x}$ jeweils über jede Zeile von A gelegt und das Skalarprodukt des Zeilenvektors mit dem Spaltenvektor $\vec{x}$ gebildet wird. Somit haben wir die Multiplikation einer Matrix A mit einem Spaltenvektor $\vec{x}$ erklärt. Auf analoge Weise wird das Produkt zweier Matrizen $A \cdot B$ definiert:

Man legt die erste *Spalte* der Matrix B über die erste *Zeile* der Matrix A , um das Element c_{11} der Produktmatrix zu erhalten. Anschließend wählt man die erste Spalte der Matrix B und legt diese über die zweite Zeile der Matrix A , um das Element c_{12} der Produktmatrix zu erhalten usw. So erhält man allgemein die folgende Definition des Matrizenproduktes:

Definition: (Matrizenprodukt)

Sei $A = (a_{ij})_{mn}$ eine $(m \times n)$ -Matrix (m Zeilen und n Spalten) und $B = (b_{jk})_{nl}$ eine $(n \times l)$ -Matrix (n Zeilen und l Spalten). Dann ist das **Produkt** $C = A \cdot B = (c_{ik})_{ml}$ eine $(m \times l)$ -Matrix definiert durch

$$c_{ik} := \sum_{j=1}^n a_{ij} b_{jk} = \begin{pmatrix} a_{i1} \\ a_{i2} \\ \vdots \\ a_{in} \end{pmatrix} \begin{pmatrix} b_{1k} \\ b_{2k} \\ \vdots \\ b_{nk} \end{pmatrix} \quad \begin{matrix} i = 1 \dots m \\ k = 1 \dots l \end{matrix}.$$

⚠ Achtung: Um das Produkt von zwei Matrizen bilden zu können, muss die Spaltenzahl von A mit der Zeilenzahl von B übereinstimmen. Das Element c_{ik} von C ist das Skalarprodukt der i -ten Zeile von A mit der k -ten Spalte von B . D.h. die Zeilenlänge der Matrix A muss mit der Spaltenlänge der Matrix B übereinstimmen!

Beispiele 3.4:

$$\begin{aligned} \textcircled{1} \quad & \begin{matrix} i = 1 \\ i = 2 \end{matrix} \begin{pmatrix} \boxed{2} & \boxed{1} \\ \boxed{0} & \boxed{4} \end{pmatrix} \cdot \begin{matrix} \begin{matrix} \uparrow & \uparrow & \uparrow \\ \boxed{3} & \boxed{-2} & \boxed{0} \\ \boxed{0} & \boxed{4} & \boxed{2} \end{matrix} \\ j = 1 & \quad j = 2 \quad j = 3 \end{matrix} \\ &= \begin{pmatrix} 2 \cdot 3 + 1 \cdot 0 & 2 \cdot (-2) + 1 \cdot 4 & 2 \cdot 0 + 1 \cdot 2 \\ 0 \cdot 3 + 4 \cdot 0 & 0 \cdot (-2) + 4 \cdot 4 & 0 \cdot 0 + 4 \cdot 2 \end{pmatrix} = \begin{pmatrix} 6 & 0 & 2 \\ 0 & 16 & 8 \end{pmatrix}. \\ \\ \textcircled{2} \quad & \begin{matrix} i = 1 \\ i = 2 \\ i = 3 \end{matrix} \begin{pmatrix} \boxed{1} & \boxed{0} & \boxed{3} \\ \boxed{4} & \boxed{2} & \boxed{1} \\ \boxed{0} & \boxed{1} & \boxed{-1} \end{pmatrix} \cdot \begin{matrix} \begin{matrix} \uparrow & \uparrow & \uparrow \\ \boxed{2} & \boxed{2} & \boxed{-1} \\ \boxed{1} & \boxed{5} & \boxed{4} \\ \boxed{3} & \boxed{1} & \boxed{2} \end{matrix} \\ j = 1 & \quad j = 2 \quad j = 3 \end{matrix} \\ &= \begin{pmatrix} 2 + 9 & 2 + 3 & -1 + 6 \\ 8 + 2 + 3 & 8 + 10 + 1 & -4 + 8 + 2 \\ 1 - 3 & 5 - 1 & 4 - 2 \end{pmatrix} = \begin{pmatrix} 11 & 5 & 5 \\ 13 & 19 & 6 \\ -2 & 4 & 2 \end{pmatrix}. \quad \square \end{aligned}$$

Die Multiplikation wird so ausgeführt, dass die Spalten von B über die Zeilen von A gelegt werden und anschließend das **Skalarprodukt** berechnet wird.

Bemerkungen:

- (1)  **Achtung:** Aus der Definition des Produktes ist ersichtlich, dass für das Produkt zweier Matrizen $A \cdot B$ die Zeilenlänge von A mit der Spaltenlänge von B übereinstimmen muss. Ist also $A \cdot B$ definiert, so muss nicht notwendigerweise $B \cdot A$ definiert sein. Falls beide Produkte definiert sind (z.B. bei quadratischen Matrizen), ist i.A. $A \cdot B \neq B \cdot A$.
- (2) Zur Berechnung des Matrizenproduktes trägt man zweckmäßigerweise die Matrizen in ein Schema (dem sog. **Falk-Schema**) ein:

			0	1	
			1	2	
			3	4	
$A =$					$= B$
	2	0	2		
	4	3	1	6	10
				6	14
					$= A \cdot B = C$

		2	0	2	
		4	3	1	
$B =$	0	1			$= A$
	1	2			
	3	4			
			4	3	1
			10	6	4
			22	12	10
					$= B \cdot A = D$

Das Element c_{12} der Produktmatrix $C = A \cdot B$ ergibt sich durch einfache Rechnung $c_{12} = 2 \cdot 1 + 0 \cdot 2 + 2 \cdot 4 = 10$; das Element d_{22} des Produktes $D = B \cdot A$ durch $d_{22} = 1 \cdot 0 + 2 \cdot 3 = 6$. Anhand dieser beiden Beispiele erkennt man, dass $A \cdot B \neq B \cdot A$!

- (3) Das Produkt einer $(m \times n)$ -Matrix A mit einer $(n \times k)$ -Matrix B kann man auch folgendermaßen interpretieren:

B besteht aus k Spaltenvektoren $\vec{b}_1, \dots, \vec{b}_k \Rightarrow B = (\vec{b}_1, \dots, \vec{b}_k)$. Dann sind die Spalten des Produktes $C = A \cdot B$ gegeben durch

$$\vec{c}_j = A \vec{b}_j \quad (j = 1, \dots, k)$$

und

$$C = (\vec{c}_1, \vec{c}_2, \dots, \vec{c}_k) = (\overrightarrow{A b_1}, \overrightarrow{A b_2}, \dots, \overrightarrow{A b_k}).$$

Rechenregeln für Matrizenprodukte

Für Produkte von Matrizen A, B, C gelten die folgenden Regeln (die jeweiligen Produkte seien definiert):

$$(A \cdot B) \cdot C = A \cdot (B \cdot C) \quad \text{Assoziativgesetz}$$

$$(A \cdot B)^t = B^t \cdot A^t$$

$$A \cdot (B + C) = A \cdot B + A \cdot C \quad \text{Distributivgesetz}$$

$$(A + B) \cdot C = A \cdot C + B \cdot C \quad \text{Distributivgesetz}$$

3.1.3 Inverse Matrix

Die Gleichung $a \cdot x = b$ besitzt für $a \neq 0$ genau eine Lösung. Mit dem zu a inversen Element a^{-1} gilt

$$x = a^{-1}b = \frac{b}{a} \in \mathbb{R}, \quad \text{da} \quad a \cdot a^{-1} = a^{-1} \cdot a = 1.$$

Diese Konstruktion wird für quadratische Matrizen verallgemeinert: Um die Gleichung

$$A \vec{x} = \vec{b}$$

nach $\vec{x}$ aufzulösen, ist eine Matrix A^{-1} gesucht, so dass

$$A \cdot A^{-1} = A^{-1} \cdot A = I_n.$$

Denn dann ist die Lösung für den Vektor $\vec{x}$ gegeben durch $\vec{x} = A^{-1} \vec{b}$. Man definiert:

Definition: Gibt es zu einer quadratischen $(n \times n)$ -Matrix A eine $(n \times n)$ -Matrix X mit

$$A \cdot X = X \cdot A = I_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix},$$

so heißt X die **zu A inverse Matrix**. Sie wird durch das Symbol A^{-1} gekennzeichnet.

Bemerkungen:

- (1) Besitzt A eine inverse Matrix A^{-1} , so heißt A *invertierbar* bzw. *regulär*. A^{-1} heißt *Umkehrmatrix* oder *Inverse*.
- (2) Eine quadratische Matrix besitzt -wenn überhaupt- genau eine *Inverse*.
- (3) Aufgrund der Definition gilt $A \cdot A^{-1} = A^{-1} \cdot A = I_n$, d.h. A und A^{-1} sind kommutativ.
- (4)  **Achtung:** Nicht jede quadratische Matrix ist umkehrbar: Z.B. $\begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$ ist nicht umkehrbar!

Im Folgenden werden wir ein Schema vorstellen, mit dem man entscheiden kann, ob eine $(n \times n)$ -Matrix A invertierbar ist und wie die Inverse A^{-1} sich dann berechnet. Dazu gehen wir zunächst davon aus, dass die zu A inverse Matrix A^{-1} existiert. A^{-1} sei gegeben durch Spaltenvektoren $\vec{s}_1, \dots, \vec{s}_n$, d.h.

$$A^{-1} = (\vec{s}_1, \vec{s}_2, \dots, \vec{s}_n).$$

Aufgrund der Eigenschaft der inversen Matrix

$$A \cdot A^{-1} = I_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}$$

ist die k -te Spalte des Produktes $A \cdot A^{-1}$ der Einheitsvektor $\vec{e}_k$. Man erhält die k -te Spalte des Produktes, indem man die Matrix A auf den k -ten Spaltenvektor von A^{-1} multipliziert:

$$A \vec{s}_k = \vec{e}_k = \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \leftarrow k \quad k = 1, 2, \dots, n.$$

Folglich sind die Spalten der inversen Matrix die Lösungen der linearen Gleichungssysteme $A \vec{s}_k = \vec{e}_k$. Man muss zur Bestimmung der inversen Matrix also n LGS mit der selben Matrix A lösen. Nur die rechten Seiten der LGS sind unterschiedlich. Da die Rechenoperationen beim Gauß-Algorithmus nur durch die Koeffizienten der Matrix A nicht aber von der rechten Seite bestimmt werden, können die n LGS simultan gelöst werden. Dies führt auf den *Gauß-Jordan-Algorithmus*:

Berechnung der inversen Matrix mit dem Gauß-Jordan-Verfahren

- (1) Man erstellt das aus A und allen rechten Seiten kombinierte Matrixschema:

$$(A | I_n) = \left(\begin{array}{cccc|cccc} a_{11} & \cdots & \cdots & a_{1n} & 1 & 0 & \cdots & 0 \\ \vdots & & & \vdots & 0 & \ddots & \ddots & \vdots \\ \vdots & & & \vdots & \vdots & \ddots & \ddots & 0 \\ a_{n1} & \cdots & \cdots & a_{nn} & 0 & \cdots & 0 & 1 \end{array} \right)$$

- (2) Mit Hilfe elementarer Zeilenumformungen wird die linke Seite des Schemas so umgeformt, dass die Einheitsmatrix I_n den Platz von A einnimmt. Die inverse Matrix A^{-1} steht dann auf dem Platz von I_n :

$$\left(\begin{array}{cccc|cccc} 1 & 0 & \cdots & 0 & b_{11} & \cdots & \cdots & b_{1n} \\ 0 & \ddots & \ddots & \vdots & \vdots & & & \vdots \\ \vdots & \ddots & \ddots & 0 & \vdots & & & \vdots \\ 0 & \cdots & 0 & 1 & b_{n1} & \cdots & \cdots & b_{nn} \end{array} \right) = (I_n | A^{-1})$$

⚠ Enthält zum Schluss die linke Seite eine Nullzeile, dann ist die Matrix A nicht invertierbar.

Beispiele 3.5:

- ① Gesucht ist die inverse Matrix zu $A = \begin{pmatrix} 1 & 0 & -2 \\ -1 & 1 & 2 \\ 0 & 4 & -1 \end{pmatrix}$.

Wir führen den Gauß-Jordan-Algorithmus durch und schreiben die jeweiligen Operationen an die rechte Seite des Schemas:

$$\begin{aligned} (A | I_3) &= \left(\begin{array}{ccc|ccc} 1 & 0 & -2 & 1 & 0 & 0 \\ -1 & 1 & 2 & 0 & 1 & 0 \\ 0 & 4 & -1 & 0 & 0 & 1 \end{array} \right) &+ Z_1 \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 0 & -2 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 4 & -1 & 0 & 0 & 1 \end{array} \right) &-4Z_2 \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 0 & -2 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & -1 & -4 & -4 & 1 \end{array} \right) &-2Z_3 \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 9 & 8 & -2 \\ 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 4 & 4 & -1 \end{array} \right) &(-1) \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 9 & 8 & -2 \\ 0 & 1 & 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 4 & 4 & -1 \end{array} \right) &= (I_3 | A^{-1}) \end{aligned}$$

Die zu A inverse Matrix lautet $A^{-1} = \begin{pmatrix} 9 & 8 & -2 \\ 1 & 1 & 0 \\ 4 & 4 & -1 \end{pmatrix}$.

② Gesucht ist die Inverse von $A = \begin{pmatrix} 1 & 1 & 2 \\ 2 & 1 & -2 \\ -3 & -2 & 0 \end{pmatrix}$:

$$\begin{aligned} (A | I_3) &= \left(\begin{array}{ccc|ccc} 1 & 1 & 2 & 1 & 0 & 0 \\ 2 & 1 & -2 & 0 & 1 & 0 \\ -3 & -2 & 0 & 0 & 0 & 1 \end{array} \right) \begin{array}{l} \\ -2Z_1 \\ 3Z_1 \end{array} \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 1 & 2 & 1 & 0 & 0 \\ 0 & -1 & -6 & -2 & 1 & 0 \\ 0 & 1 & 6 & 3 & 0 & 1 \end{array} \right) +Z_2 \\ \Rightarrow &\left(\begin{array}{ccc|ccc} 1 & 1 & 2 & 1 & 0 & 0 \\ 0 & -1 & -6 & -2 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 \end{array} \right) \end{aligned}$$

Die linke Seite enthält eine Nullzeile. Damit sind die linearen Gleichungssysteme nicht lösbar. A ist **nicht invertierbar**. $\square$

Rechenregeln bei der Inversion von Matrizen: Seien A und B invertierbare $(n \times n)$ -Matrizen. Dann gilt:

(1) Die Inverse einer invertierbaren Matrix ist invertierbar mit

$$(A^{-1})^{-1} = A$$

(2) Das Produkt zweier invertierbarer Matrizen ist invertierbar:

$$(A \cdot B)^{-1} = B^{-1} \cdot A^{-1}$$

(3) Die Transponierte einer invertierbaren Matrix ist invertierbar:

$$(A^t)^{-1} = (A^{-1})^t$$

Die Bedeutung von invertierbaren Matrizen werden wir noch ausführlicher im Kapitel über die Lösbarkeit von LGS ($\rightarrow$ Abschnitt 3.3) diskutieren.

► 3.1.4 Lineare Abbildungen

Wir werden nun eine geometrische Deutung von Matrizen angeben. Sei $A = (a_{ij})$ eine $(m \times n)$ -Matrix und $\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n$ ein Vektor. Dann ist das Produkt

$$A\vec{x} = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} \sum_{j=1}^n a_{1j} x_j \\ \vdots \\ \sum_{j=1}^n a_{mj} x_j \end{pmatrix} = \begin{pmatrix} y_1 \\ \vdots \\ y_m \end{pmatrix}$$

ein Vektor $\vec{y} \in \mathbb{R}^m$. A legt also eine Abbildung $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ fest, die jedem Vektor $\vec{x} \in \mathbb{R}^n$ genau einen Vektor $\vec{y} \in \mathbb{R}^m$ zuordnet. Es gilt folgender Satz:

Satz: Die Gesamtheit aller $(m \times n)$ -Matrizen entspricht in umkehrbar eindeutiger Weise der Gesamtheit der linearen Abbildungen $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$.

Begründung:

- (1) Die Matrix $A = (a_{ij})_{mn}$ definiert eine **lineare Abbildung** von $\mathbb{R}^n$ nach $\mathbb{R}^m$, denn man rechnet sofort nach, dass

$$A(\alpha \vec{x}_1 + \beta \vec{x}_2) = \alpha A\vec{x}_1 + \beta A\vec{x}_2$$

für beliebige Vektoren $\vec{x}_1, \vec{x}_2$ und Zahlen α, β .

- (2) Ist $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine lineare Abbildung. Dann wird diese lineare Abbildung eindeutig festgelegt durch die Bilder der Einheitsvektoren $\vec{e}_j$:

$$\varphi(\vec{e}_j) = \vec{a}_j = \begin{pmatrix} a_{1j} \\ \vdots \\ a_{mj} \end{pmatrix} \quad j = 1, \dots, n.$$

Mit diesen Vektoren bilden wir die Matrix

$$A = (a_{ij})_{mn} = (\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n).$$

Man rechnet unmittelbar nach, dass $A\vec{e}_j = \vec{a}_j$. Um die Matrix einer linearen Abbildung aufzustellen, ist es wichtig sich zu merken:

Die Spalten von A sind genau die Bildvektoren der Einheitsvektoren.

Beispiel 3.6. Die Abbildung $f : \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \mapsto \begin{pmatrix} x_1 + x_2 - x_3 \\ x_2 + x_3 - x_1 \end{pmatrix}$ ist eine lineare Abbildung von $\mathbb{R}^3$ nach $\mathbb{R}^2$. Die Bilder der Basisvektoren sind

$$\begin{aligned} f(\vec{e}_1) &= f\left(\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}\right) = \begin{pmatrix} 1 \\ -1 \end{pmatrix}; \\ f(\vec{e}_2) &= \begin{pmatrix} 1 \\ 1 \end{pmatrix}; \quad f(\vec{e}_3) = \begin{pmatrix} -1 \\ 1 \end{pmatrix}. \end{aligned}$$

Damit wird die lineare Abbildung dargestellt durch die Matrix

$$A = \begin{pmatrix} 1 & 1 & -1 \\ -1 & 1 & 1 \end{pmatrix},$$

denn die Spalte der Matrix A sind die Bilder der Einheitsvektoren. $\square$

► 3.1.5 Anwendungsbeispiele

Beispiel 3.7 (Anwendungsbeispiel: Rohstoffkette).

In einem Betrieb werden aus vier Rohstoffen mit Einheiten r_1, r_2, r_3, r_4 drei Zwischenprodukte mit Einheiten z_1, z_2, z_3 hergestellt. Aus den Zwischenprodukten entstehen drei Endprodukte mit den Einheiten p_1, p_2, p_3 . Der Materialverbrauch sei gegeben durch die linearen Gleichungen

$$\begin{array}{lll} r_1 = & z_1 & + \quad z_3 \\ r_2 = & 2z_1 + z_2 + z_3 & \text{bzw.} \\ r_3 = & & z_2 + z_3 \\ r_4 = & z_1 + z_2 + 2z_3 & \end{array} \quad \begin{array}{l} z_1 = p_1 + 2p_2 + p_3 \\ z_2 = 2p_1 + 3p_2 + p_3 \\ z_3 = 4p_1 + 2p_2 + 2p_3 \end{array}$$

Führt man die Matrizen $A = \begin{pmatrix} 1 & 0 & 1 \\ 2 & 1 & 1 \\ 0 & 1 & 1 \\ 1 & 1 & 2 \end{pmatrix}$ und $B = \begin{pmatrix} 1 & 2 & 1 \\ 2 & 3 & 1 \\ 4 & 2 & 2 \end{pmatrix}$ ein, so kann

man die beiden Gleichungssysteme schreiben in der Form

$$\vec{r} = A \vec{z}; \quad \vec{z} = B \vec{p}.$$

Gesucht ist der Rohstoffbedarf, wenn 100 Einheiten von p_1 , 80 Einheiten von p_2 und 60 Einheiten von p_3 hergestellt werden sollen. Es ist

$$\vec{r} = A \vec{z} = A (B \vec{p}) = (A \cdot B) \vec{p}.$$

Zur Berechnung der Menge der Rohstoffe führt man das Matrizenprodukt $A \cdot B$ aus und wendet das Produkt auf den Vektor $\vec{p} = \begin{pmatrix} 100 \\ 80 \\ 60 \end{pmatrix}$ an:

$$\vec{r} = \begin{pmatrix} 5 & 4 & 3 \\ 8 & 9 & 5 \\ 6 & 5 & 3 \\ 11 & 9 & 6 \end{pmatrix} \begin{pmatrix} 100 \\ 80 \\ 60 \end{pmatrix} = \begin{pmatrix} 1000 \\ 1820 \\ 1180 \\ 2180 \end{pmatrix}.$$

Es werden also 1000 Einheiten des ersten, 1820 Einheiten des zweiten, 1180 Einheiten des dritten und 2180 Einheiten des vierten Rohstoffes benötigt. $\square$

Beispiel 3.8 (Anwendungsbeispiel: Beschreibung eines Vierpols).

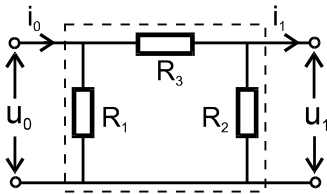


Abb. 3.1. Elektrischer Vierpol

Die Schaltung in Abb. 3.1 heißt linearer, *elektrischer Vierpol*. Gesucht ist der Zusammenhang zwischen den Eingangsgrößen i_0, u_0 und den Ausgangsgrößen i_1, u_1 .

Um die Modellgleichungen aufzustellen, verwenden wir die **Kirchhoffschen Gesetze**: Der *Knotensatz* besagt, dass die Summe der in einem Knoten zu- und abfließenden Ströme gleich Null ist. Der *Maschensatz* besagt, dass in einer Masche die Summe aller Spannungen Null ergibt:

$$i_0 = i_{R_1} + i_{R_3}$$

$$u_0 = R_3 i_{R_3} + u_1$$

$$i_{R_3} = i_{R_2} + i_1.$$

Um die gesuchte Abhängigkeit der Eingangsgrößen von der Ausgangsgrößen zu erkennen, ersetzt man $i_{R_1} = \frac{1}{R_1} u_0$, $i_{R_2} = \frac{1}{R_2} u_1$ und eliminiert die Größe i_{R_3} aus diesem System, so dass nur noch i_0, u_0 und i_1, u_1 vorkommen:

$$u_0 = \frac{R_2 + R_3}{R_2} u_1 + R_3 i_1$$

$$i_0 = \frac{R_1 + R_2 + R_3}{R_1 R_2} u_1 + \frac{R_1 + R_3}{R_1} i_1$$

Die Eingangsgrößen stehen nur auf der linken Seite der Gleichung; die Ausgangsgrößen auf der rechten Seite. Man hat nur noch zwei Gleichungen für zwei Unbekannte. Führt man die Verknüpfungsmatrix

$$M = \begin{pmatrix} \frac{R_2 + R_3}{R_2} & R_3 \\ \frac{R_1 + R_2 + R_3}{R_1 R_2} & \frac{R_1 + R_3}{R_1} \end{pmatrix}$$

ein, so gilt die folgende Beziehung zwischen den Eingangsgrößen i_0, u_0 und den Ausgangsgrößen i_1, u_1 :

$$\begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = M \begin{pmatrix} u_1 \\ i_1 \end{pmatrix}. \quad \square$$

Beispiel 3.9 (Anwendungsbeispiel: Kettenschaltung von Vierpolen).

Schaltet man einen zweiten Vierpol mit gleichen Widerständen hinzu, gilt

$$\begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = M \begin{pmatrix} u_1 \\ i_1 \end{pmatrix} \quad \text{und} \quad \begin{pmatrix} u_1 \\ i_1 \end{pmatrix} = M \begin{pmatrix} u_2 \\ i_2 \end{pmatrix},$$

d.h. zwischen den Eingangsgrößen u_0, i_0 und den Ausgangsgrößen u_2, i_2 besteht die Beziehung

$$\begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = M \cdot M \begin{pmatrix} u_2 \\ i_2 \end{pmatrix}.$$

Für eine lineare Kette von n identischen Vierpolen (siehe Abb. 3.2) gilt folglich

$$\begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = M^n \begin{pmatrix} u_n \\ i_n \end{pmatrix}.$$

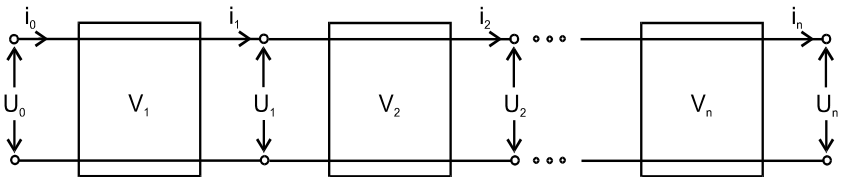


Abb. 3.2. Lineare Kette von Vierpolen

□

Beispiel 3.10 (Zahlenbeispiel, mit [MAPLE-Worksheet](#)).

Der in Beispiel 3.8 diskutierte Vierpol habe die Widerstände $R_1 = R_2 = 100\Omega$, $R_3 = 500\Omega$. Das Übertragungsverhalten des Vierpols ist dann gegeben durch die Übertragungsmatrix

$$M := \begin{pmatrix} 6 & 500 \\ 0.07 & 6 \end{pmatrix}.$$

- (i) Wie groß sind u_1 und i_1 , wenn $u_0 = 2V$ und $i_0 = 0.1A$?

Wegen $\begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = M \begin{pmatrix} u_1 \\ i_1 \end{pmatrix}$ folgt $\begin{pmatrix} u_1 \\ i_1 \end{pmatrix} = M^{-1} \begin{pmatrix} u_0 \\ i_0 \end{pmatrix}$. Die zu M inverse Matrix M^{-1} lautet

$$M^{-1} = \begin{pmatrix} 6 & -500 \\ -0.07 & 6 \end{pmatrix}$$

und damit ist

$$\begin{pmatrix} u_1 \\ i_1 \end{pmatrix} = M^{-1} \begin{pmatrix} u_0 \\ i_0 \end{pmatrix} = \begin{pmatrix} 6 & -500 \\ -0.07 & 6 \end{pmatrix} \begin{pmatrix} 2 \\ 0.1 \end{pmatrix} = \begin{pmatrix} -38 \\ 0.46 \end{pmatrix}.$$

Somit beträgt $u_1 = -38V$ und $i_1 = 0.46A$.

- (ii) Für eine Schaltung von drei Vierpolen erhalten wir die Transfermatrix T :

$$T = M^3 = \begin{pmatrix} 846.00 & 71500 \\ 10.01 & 846.00 \end{pmatrix} \quad \square$$

3.2 Determinanten

Um zu entscheiden, ob eine Matrix invertierbar ist oder nicht, muss zunächst das Gauß-Jordan-Verfahren angewendet werden. Enthält nach der Umformung die linke Seite eine Nullzeile, so ist die Matrix nicht invertierbar. Es wäre schön, wenn man vor der Anwendung des Verfahrens entscheiden könnte, ob die inverse Matrix existiert oder nicht. Dies führt auf den Begriff der Determinante.

3.2.1 Einführung

In diesem Kapitel diskutieren wir, unter welchen Bedingungen ein lineares, quadratisches Gleichungssystem eindeutig lösbar ist. Bei der expliziten Lösung eines Gleichungssystems mit dem Gauß-Algorithmus kann zum Schluss der Rechnung festgestellt werden, ob das Gleichungssystem lösbar ist oder nicht. Im Falle einer numerischen Lösung des LGS muss aber im Voraus bekannt sein, ob das LGS eindeutig lösbar ist. Im Folgenden wird ein Formalismus entwickelt (Determinanten-Berechnung), mit dem entschieden werden kann, ob quadratische LGS eindeutig lösbar sind. Alle in diesem Paragraphen auftretenden Matrizen sind **quadratisch**.

(1) Wir betrachten zunächst das Problem: Unter welchen Voraussetzungen ist die Gleichung

$$a_{11} x_1 = c_1 \quad (\text{I})$$

mit einer Unbekannten x_1 für beliebige c_1 eindeutig lösbar. Die Antwort lautet: System (I) ist genau dann eindeutig lösbar, falls $a_{11} \neq 0$: $x_1 = \frac{c_1}{a_{11}}$.

(2) Nun betrachten wir das LGS

$$\begin{aligned} a_{11} x_1 + a_{12} x_2 &= c_1 \\ a_{21} x_1 + a_{22} x_2 &= c_2 \end{aligned} \quad (\text{II})$$

mit zwei Unbekannten x_1 und x_2 . Welche Bedingungen müssen die Koeffizienten a_{11} , a_{12} , a_{21} , a_{22} erfüllen, damit das System (II) für beliebige c_1 , c_2 eindeutig lösbar ist? Die Lösung des Systems erhalten wir, indem wir Gl. II.1 mit $a_{22} \neq 0$ und Gl. II.2 mit $(-a_{12}) \neq 0$ multiplizieren und beide Gleichungen addieren:

$$\begin{aligned} &\left. \begin{aligned} a_{11} a_{22} x_1 + a_{12} a_{22} x_2 &= c_1 a_{22} \\ -a_{12} a_{21} x_1 - a_{12} a_{22} x_2 &= -c_2 a_{12} \end{aligned} \right\} \\ \Rightarrow (a_{11} a_{22} - a_{12} a_{21}) x_1 &= c_1 a_{22} - c_2 a_{12}. \end{aligned}$$

Indem wir Gl. II.1 mit $-a_{21} \neq 0$ und Gl. II.2 mit $a_{11} \neq 0$ multiplizieren und beide resultierenden Gleichungen addieren, erhalten wir analog

$$(a_{11} a_{22} - a_{12} a_{21}) x_2 = c_2 a_{11} - c_1 a_{21}.$$

Folglich ist das LGS (II) eindeutig für beliebige c_1 und c_2 lösbar, falls

$$a_{11} a_{22} - a_{12} a_{21} \neq 0.$$

Diese Bedingung beinhaltet auch die hier nicht aufgeführten Sonderfälle $a_{11} = 0$, $a_{12} = 0$, $a_{21} = 0$ oder $a_{22} = 0$.

Definition: Die aus der Koeffizientenmatrix $A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ berechnete Größe

$$D := \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11} a_{22} - a_{12} a_{21}$$

heißt **zweireihige Determinante** oder kurz **Determinante**.

Für Determinanten sind die Symbole D , $\det(A)$, $|A|$, $\det(a_{ij})$ gebräuchlich.

Beispiele 3.11:

$$\textcircled{1} \quad \det \begin{pmatrix} 5 & 2 \\ 3 & 1 \end{pmatrix} = 5 \cdot 1 - 3 \cdot 2 = -1.$$

$$\textcircled{2} \quad \det \begin{pmatrix} -2 & -2 \\ 1 & 1 \end{pmatrix} = -2 \cdot 1 - 1 \cdot (-2) = 0. \quad \square$$

Eine zweireihige Determinante wird also berechnet, indem die Differenz des Produktes der Hauptdiagonalelemente $a_{11} a_{22}$ mit dem Produkt der Nebendiagonalelemente $a_{12} a_{21}$ gebildet wird. Die Determinante einer Matrix ist damit eine reelle Zahl! Für zweireihige Determinanten rechnet man unmittelbar die folgenden Regeln nach:

► 3.2.2 Rechenregeln für zweireihige Determinanten

Regel 1: Der Wert einer Determinante ändert sich **nicht**, wenn man Zeilen und Spalten vertauscht:

$$\det A = \det A^t.$$

Beispiel: $\det \begin{pmatrix} 2 & 3 \\ -1 & 4 \end{pmatrix} = 8 + 3 = 11 = \det \begin{pmatrix} 2 & -1 \\ 3 & 4 \end{pmatrix}.$ □

Regel 2: Der Wert der Determinante **wechselt das Vorzeichen**, wenn man zwei Zeilen (oder zwei Spalten) miteinander vertauscht.

Beispiel: $\det \begin{pmatrix} 3 & 8 \\ 4 & -7 \end{pmatrix} = -21 - 32 = -53$; $\det \begin{pmatrix} 4 & -7 \\ 3 & 8 \end{pmatrix} = 32 + 21 = 53$. $\square$

Regel 3: Werden die Elemente **einer** (!) Zeile (oder Spalte) mit einem Skalar λ multipliziert, so wird der **Determinantenwert** mit λ multipliziert.

Beispiel: $\det \begin{pmatrix} \lambda a_{11} & \lambda a_{12} \\ a_{21} & a_{22} \end{pmatrix} = \lambda a_{11} a_{22} - a_{21} \lambda a_{12} = \lambda(a_{11} a_{22} - a_{21} a_{12})$
 $= \lambda \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$. $\square$

Regel 4: Die Determinante hat den Wert 0, wenn eine Spalte (oder Zeile) der Nullvektor ist. Allgemeiner, wenn die Spalten (oder Zeilen) linear abhängig sind.

Beispiele: $\det \begin{pmatrix} 5 & 0 \\ 3 & 0 \end{pmatrix} = 0$; $\det \begin{pmatrix} 5 & -10 \\ 3 & -6 \end{pmatrix} = 0$, $\det \begin{pmatrix} a_{11} & a_{12} \\ \lambda a_{11} & \lambda a_{12} \end{pmatrix} = 0$. $\square$

Regel 5: Der Wert einer Determinante ändert sich **nicht**, wenn man zu einer Zeile (oder Spalte) ein beliebiges Vielfaches einer anderen Zeile (oder Spalte) elementweise addiert.

Beweis:

$$\begin{aligned} \det \begin{pmatrix} a_{11} + \lambda a_{21} & a_{12} + \lambda a_{22} \\ a_{21} & a_{22} \end{pmatrix} &= (a_{11} + \lambda a_{21}) a_{22} - a_{21} (a_{12} + \lambda a_{22}) \\ &= a_{11} a_{22} - a_{21} a_{12} + \underbrace{\lambda a_{21} a_{22} - \lambda a_{21} a_{22}}_{=0} \\ &= \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}. \end{aligned} \quad \square$$

Regel 6: Multiplikationstheorem für Determinanten.

Für zwei Matrizen A, B gilt stets:

$$\det(A \cdot B) = \det(A) \cdot \det(B).$$

Beispiel: $\det \left(\begin{pmatrix} 0 & 4 \\ 3 & -1 \end{pmatrix} \begin{pmatrix} 4 & 2 \\ 3 & 2 \end{pmatrix} \right) = \det \begin{pmatrix} 12 & 8 \\ 9 & 4 \end{pmatrix} = 48 - 72 = -24,$
 $\det \begin{pmatrix} 0 & 4 \\ 3 & -1 \end{pmatrix} \cdot \det \begin{pmatrix} 4 & 2 \\ 3 & 2 \end{pmatrix} = (-12) \cdot 2 = -24. \quad \square$

Regel 7: Die Determinante einer **Dreiecksmatrix** besitzt den Wert des Produktes der Hauptdiagonalelemente.

Beispiel: $\det \begin{pmatrix} 4 & 7 \\ 0 & 2 \end{pmatrix} = 4 \cdot 2 = 8; \quad \det \begin{pmatrix} 3 & 0 \\ 5 & -1 \end{pmatrix} = 3(-1) = -3. \quad \square$

Bemerkung: Die aufgelisteten Regeln sind so formuliert, dass sie auch für allgemeine, n -reihige Determinanten gültig sind.

➤ **Übergang zu den Formeln für $(n \times n)$ -Matrizen:**

Übertragen wir die Überlegungen aus der Einführung auf ein lineares Gleichungssystem mit drei Unbekannten x_1, x_2, x_3 ,

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 &= b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 &= b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 &= b_3, \end{aligned}$$

so kann dieses LGS mit allgemeinen Koeffizienten nach den Unbekannten formal aufgelöst werden. Für x_1 erhält man unter Zuhilfenahme von MAPLE als Lösung

$$x_1 = \frac{a_{22}a_{33}b_1 - a_{22}b_3a_{13} + a_{23}a_{12}b_3 - a_{23}a_{32}b_1 + b_2a_{32}a_{13} - b_2a_{12}a_{33}}{a_{21}a_{32}a_{13} - a_{21}a_{12}a_{33} - a_{31}a_{22}a_{13} + a_{12}a_{31}a_{23} - a_{32}a_{11}a_{23} + a_{11}a_{22}a_{33}}.$$

Ähnliche Lösungsformeln erhält man für x_2 und x_3 . Auch bei der Berechnung von x_2 und x_3 wird durch den Term

$$a_{21}a_{32}a_{13} - a_{21}a_{12}a_{33} - a_{31}a_{22}a_{13} + a_{12}a_{31}a_{23} - a_{32}a_{11}a_{23} + a_{11}a_{22}a_{33}$$

dividiert. Das LGS ist also für jede rechte Seite eindeutig lösbar, sofern

$$a_{21}a_{32}a_{13} - a_{21}a_{12}a_{33} - a_{31}a_{22}a_{13} + a_{12}a_{31}a_{23} - a_{32}a_{11}a_{23} + a_{11}a_{22}a_{33} \neq 0.$$

Diese Zahl heißt die Determinante von $A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$.

Analog gibt es eine entsprechende Zahl für ein LGS mit vier Unbekannten. Diese Zahl wird wieder als Determinante der zugehörigen Matrix bezeichnet. Die allgemeine Berechnung der Determinante erfolgt mit dem folgenden Schema:

3.2.3 n -reihige Determinanten

Definition / Satz: Entwicklungssatz nach Laplace.

Sei $A = (a_{ij}) = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{pmatrix}$ eine $(n \times n)$ -Matrix. Dann ist die **Determinante der Matrix A** definiert durch

$$\det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A'_{ij} \quad \text{für ein festes } i \in \{1, \dots, n\}.$$

Dabei ist A'_{ij} die $(n-1) \times (n-1)$ -Untermatrix, die aus A entsteht, indem die i -te Zeile und j -te Spalte gestrichen wird:

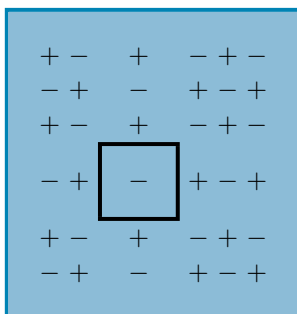
$$A'_{ij} = \begin{pmatrix} a_{11} & \cdots & a_{1j} & \cdots & a_{1n} \\ \vdots & & \vdots & & \vdots \\ a_{i1} & \cdots & a_{ij} & \cdots & a_{in} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \cdots & a_{nj} & \cdots & a_{nn} \end{pmatrix} \leftarrow i\text{-te Zeile}$$

Man bezeichnet dieses Vorgehen als **Entwicklung der Determinante nach der i -ten Zeile**.

Da A'_{ij} eine $(n-1) \times (n-1)$ -Matrix ist, bezeichnet man $\det(A'_{ij})$ als **Unterdeterminante**. Durch wiederholtes Anwenden der Entwicklungsformel lässt sich eine n -reihige Determinante auf 2-reihige Determinanten zurückführen und damit berechnen.

Bemerkungen:

- (1)  **Achtung:** Die Determinante ist nur für **quadratische** Matrizen definiert.
- (2) Das Vorzeichen $(-1)^{i+j}$ kann man sich als Vorzeichen im folgenden Schachbrettmuster vorstellen. Z.B. ist das Vorzeichen von $a_{43} : (-1)^{4+3} = -1$.



- (3) Die Zeile $i \in \{1, \dots, n\}$, nach der die Determinante entwickelt wird, ist beliebig; der Determinantenwert hängt nicht von der Wahl von i ab.
- (4) Die Entwicklung der Determinante nach der i -ten Zeile kann man sich so vorstellen, dass die i -te Zeile gestrichen wird und nacheinander ein Fadenzkreuz jeweils eine Spalte streicht:

1. Spalte: $(-1)^{i+1} a_{i1} \det A'_{i1}$

2. Spalte: $(-1)^{i+2} a_{i2} \det A'_{i2}$

$\vdots$

n -te Spalte: $(-1)^{i+n} a_{in} \det A'_{in}$.

Das Aufsummieren aller Unterdeterminanten mit dem zugehörigen Faktor a_{ij} und dem Vorzeichen entsprechend dem Schachbrettmuster liefert den Determinantenwert.

- (5) Die Determinante lässt sich auch nach der k -ten **Spalte** entwickeln:

$$\det A = \sum_{i=1}^n (-1)^{i+k} a_{ik} \det A'_{ik} \quad \text{für festes } k \in \{1, 2, \dots, n\}.$$

- (6) Für n -reihige Determinanten gelten die gleichen Rechenregeln, wie die in Abschnitt 3.2.2 angegebenen Rechenregeln für 2-reihige Determinanten.

Spezialfälle:

$$n = 1: \quad \det(a_{11}) = a_{11}.$$

$n = 2$: Entwicklung nach der ersten Spalte

$$\begin{aligned} \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} &= a_{11} - \det(a_{22}) - a_{21} \det(a_{12}) \\ &= a_{11} a_{22} - a_{21} a_{12}. \end{aligned}$$

$n = 3$: Entwicklung nach der ersten Spalte

$$\begin{aligned} \det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} &= (+1) a_{11} \begin{vmatrix} \cancel{a_{11}} & \cancel{a_{12}} & \cancel{a_{13}} \\ \cancel{a_{21}} & a_{22} & a_{23} \\ \cancel{a_{31}} & a_{32} & a_{33} \end{vmatrix} \\ &\quad + (-1) a_{21} \begin{vmatrix} \cancel{a_{11}} & a_{12} & a_{13} \\ \cancel{a_{21}} & \cancel{a_{22}} & \cancel{a_{23}} \\ \cancel{a_{31}} & a_{32} & a_{33} \end{vmatrix} \\ &\quad + (+1) a_{31} \begin{vmatrix} \cancel{a_{11}} & a_{12} & a_{13} \\ \cancel{a_{21}} & a_{22} & a_{23} \\ \cancel{a_{31}} & \cancel{a_{32}} & \cancel{a_{33}} \end{vmatrix} \end{aligned}$$

$$\begin{aligned}
&= a_{11} \begin{vmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{vmatrix} - a_{21} \begin{vmatrix} a_{12} & a_{13} \\ a_{32} & a_{33} \end{vmatrix} + a_{31} \begin{vmatrix} a_{12} & a_{13} \\ a_{22} & a_{23} \end{vmatrix} \\
&= a_{11} (a_{22} a_{33} - a_{32} a_{23}) - a_{21} (a_{12} a_{33} - a_{32} a_{13}) \\
&\quad + a_{31} (a_{12} a_{23} - a_{22} a_{13}).
\end{aligned}$$

Beispiele 3.12:

① Wir berechnen die Determinante durch Entwicklung nach der zweiten Zeile:

$$\begin{aligned}
\det \begin{pmatrix} 2 & 3 & 5 \\ \mathbf{0} \mathbf{-4 -1} \\ 1 & -2 & 0 \end{pmatrix} &= (-1) \cdot \mathbf{0} \begin{vmatrix} 3 & 5 \\ -2 & 0 \end{vmatrix} + \mathbf{4} \begin{vmatrix} 2 & 5 \\ 1 & 0 \end{vmatrix} + (-1) \cdot (-1) \begin{vmatrix} 2 & 3 \\ 1 & -2 \end{vmatrix} \\
&= 4 \cdot (-5) + 1 \cdot (-7) = -27.
\end{aligned}$$

② Wir entwickeln die nachfolgende Determinante nach der ersten Spalte, da sie zwei Nullen als Elemente aufweist (man hätte aber auch eine andere Zeile oder Spalte mit zwei Nullen wählen können):

$$\det \begin{pmatrix} 1 & 2 & 0 & 4 \\ \mathbf{0} & 0 & 1 & 2 \\ \mathbf{3} & 3 & 2 & 1 \\ \mathbf{0} & 1 & 1 & 0 \end{pmatrix} = 1 \cdot \begin{vmatrix} 0 & 1 & 2 \\ 3 & 2 & 1 \\ 1 & 1 & 0 \end{vmatrix} + \mathbf{3} \begin{vmatrix} 2 & 0 & 4 \\ 0 & 1 & 2 \\ 1 & 1 & 0 \end{vmatrix} = 3 + 3 \cdot (-8) = -21.$$

③ Die Einheitsmatrix $I_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}$ hat als Determinante: $\det I_n = 1.$

④ Eine Dreiecksmatrix $D = \begin{pmatrix} \lambda_1 & & * \\ & \ddots & \\ 0 & & \lambda_n \end{pmatrix}$ hat als Determinante:

$$\det D = \lambda_1 \cdot \lambda_2 \cdot \dots \cdot \lambda_n.$$

□

➤ **Berechnung einer 3-reihigen Determinante nach Sarrus (Sarrussche Regel)**

⚠ Für den Spezialfall $n = 3$ (und nur für diesen Spezialfall !) kann man den Determinantenwert durch die auf Sarrus zurückgehende Regel berechnen:

$$\det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} \rightarrow$$

Die erste und zweite Spalte der Matrix A wird nochmals neben die Determinante gesetzt. Den Wert der Determinante erhält man dann, indem man die drei Diagonalprodukte addiert und die Antidiagonalprodukte davon subtrahiert:

$$\det A = \begin{array}{c} a_{11} a_{22} a_{33} + a_{12} a_{23} a_{31} + a_{13} a_{21} a_{32} \\ - a_{13} a_{22} a_{31} - a_{11} a_{23} a_{32} - a_{12} a_{21} a_{33}. \end{array}$$

Beispiel 3.13.

Wir berechnen die Determinante von $A = \begin{pmatrix} 1 & 2 & -1 \\ 1 & 0 & 1 \\ -1 & -2 & 1 \end{pmatrix}$ nach Sarrus:

$$\left| \begin{array}{ccc|cc} 1 & 2 & -1 & 1 & 2 \\ 1 & 0 & 1 & 1 & 0 \\ -1 & -2 & 1 & -1 & -2 \end{array} \right| :$$

$$\det A = 1 \cdot 0 \cdot 1 + 2 \cdot 1 \cdot (-1) + (-1) \cdot 1 \cdot (-2) - (-1) \cdot 0 \cdot (-1) - 1 \cdot 1 \cdot (-2) - 2 \cdot 1 \cdot 1 = -2 + 2 + 2 - 2 = 0. \quad \square$$

⊗ **Determinante der inversen Matrix**

Sei A eine invertierbare Matrix. Dann gibt es zu dieser Matrix eine Inverse A^{-1} mit der Eigenschaft

$$A \cdot A^{-1} = I_n$$

wenn I_n die Einheitsmatrix ist. Nach dem Multiplikationssatz für Determinanten (Regel 6) gilt dann

$$\det(A \cdot A^{-1}) = \det(A) \cdot \det(A^{-1}) = \det I_n = 1.$$

Damit folgt für den Determinantenwert der Inversen:

Satz: (Determinante der Inversen)

Ist A eine **invertierbare** Matrix, dann ist

$$\det A \neq 0$$

und die Determinante der Inversen A^{-1} ist gegeben durch

$$\det A^{-1} = \frac{1}{\det A}.$$

⊗ **Praktische Berechnung n -reihiger Determinanten**

Bei der praktischen Berechnung n -reihiger Determinanten wächst der Rechenaufwand mit zunehmender Ordnung rasch an. Denn aus einer n -reihigen Determinante entstehen durch Entwicklung $\prod_{k=3}^n k = 3 \cdot 4 \cdot \dots \cdot n$ zweireihige Unterdeterminanten bzw. $\prod_{k=4}^n k = 4 \cdot 5 \cdot \dots \cdot n$ 3-reihige Unterdeterminanten. Daher führt man gemäß der aufgelisteten Rechenregeln die Matrix auf eine einfachere Form über und berechnet dann erst den Determinantenwert.

Da sich der Wert der Determinante nicht ändert, wenn man zu einer Zeile (oder Spalte) ein Vielfaches einer anderen Zeile (oder Spalte) hinzu addiert, ist es zur Berechnung von $\det A$ (für $n > 3$) zweckmäßig, zuerst mit elementaren Umformungen möglichst viele Nullen in einer Spalte oder Zeile zu erzeugen. Nach dieser Spalte oder Zeile wird dann entwickelt. Die 3-reihigen Determinanten können auch nach der Sarrusschen Regel berechnet werden.

Beispiel 3.14. Wir subtrahieren die vierte Zeile von der dritten Zeile und entwickeln anschließend nach der zweiten Spalte. Bei der verbleibenden 3-reihigen Determinante subtrahieren wir von der zweiten und dritten Zeile jeweils die erste Zeile. Anschließend entwickeln wir nach der dritten Spalte:

$$\begin{aligned} \det \begin{pmatrix} 1 & 0 & 2 & 1 \\ 2 & 0 & 1 & 1 \\ -1 & 1 & 3 & 2 \\ 2 & 1 & -1 & 1 \end{pmatrix} &= \det \begin{pmatrix} 1 & \mathbf{0} & 2 & 1 \\ 2 & \mathbf{0} & 1 & 1 \\ -3 & \mathbf{0} & 4 & 1 \\ 2 & \mathbf{1} & -1 & 1 \end{pmatrix} = (+1) \det \begin{pmatrix} 1 & 2 & 1 \\ 2 & 1 & 1 \\ -3 & 4 & 1 \end{pmatrix} \\ &= \det \begin{pmatrix} 1 & 2 & \mathbf{1} \\ 1 & -1 & \mathbf{0} \\ -4 & 2 & \mathbf{0} \end{pmatrix} = (+1)(2 - 4) = -2. \quad \square \end{aligned}$$

Beispiel 3.15. Wir addieren vor der Berechnung der Determinante die erste Zeile zur dritten und das Zweifache der ersten Zeile zur vierten und entwickeln danach die 5-reihige Determinante nach der ersten Spalte. Bei der 4-reihigen Determinante subtrahieren wir von der ersten Spalte die dritte und entwickeln anschließend nach der letzten Zeile:

$$\begin{aligned} \det \begin{pmatrix} -1 & 1 & 0 & -2 & 0 \\ 0 & 2 & 1 & 1 & 4 \\ 1 & 2 & 4 & 3 & 2 \\ 2 & 1 & 0 & 0 & 1 \\ 0 & 4 & 0 & 4 & 0 \end{pmatrix} &= \det \begin{pmatrix} -1 & 1 & 0 & -2 & 0 \\ \mathbf{0} & 2 & 1 & 1 & 4 \\ \mathbf{0} & 3 & 4 & 1 & 2 \\ \mathbf{0} & 3 & 0 & -4 & 1 \\ \mathbf{0} & 4 & 0 & 4 & 0 \end{pmatrix} = (-1) \det \begin{pmatrix} 2 & 1 & 1 & 4 \\ 3 & 4 & 1 & 2 \\ 3 & 0 & -4 & 1 \\ 4 & 0 & 4 & 0 \end{pmatrix} \\ &= (-1) \det \begin{pmatrix} 1 & 1 & 1 & 4 \\ 2 & 4 & 1 & 2 \\ 7 & 0 & -4 & 1 \\ \mathbf{0} & \mathbf{0} & \mathbf{4} & \mathbf{0} \end{pmatrix} = (-1)(-4) \det \begin{pmatrix} 1 & 1 & 4 \\ 2 & 4 & 2 \\ 7 & 0 & 1 \end{pmatrix} = 4(-96) = -384. \quad \square \end{aligned}$$

► 3.2.4 Anwendungen von Determinanten

⊗ (1.) Lösen von linearen Gleichungssystemen: Cramersche Regel.

Cramersche Regel: Sei $A = (a_{ij})$ eine $(n \times n)$ -Matrix mit den Spaltenvektoren $(a^1, a^2, \dots, a^n)$ und $\det(A) \neq 0$. Dann ist die Lösung des linearen Gleichungssystems

$$A \cdot \vec{x} = \vec{b}$$

mit $\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$, $\vec{b} = \begin{pmatrix} b_1 \\ \vdots \\ b_n \end{pmatrix}$ gegeben durch

$$x_i = \frac{\det(a^1, \dots, a^{i-1}, \vec{b}, a^{i+1}, \dots, a^n)}{\det(A)},$$

Man ersetzt die i -te Spalte von A durch die rechte Seite $\vec{b}$ des LGS. Dann ist die i -te Komponente des Lösungsvektors $\vec{x}$ der Quotient der Determinante der so entstandenen Matrix und der Determinante von A .

Beispiel 3.16. Gesucht ist die Lösung des LGS

$$\begin{aligned} x_1 + x_2 &= 1 \\ x_2 + x_3 &= 1 \\ 3x_1 + 2x_2 + x_3 &= 0 \end{aligned}$$

Es ist $A = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 3 & 2 & 1 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ mit $\det A = 2 \neq 0$. Nach der Cramerschen Regel ergeben sich die Komponenten des Lösungsvektors aus

$$x_1 = \frac{1}{2} \det \begin{pmatrix} 1 & 1 & 0 \\ 1 & 1 & 1 \\ 0 & 2 & 1 \end{pmatrix} = -1; \quad x_2 = \frac{1}{2} \det \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 3 & 0 & 1 \end{pmatrix} = 2$$

$$x_3 = \frac{1}{2} \det \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 3 & 2 & 0 \end{pmatrix} = -1 \quad \Rightarrow \quad \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -1 \\ 2 \\ -1 \end{pmatrix}. \quad \square$$

⚠ Achtung: Man sollte allerdings die Cramersche Regel nicht zur Auflösung von linearen Gleichungssystemen mit vielen Unbekannten verwenden, da der Rechenaufwand erheblich wird. Die Formel ist aber nützlich für die Berechnung einzelner Unbekannter bzw. wenn das Gleichungssystem einen Parameter enthält.

⊗ (2.) Vektorprodukt

$\vec{a} \times \vec{b}$ zweier Vektoren $\vec{a}$ und $\vec{b}$:

In Kapitel 2.2.3 wird eine Formel für das Vektorprodukt $\vec{c} = \vec{a} \times \vec{b}$ zweier Vektoren $\vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$ entwickelt.

Eine **einfache Merkregel** für diese Formel erhält man, wenn man sie formal durch eine 3-reihige Determinante darstellt: Sind $\vec{e}_1, \vec{e}_2, \vec{e}_3$ die Einheitsvektoren, dann ist:

$$\vec{a} \times \vec{b} = \begin{vmatrix} \vec{e}_1 & a_x & b_x \\ \vec{e}_2 & a_y & b_y \\ \vec{e}_3 & a_z & b_z \end{vmatrix} = \vec{e}_1 \begin{vmatrix} a_y & b_y \\ a_z & b_z \end{vmatrix} - \vec{e}_2 \begin{vmatrix} a_x & b_x \\ a_z & b_z \end{vmatrix} + \vec{e}_3 \begin{vmatrix} a_x & b_x \\ a_y & b_y \end{vmatrix}.$$

Beispiel 3.17. Berechnung des Drehmomentes $\vec{M}$ auf einen Körper mit

$$\begin{aligned} \vec{r} &= \begin{pmatrix} 1 \\ -2 \\ 5 \end{pmatrix} m \text{ und } \vec{F} = \begin{pmatrix} 10 \\ 20 \\ 50 \end{pmatrix} N: \quad \vec{M} = \vec{r} \times \vec{F} = \begin{vmatrix} \vec{e}_1 & 1 & 10 \\ \vec{e}_2 & -2 & 20 \\ \vec{e}_3 & 5 & 50 \end{vmatrix} \\ &= \vec{e}_1 \begin{vmatrix} -2 & 20 \\ 5 & 50 \end{vmatrix} - \vec{e}_2 \begin{vmatrix} 1 & 10 \\ 5 & 50 \end{vmatrix} + \vec{e}_3 \begin{vmatrix} 1 & 10 \\ -2 & 20 \end{vmatrix} = \begin{pmatrix} -200 \\ 0 \\ 40 \end{pmatrix} [Nm]. \quad \square \end{aligned}$$

3.3

3.3 Lösbarkeit von linearen Gleichungssystemen

Die Determinante entscheidet bei linearen Gleichungssystemen mit quadratischen Matrizen, ob es für eine beliebige rechte Seite lösbar ist. Für LGS mit nicht quadratischen Matrizen ist die Determinante nicht definiert. Wir führen den Rang einer Matrix ein, um in diesem Fall Aussagen über die Lösbarkeit bei konkreter gegebener rechter Seite treffen zu können.

⊗ 3.3.1 Lineare Gleichungssysteme, Rang

Wir kommen nun auf die Fragestellung zurück, unter welchen Bedingungen ein lineares Gleichungssystem

$$A \vec{x} = \vec{b}$$

lösbar ist, wenn A keine quadratische, sondern allgemeiner eine $(m \times n)$ -Matrix ist. Und wenn es lösbar ist, ob es eine eindeutige Lösung oder unendlich viele Lösungen besitzt. Zur Charakterisierung von LGS bzw. linearen Abbildungen führt man die folgenden Begriffe ein:

Definition: Sei A eine $(m \times n)$ -Matrix.

- (1) Die Menge aller Lösungen des homogenen LGS $A\vec{x} = \vec{0}$ bezeichnet man als **Kern** oder **Nullraum** von A .
- (2) Die Menge der Vektoren $A\vec{x}$, $\vec{x} \in \mathbb{R}^n$, heißt **Spaltenraum** von A .
- (3) Die maximale Anzahl linear unabhängiger Spalten von A heißt der **Spaltenrang** von A .
- (4) Die maximale Anzahl linear unabhängiger Zeilen von A heißt der **Zeilenrang** von A .

Bemerkung: Da jede $(m \times n)$ -Matrix einer linearen Abbildung $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ entspricht, sind sowohl der Kern als auch der Spaltenraum von A Vektorräume. Der Kern ist ein Untervektorraum von $\mathbb{R}^n$ und der Spaltenraum von $\mathbb{R}^m$.

Folgender Satz gibt Auskunft darüber, wie man den Spalten- und Zeilenrang einer Matrix bestimmen kann.

Satz: Elementare Zeilen- und Spaltenumformungen einer Matrix lassen sowohl den Zeilenrang als auch den Spaltenrang unverändert.

Beispiel 3.18. Wir berechnen den Zeilen- und Spaltenrang der unten angegebenen Matrix, indem wir durch elementare Zeilenumformungen A zunächst auf obere Dreiecksform bringen und dann die erste und zweite Zeile ausräumen:

$$\begin{aligned}
 A &= \begin{pmatrix} 1 & 2 & -1 & 1 & 2 \\ 1 & 4 & -3 & -1 & -1 \\ 2 & 6 & -4 & 0 & 1 \\ 1 & 0 & 1 & 3 & 5 \\ -1 & 4 & -1 & 1 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 2 & -1 & 1 & 2 \\ 0 & 2 & -2 & -2 & -3 \\ 0 & 2 & -2 & -2 & -3 \\ 0 & -2 & 2 & 2 & 3 \\ 0 & 6 & -2 & 2 & 3 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 2 & -1 & 1 & 2 \\ 0 & 2 & -2 & -2 & -3 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 4 & 8 & 12 \end{pmatrix} \\
 &\rightarrow \begin{pmatrix} 1 & 0 & 1 & 3 & 5 \\ 0 & 1 & -1 & -1 & -\frac{3}{2} \\ 0 & 0 & 1 & 2 & 3 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 0 & 1 & 2 \\ 0 & 1 & 0 & 1 & \frac{3}{2} \\ 0 & 0 & 1 & 2 & 3 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} = A_2.
 \end{aligned}$$

Man erkennt, dass A_2 den Zeilenrang = Spaltenrang = 3 besitzt. Demnach hat auch A den Zeilenrang = Spaltenrang = 3. Allgemein gilt: $\square$

Satz / Definition: Der Spaltenrang und der Zeilenrang einer $(m \times n)$ -Matrix A sind immer gleich. Man spricht daher nur vom **Rang einer Matrix**:

$$\text{Rang}(A) := \text{Spaltenrang}(A) = \text{Zeilenrang}(A).$$

Der Rang einer Matrix kann nun als Kriterium herangezogen werden, um zu entscheiden, ob ein LGS $A\vec{x} = \vec{b}$ lösbar ist:

Satz: Gegeben sei ein LGS

$$A\vec{x} = \vec{b} \quad (*)$$

mit einer $(m \times n)$ -Matrix A und rechten Seite $\vec{b} \in \mathbb{R}^n$. $(A | \vec{b})$ bezeichnet die Matrix, die aus A entsteht, indem der Vektor $\vec{b}$ als Spalte zu A hinzugefügt wird. Ist

$$\text{Rang}(A) = \text{Rang}(A | \vec{b}),$$

dann ist das LGS $(*)$ lösbar.

Folgerungen:

Ist $A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$ eine $(m \times n)$ -Matrix. Dann gilt:

- (1) Das homogene LGS $A\vec{x} = \vec{0}$ besitzt genau dann von Null verschiedene Lösungen, wenn der Rang $(A) < n$.
- (2) $A\vec{x} = \vec{0}$ besitzt stets von Null verschiedene Lösungen, wenn $m < n$, d.h. die Anzahl der Gleichungen kleiner ist als die Anzahl der Unbekannten.
- (3) Ist A eine quadratische Matrix, so ist nach der Determinantenregel 4 die Aussage (1) gleichbedeutend mit

$$\det A = 0.$$

- (4) Das inhomogene LGS $A\vec{x} = \vec{b}$ ist genau dann lösbar, wenn

$$\text{Rang}(A) = \text{Rang}(A | \vec{b}).$$

- (5) Das inhomogene LGS $A\vec{x} = \vec{b}$ ist genau dann **eindeutig lösbar**, wenn es lösbar ist und $A\vec{x} = \vec{0}$ nur den Nullvektor als Lösung besitzt:

$$\text{Rang } A = \text{Rang}(A|\vec{b}) = n.$$

- (6) Ist A eine quadratische Matrix, so ist nach der Determinantenregel 4 die Aussage (5) gleichbedeutend mit:

$$A\vec{x} = \vec{b} \text{ ist eindeutig lösbar} \iff \det(A) \neq 0.$$

Beispiele 3.19 (Mit MAPLE-Worksheet).

- ① Gegeben ist das LGS $A\vec{x} = \vec{b}$ mit der Matrix

$$A = \begin{pmatrix} -3 & 0 & 6 & 0 \\ 1 & 1 & -2 & 5 \\ 1 & 0 & -2 & 0 \\ -2 & -2 & 4 & -10 \end{pmatrix} \text{ und dem Vektor } \vec{b} = \begin{pmatrix} -3 \\ 2 \\ 1 \\ -4 \end{pmatrix}.$$

Wir vergleichen den Rang von A mit dem Rang der um $\vec{b}$ erweiterten Matrix. Durch Zeilenumformungen mit der erweiterten Matrix erhalten wir

$$\left(\begin{array}{cccc|c} -3 & 0 & 6 & 0 & -3 \\ 1 & 1 & -2 & 5 & 2 \\ 1 & 0 & -2 & 0 & 1 \\ -2 & -2 & 4 & -10 & -4 \end{array} \right) \hookrightarrow \left(\begin{array}{cccc|c} -1 & 0 & 2 & 0 & -1 \\ 0 & 1 & 0 & 5 & 1 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right).$$

Die Anzahl der linear unabhängigen Zeilen von A beträgt 2, d.h. $\text{Rang}(A) = 2$. Die Anzahl der linear unabhängigen Zeilen der erweiterten Matrix $(A|\vec{b})$ beträgt ebenfalls 2, d.h. $\text{Rang}(A|\vec{b}) = 2$. Beide Ränge stimmen überein und das LGS ist daher lösbar.

Da $\text{Rang}(A) = 2 < 4$, ist das LGS nicht eindeutig lösbar und die Lösung besitzt freie Parameter t_1 und t_2 :

$$x_4 = t_1, \quad x_3 = t_2, \quad x_2 = 1 - 5t_1, \quad x_1 = 1 + 2t_2.$$

- ② Für $\vec{b} = \begin{pmatrix} -3 \\ 2 \\ 1 \\ -2 \end{pmatrix}$ ist das LGS $A\vec{x} = \vec{b}$ nicht lösbar, da der Rang der erweiterten Matrix $(A|\vec{b}) = 3$. □

Wir fassen die Ergebnisse für quadratische, lineare Gleichungssysteme im folgenden Satz zusammen:

Satz: (Fundamentalsatz für quadratische LGS). Gegeben ist das LGS für die Unbekannten $(x_1, x_2, \dots, x_n)$:

$$\begin{aligned} a_{11}x_1 + \dots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + \dots + a_{2n}x_n &= b_2 \\ &\vdots \\ a_{n1}x_1 + \dots + a_{nn}x_n &= b_n \end{aligned}$$

Definieren wir die Matrix $A = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix}$, die rechte Seite $\vec{b} = \begin{pmatrix} b_1 \\ \vdots \\ b_n \end{pmatrix}$ und den Vektor $\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$, so lautet das LGS $A\vec{x} = \vec{b}$. Dann sind die folgenden Aussagen äquivalent:

- (1) Das LGS $A\vec{x} = \vec{b}$ besitzt eine eindeutig bestimmte Lösung.
- (2) A ist invertierbar.
- (3) $\det A \neq 0$.
- (4) Die Lösung des LGS ist gegeben durch $\vec{x} = A^{-1}\vec{b}$.

Beispiel 3.20 (Mit MAPLE-Worksheet). Gegeben ist das inhomogene LGS

$$A\vec{x} = \vec{b} \text{ mit der Matrix } A = \begin{pmatrix} 1 & 2 & 0 \\ 1 & 7 & 4 \\ 3 & 13 & 4 \end{pmatrix} \text{ und den rechten Seiten } \vec{b}_1 = \begin{pmatrix} -4 \\ 3 \\ 1 \end{pmatrix}, \vec{b}_2 = \begin{pmatrix} 1 \\ 8 \\ 8 \end{pmatrix}, \vec{b}_3 = \begin{pmatrix} 1 \\ -4 \\ 0 \end{pmatrix}.$$

- (i) Das LGS hat für jeden Vektor $\vec{b}$ eine eindeutig bestimmte Lösung, da die Determinante der Matrix $\det(A) = -8$ ungleich Null ist.
- (ii) Da $\det(A) \neq 0$ existiert die inverse Matrix

$$A^{-1} = \begin{pmatrix} 3 & 1 & -1 \\ -1 & -\frac{1}{2} & \frac{1}{2} \\ 1 & \frac{7}{8} & -\frac{5}{8} \end{pmatrix}.$$

(iii) Die eindeutige Lösung des LGS ist gegeben durch

$$\vec{x} = A^{-1} \vec{b}.$$

Es gilt für die drei Inhomogenitäten

$$\vec{x}_1 = A^{-1} \vec{b}_1 = \begin{pmatrix} 3 & 1 & -1 \\ -1 & -\frac{1}{2} & \frac{1}{2} \\ 1 & \frac{7}{8} & -\frac{5}{8} \end{pmatrix} \begin{pmatrix} -4 \\ 3 \\ 1 \end{pmatrix} = \begin{pmatrix} -10 \\ 3 \\ -2 \end{pmatrix}$$

$$\vec{x}_2 = A^{-1} \vec{b}_2 = \begin{pmatrix} 3 & 1 & -1 \\ -1 & -\frac{1}{2} & \frac{1}{2} \\ 1 & \frac{7}{8} & -\frac{5}{8} \end{pmatrix} \begin{pmatrix} 1 \\ 8 \\ 8 \end{pmatrix} = \begin{pmatrix} 3 \\ -1 \\ 3 \end{pmatrix}$$

$$\vec{x}_3 = A^{-1} \vec{b}_3 = \begin{pmatrix} 3 & 1 & -1 \\ -1 & -\frac{1}{2} & \frac{1}{2} \\ 1 & \frac{7}{8} & -\frac{5}{8} \end{pmatrix} \begin{pmatrix} 1 \\ -4 \\ 0 \end{pmatrix} = \begin{pmatrix} -1 \\ 1 \\ -\frac{5}{2} \end{pmatrix}$$

Der Vorteil der Lösung des LGS über die inverse Matrix gegenüber der Lösung mit dem Gauß-Algorithmus besteht darin, dass die Inverse nur einmal berechnet wird und das Lösen des LGS für verschiedene Inhomogenitäten dann nur noch eine Matrix-Vektor-Multiplikation darstellt. $\square$

► 3.3.2 Anwendungen

Anwendungsbeispiel 3.21 (Lineare Unabhängigkeit von Vektoren).

n Vektoren des n -dimensionalen Vektorraums $\mathbb{R}^n$ bilden nach Kapitel 2.4.5 eine Basis von $\mathbb{R}^n$ genau dann, wenn sie **linear unabhängig** sind. Die n Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{R}^n$ bilden also eine Basis, wenn

$$\lambda_1 \vec{a}_1 + \lambda_2 \vec{a}_2 + \dots + \lambda_n \vec{a}_n = \vec{0}$$

nur durch $\lambda_1 = \lambda_2 = \dots = \lambda_n = 0$ eindeutig lösbar ist. Nach dem Fundamentalsatz über LGS ist dies aber gleichbedeutend, dass

$$\det(A) = \det(\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n) \neq 0.$$

Man erhält dadurch ein sehr einfaches Kriterium, um nachzuprüfen ob n Vektoren des $\mathbb{R}^n$ eine Basis bilden:

Satz: Für n Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \in \mathbb{R}^n$ gilt:

$\det A = \det(\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n) \neq 0 \Leftrightarrow (\vec{a}_1, \dots, \vec{a}_n)$ ist eine Basis von $\mathbb{R}^n$.

Beispiel 3.22 (Zahlenbeispiel, mit MAPLE-Worksheet).

Gegeben sind die Vektoren

$$\vec{a}_1 = \begin{pmatrix} 4 \\ 3 \\ 2 \\ 1 \end{pmatrix}, \vec{a}_2 = \begin{pmatrix} 0 \\ 1 \\ 5 \\ 4 \end{pmatrix}, \vec{a}_3 = \begin{pmatrix} -1 \\ -1 \\ 0 \\ 1 \end{pmatrix}, \vec{a}_4 = \begin{pmatrix} 3 \\ 5 \\ 0 \\ 1 \end{pmatrix} \text{ und } \vec{b} = \begin{pmatrix} 0 \\ 1 \\ 0 \\ 1 \end{pmatrix}.$$

Zeigen Sie, dass die Vektoren $\vec{a}_1, \vec{a}_2, \vec{a}_3$ und $\vec{a}_4$ eine Basis des $\mathbb{R}^4$ bilden und stellen Sie den Vektor $\vec{b}$ als Linearkombination dieser Basis dar.

Zunächst prüfen wir, dass die 4 Vektoren linear unabhängig sind, indem wir von der Matrix $A = (\vec{a}_1, \vec{a}_2, \vec{a}_3, \vec{a}_4)$ die Determinante bestimmen:

$$\det(A) = \begin{vmatrix} 4 & 0 & -1 & 3 \\ 3 & 1 & -1 & 5 \\ 2 & 5 & 0 & 0 \\ 1 & 4 & 1 & 1 \end{vmatrix} = 62.$$

Da die Determinante ungleich Null ist, sind die Vektoren linear unabhängig und bilden eine Basis von $\mathbb{R}^4$. Der Vektor $\vec{b}$ lässt sich damit als Linearkombination der Vektoren $(\vec{a}_1, \vec{a}_2, \vec{a}_3, \vec{a}_4)$ darstellen. Indem wir das LGS

$$\begin{pmatrix} 4 & 0 & -1 & 3 & | & 0 \\ 3 & 1 & -1 & 5 & | & 1 \\ 2 & 5 & 0 & 0 & | & 0 \\ 1 & 4 & 1 & 1 & | & 1 \end{pmatrix}$$

mit dem Gauß-Algorithmus lösen, erhalten wir als Lösung

$$\left(-\frac{5}{31}, \frac{2}{31}, \frac{16}{31}, \frac{12}{31} \right).$$

$$\Rightarrow \vec{b} = -\frac{5}{31} \vec{a}_1 + \frac{2}{31} \vec{a}_2 + \frac{16}{31} \vec{a}_3 + \frac{12}{31} \vec{a}_4$$

□

Anwendungsbeispiel 3.23 (Kreis durch 3 Punkte).

In CAD-Systemen werden ebene Flächen zumeist durch Geraden- und Kreisstücke zusammengesetzt. Die Erfassung der Geometrie erfolgt meist interaktiv per Mouse-Klick, indem für die Geometrie charakteristische Punkte eingegeben und diese durch Geraden oder Kreissegmente verbunden werden. Ein Geradenstück wird durch zwei Punkte festgelegt; ein Kreissegment durch die Angabe von drei Punkten, falls die Punkte nicht auf einer Geraden liegen.

Wir behandeln im Folgenden die Fragestellung, ob durch drei vorgegebene Punkte (x_1, y_1) , (x_2, y_2) , (x_3, y_3) ein Kreis gelegt werden kann, und wenn ja, welches die Mittelpunktskoordinaten und der Radius des zugehörigen Kreises sind. Die allgemeine Kreisgleichung lautet

$$(x - x_0)^2 + (y - y_0)^2 = R^2$$

bzw.

$$A(x^2 + y^2) + Bx + Cy + D = 0.$$

Da der Kreis durch die gegebenen Punkte gehen soll, muss noch zusätzlich gelten:

$$\begin{aligned} A(x_1^2 + y_1^2) + Bx_1 + Cy_1 + D &= 0 \\ A(x_2^2 + y_2^2) + Bx_2 + Cy_2 + D &= 0 \\ A(x_3^2 + y_3^2) + Bx_3 + Cy_3 + D &= 0. \end{aligned}$$

Dies ist ein homogenes LGS für die Größen (A, B, C, D) mit der Koeffizientenmatrix

$$M = \begin{pmatrix} x^2 + y^2 & x & y & 1 \\ x_1^2 + y_1^2 & x_1 & y_1 & 1 \\ x_2^2 + y_2^2 & x_2 & y_2 & 1 \\ x_3^2 + y_3^2 & x_3 & y_3 & 1 \end{pmatrix}.$$

Damit das homogene LGS nicht nur die Null-Lösung besitzt, muss

$\det M = 0$

sein. Entwickelt man die Determinante nach der ersten Zeile, erhält man

$$\begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix} (x^2 + y^2) - \begin{vmatrix} x_1^2 + y_1^2 & y_1 & 1 \\ x_2^2 + y_2^2 & y_2 & 1 \\ x_3^2 + y_3^2 & y_3 & 1 \end{vmatrix} x + \begin{vmatrix} x_1^2 + y_1^2 & x_1 & 1 \\ x_2^2 + y_2^2 & x_2 & 1 \\ x_3^2 + y_3^2 & x_3 & 1 \end{vmatrix} y - \begin{vmatrix} x_1^2 + y_1^2 & x_1 & y_1 \\ x_2^2 + y_2^2 & x_2 & y_2 \\ x_3^2 + y_3^2 & x_3 & y_3 \end{vmatrix} = 0.$$

Die 3-reihigen Unterdeterminanten sind also gerade die Koeffizienten in der Kreisgleichung. Insbesondere folgt aus dieser Darstellung, dass

$$K := \begin{vmatrix} x_1 & y_1 & 1 \\ x_2 & y_2 & 1 \\ x_3 & y_3 & 1 \end{vmatrix} \neq 0$$

sein muss, damit der Term $(x^2 + y^2)$ in der Kreisgleichung vorkommt. Für $K = 0$ liegen die 3 Punkte also auf einer Geraden.

Beispiel 3.24 (Mit MAPLE-Worksheet). Gesucht ist die Kreisgleichung

$$A(x^2 + y^2) + Bx + Cy + D = 0$$

durch die Punkte $(0, 0)$, $(1, 3)$ und $(2, -1)$. Wir setzen die drei Punkte in die Kreisgleichung ein

$$\begin{array}{rccccccccccc} A \cdot & 0 & + & B \cdot 0 & + & C \cdot & 0 & + & D & = & 0 \\ A \cdot & 10 & + & B \cdot 1 & + & C \cdot & 3 & + & D & = & 0 \\ A \cdot & 5 & + & B \cdot 2 & + & C \cdot & (-1) & + & D & = & 0. \end{array}$$

und erhalten zusammen mit der Kreisgleichung die Matrix

$$M = \begin{pmatrix} x^2 + y^2 & x & y & 1 \\ 0 & 0 & 0 & 1 \\ 10 & 1 & 3 & 1 \\ 5 & 2 & -1 & 1 \end{pmatrix}.$$

Setzen wir $\det(M) = 0$, erhalten wir die Kreisgleichung

$$7x^2 + 7y^2 - 25x - 15y = 0.$$

Die Punkte liegen auf dem Kreis, dessen Scheitelgleichung gegeben ist durch

$$\left(x - \frac{25}{14}\right)^2 + \left(y - \frac{15}{14}\right)^2 = \frac{425}{98}.$$

Der Kreis besitzt den Radius $R = \sqrt{\frac{425}{98}}$ und hat den Kreismittelpunkt

$$(x_0, y_0) = (4, -5).$$

□

Anwendungsbeispiel 3.25 (Eigenfrequenzen eines gekoppelten Systems).

Gegeben sind zwei *Fadenpendel* (Länge l) an deren Enden zwei Massen ($m_1 = m_2 = m$) angebracht sind. Die Massen werden durch eine Feder (Federkonstante D) gekoppelt (siehe Abb. 3.3).

Ist $\varphi_1(t)$ die Winkelauslenkung des linken und $\varphi_2(t)$ die Winkelauslenkung des rechten Pendels, so lauten die Newtonschen Bewegungsgleichungen für die Massen m_1 und m_2 unter Vernachlässigung von Reibung und für kleine

Auslenkungen:

$$\begin{aligned} m_1 l \ddot{\varphi}_1(t) &= -m_1 g \varphi_1(t) + D l (\varphi_2(t) - \varphi_1(t)) \\ m_2 l \ddot{\varphi}_2(t) &= -m_2 g \varphi_2(t) + D l (\varphi_1(t) - \varphi_2(t)). \end{aligned}$$

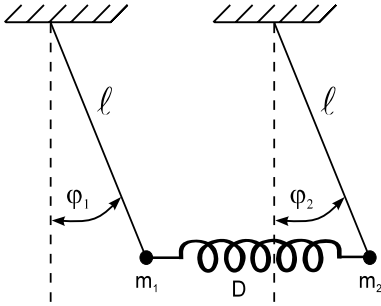


Abb. 3.3. Gekoppelte Pendel

Da $m_1 = m_2 = m$, unterdrücken wir im Folgenden den Index und erhalten mit den Abkürzungen $w_0^2 := \frac{g}{l}$ und $\Delta^2 := \frac{D}{m}$

$$\left. \begin{aligned} \ddot{\varphi}_1(t) + (w_0^2 + \Delta^2) \varphi_1(t) &= \Delta^2 \varphi_2(t) \\ \ddot{\varphi}_2(t) + (w_0^2 + \Delta^2) \varphi_2(t) &= \Delta^2 \varphi_1(t) \end{aligned} \right\} (*)$$

Dies ist ein gekoppeltes System für die Auslenkungen φ_1 und φ_2 . Gesucht sind die Schwingungsformen, bei denen beide Pendel mit gleicher Frequenz schwingen.

Da beide Pendel Schwingungen mit gleicher Frequenz aber gegebenenfalls unterschiedlichen Auslenkungen vollführen sollen, wählen wir für die Winkel den Ansatz

$$\begin{aligned} \varphi_1(t) &= A \cos(wt) \\ \varphi_2(t) &= B \cos(wt). \end{aligned}$$

Gesucht sind also die Amplituden A, B der Schwingungen sowie die Schwingungsfrequenzen w .

Setzen wir $\varphi_1(t)$ und $\varphi_2(t)$ bzw. $\ddot{\varphi}_1(t)$ und $\ddot{\varphi}_2(t)$ in die Gleichungen $(*)$ ein, folgt

$$\begin{aligned} -A w^2 \cos(wt) + (w_0^2 + \Delta^2) A \cos(wt) &= \Delta^2 B \cos(wt) \\ -B w^2 \cos(wt) + (w_0^2 + \Delta^2) B \cos(wt) &= \Delta^2 A \cos(wt). \end{aligned}$$

Dies ist ein LGS für die Amplituden A und B . Mit der Matrixschreibweise bringen wir es auf die Form

$$\begin{pmatrix} -w^2 + w_0^2 + \Delta^2 & -\Delta^2 \\ -\Delta^2 & -w^2 + w_0^2 + \Delta^2 \end{pmatrix} \begin{pmatrix} A \\ B \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}.$$

Damit dieses homogene LGS nicht nur die Null-Lösung als einzige (eindeutige) Lösung besitzt, muss nach dem Fundamentalsatz für LGS die Determinante der Koeffizientenmatrix gleich Null sein:

$$\det \begin{pmatrix} -w^2 + w_0^2 + \Delta^2 & -\Delta^2 \\ -\Delta^2 & -w^2 + w_0^2 + \Delta^2 \end{pmatrix} = (-w^2 + w_0^2 + \Delta^2)^2 - (\Delta^2)^2 = 0$$

$$\Rightarrow w_{1/2} = \pm \sqrt{w_0^2 + 2\Delta^2} \quad \text{oder} \quad w_{3/4} = \pm w_0.$$

Da physikalisch nur die positiven Frequenzen von Interesse sind, gibt es nur zwei Schwingungsfrequenzen, mit denen die Pendel schwingen können

$$w_1 = +\sqrt{w_0^2 + 2\Delta^2} \quad \text{und} \quad \omega_3 = \omega_0.$$

Diskussion:

(1) Für $w = w_0$ erhalten wir das LGS

$$\begin{pmatrix} \Delta^2 & -\Delta^2 \\ -\Delta^2 & \Delta^2 \end{pmatrix} \begin{pmatrix} A \\ B \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \Rightarrow A = B.$$

Dies bedeutet, dass Pendel (1) und (2) in die gleiche Richtung ausgelenkt werden müssen, um mit der Frequenz w_0 zu schwingen (gleichphasiger Fall).

(2) Für $w = \sqrt{w_0^2 + 2\Delta^2}$ erhalten wir das LGS

$$\begin{pmatrix} -\Delta^2 & -\Delta^2 \\ -\Delta^2 & -\Delta^2 \end{pmatrix} \begin{pmatrix} A \\ B \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \Rightarrow A = -B.$$

Dies bedeutet, dass Pendel (1) und (2) in entgegengesetzter Richtung ausgelenkt werden müssen, um mit der Frequenz $\sqrt{w_0^2 + 2\Delta^2}$ zu schwingen (gegenphasiger Fall). $\square$

MAPLE-Worksheets zu Kapitel 3



Für Kapitel 3 steht ein MAPLE-Worksheet zur Verfügung, mit dem sowohl das Arbeiten mit [Matrizen und Determinanten](#) als auch die Anwendungen bearbeitet werden können.

- [Matrizen und Determinanten](#)
- [MAPLE-Lösungen zu den Aufgaben](#)

3.4 Aufgaben zu Matrizen und Determinanten

3.1 Gegeben sind die Matrizen

$$A = \begin{pmatrix} 3 & 1 \\ 5 & 2 \\ 2 & 2 \end{pmatrix}, \quad B = \begin{pmatrix} 5 & 2 & 1 \\ 2 & 3 & -1 \end{pmatrix}, \quad C = \begin{pmatrix} 0 & 10 \\ 1 & 3 \\ -4 & 1 \end{pmatrix}$$

Man berechne

$$\begin{array}{lll} \text{a) } 2A + 3C & \text{b) } A - 2B^t & \text{c) } 3A^t + 4B \\ \text{d) } A + B^t + C & \text{e) } A^t - B + C^t & \text{f) } (A^t - B + C^t)^t \end{array}$$

3.2 Man berechne $A^2 = A \cdot A$, $B^2 = B \cdot B$, $A \cdot B$, $B \cdot A$ (soweit diese existieren) und zeige, dass in der Regel $A \cdot B \neq B \cdot A$.

$$\text{a) } A = \begin{pmatrix} 2 & 3 & 0 \\ 0 & 4 & 1 \\ 0 & 0 & 1 \end{pmatrix}, \quad B = \begin{pmatrix} 1 & 0 & 2 \\ -1 & 2 & -1 \\ 1 & -3 & 0 \end{pmatrix}$$

$$\text{b) } A = \begin{pmatrix} 2 & 3 & 0 & 4 \\ 0 & 1 & 0 & 1 \end{pmatrix}, \quad B = \begin{pmatrix} 0 & 2 \\ -1 & 2 \\ 5 & 3 \\ 1 & 4 \end{pmatrix}$$

3.3 Bestimmen Sie mit dem Gauß-Jordan-Verfahren die Inversen der folgenden Matrizen

$$A = \begin{pmatrix} 4 & 5 & -1 \\ 2 & 0 & 1 \\ 3 & 1 & 0 \end{pmatrix}, \quad B = \begin{pmatrix} 4 & 1 \\ -1 & -1 \end{pmatrix}, \quad C = \begin{pmatrix} 1 & -1 & 1 & -1 \\ -1 & 1 & 2 & 1 \\ -2 & 1 & -1 & 0 \\ 0 & 0 & 1 & 1 \end{pmatrix}$$

3.4 Überprüfen Sie obige Ergebnisse mit MAPLE.

3.5 Eine lineare Abbildung $f: \mathbb{R}^4 \rightarrow \mathbb{R}^3$ ist definiert durch die Bilder der Basisvektoren $f(\vec{e}_1) = (3, 5, 4)$, $f(\vec{e}_2) = (2, 1, 5)$, $f(\vec{e}_3) = (-1, 2, 1)$, $f(\vec{e}_4) = \vec{0}$. Erstellen Sie die Abbildungsmatrix und bestimmen Sie das Bild des Punktes $(3, 5, 1, 2)$.

3.6 Man berechne die Determinanten der folgenden zweireihigen Matrizen

$$\text{a) } A = \begin{pmatrix} 3 & 1 \\ 1 & 3 \end{pmatrix} \quad \text{b) } B = \begin{pmatrix} 4 & 2 \\ 2 & 1 \end{pmatrix} \quad \text{c) } C = \begin{pmatrix} 2 & 3 \\ \lambda & 2\lambda \end{pmatrix}$$

3.7 Man prüfe Determinantenregeln (1) und (2) bei der Matrix $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ nach.

3.8 Welchen Wert besitzen die folgenden dreireihigen Determinanten

$$\text{a) } \begin{vmatrix} 5 & 4 & 0 \\ 2 & 2 & 1 \\ 1 & 4 & 2 \end{vmatrix} \quad \text{b) } \begin{vmatrix} -2 & 8 & 1 \\ 2 & 3 & 5 \\ -1 & 5 & 2 \end{vmatrix} \quad \text{c) } \begin{vmatrix} 3 & 4 & 2 \\ 5 & -1 & 0 \\ -2 & 2 & 3 \end{vmatrix}$$

3.9 Für welche reellen Parameter λ verschwinden die Determinanten

$$\text{a) } \begin{vmatrix} 1 - \lambda & -2 \\ 1 & -2 - \lambda \end{vmatrix} \quad \text{b) } \begin{vmatrix} 1 - \lambda & 2 & 1 \\ 0 & 3 - \lambda & 2 \\ 0 & 0 & 2 - \lambda \end{vmatrix}$$

3.10 Man berechne die Determinanten von

$$A = \begin{pmatrix} 1 & 0 & 3 & 4 \\ -2 & 1 & 0 & 3 \\ 1 & 4 & 1 & 5 \\ 0 & 2 & 2 & 0 \end{pmatrix} \quad \text{und} \quad B = \begin{pmatrix} 1 & 0 & 5 & 3 & 4 \\ 1 & 2 & 2 & 1 & 0 \\ 0 & 1 & 3 & 4 & 1 \\ -4 & 0 & -1 & 2 & 3 \\ -2 & 1 & 0 & 0 & 0 \end{pmatrix}$$

3.11 Man bestimme mit der Cramerschen Regel die Lösung des LGS

$$\begin{array}{rrcr} x_1 & + & 2x_2 & = & 3 \\ x_1 & + & 7x_2 & + & 4x_3 & = & 18 \\ 3x_1 & + & 13x_2 & + & 4x_3 & = & 30 \end{array}$$

3.12 Sind die Matrizen

$$A = \begin{pmatrix} 4 & 5 & -1 \\ 2 & 0 & 1 \\ 3 & 1 & 0 \end{pmatrix}, \quad B = \begin{pmatrix} 3 & 4 & 2 \\ 1 & 5 & 3 \\ 0 & 1 & 0 \end{pmatrix}, \quad C = \begin{pmatrix} 1 & 1 & 2 \\ 3 & 2 & 4 \\ 2 & -5 & -1 \end{pmatrix}$$

reguläre Matrizen? Bestimmen Sie gegebenenfalls die Inverse.

3.13 Man bestimmen den Rang der folgenden Matrizen

$$A = \begin{pmatrix} 3 & 4 & 2 \\ 1 & 5 & 3 \\ 0 & 1 & 0 \end{pmatrix}, \quad B = \begin{pmatrix} 2 & 1 & 0 & 4 \\ 3 & 1 & 0 & 2 \\ 2 & 0 & 4 & 1 \end{pmatrix},$$

$$C = \begin{pmatrix} 1 & 1 & 2 & -3 & 0 & 1 \\ -2 & 3 & 1 & 0 & 1 & 2 \\ -5 & 10 & 5 & -3 & 3 & 1 \\ -1 & 4 & 3 & -3 & 1 & 1 \end{pmatrix}, \quad D = \begin{pmatrix} 1 & 1 & 2 \\ 3 & 2 & 4 \\ 2 & -5 & -1 \end{pmatrix}.$$

3.14 Man zeige, dass das LGS

$$\begin{array}{rrcr} x_1 & + & x_2 & + & x_3 & = & 0 \\ x_1 & + & 2x_2 & + & 4x_3 & = & 5 \\ x_1 & + & 3x_2 & + & 9x_3 & = & 12 \end{array}$$

eindeutig lösbar ist und bestimme diese Lösung.

3.15 Gegeben ist die Matrix

$$A = \begin{pmatrix} 3 & -1 & 2 \\ 6 & 2 & 5 \\ -2 & 6 & 0 \end{pmatrix} \quad \text{sowie der Vektor } \vec{b} = \begin{pmatrix} 2 \\ -2 \\ 1 \end{pmatrix}.$$

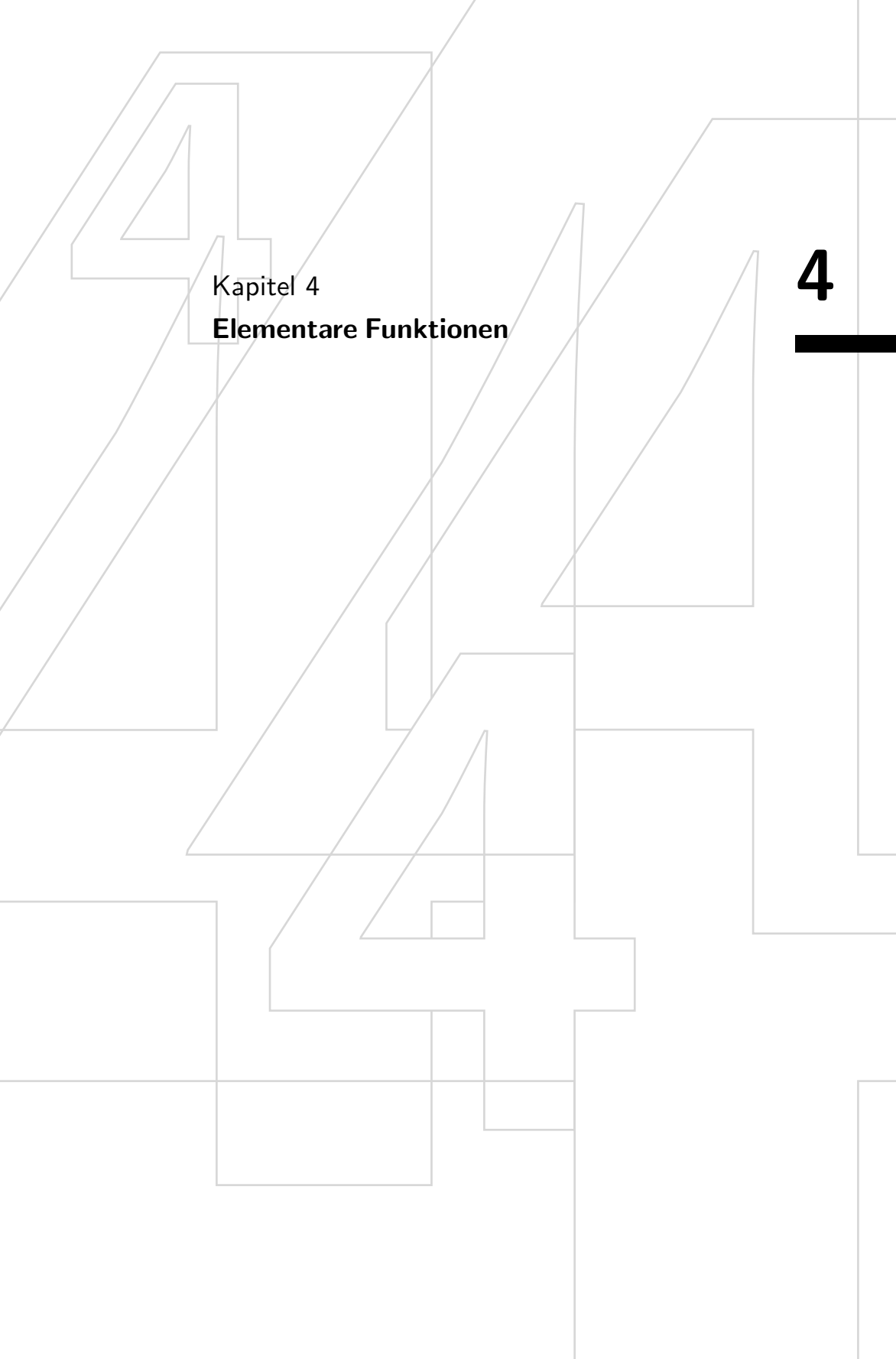
Welches der Gleichungssysteme $A\vec{x} = \vec{b}$, $A^T\vec{x} = \vec{b}$ besitzt eine Lösung? Ist diese eindeutig?

3.16 a) Sind die Vektoren

$$\vec{a} = \begin{pmatrix} 0 \\ 6 \\ -2 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 5 \\ 2 \\ 6 \end{pmatrix}, \quad \vec{c} = \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix} \quad \text{linear abhängig?}$$

$$\text{b) Zeigen Sie, dass die Vektoren } \vec{a} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 9 \\ 1 \\ 4 \end{pmatrix}, \quad \vec{c} = \begin{pmatrix} 3 \\ 1 \\ 2 \end{pmatrix} \quad \text{eine Basis}$$

des $\mathbb{R}^3$ bilden und stellen Sie $\vec{d} = \begin{pmatrix} 12 \\ 0 \\ 5 \end{pmatrix}$ als Linearkombination von $\vec{a}, \vec{b}, \vec{c}$ dar.



Kapitel 4 **Elementare Funktionen**

4

4	Elementare Funktionen	137
4.1	Allgemeine Funktionseigenschaften	139
4.1.1	Grundbegriffe	139
4.1.2	Nullstellen	144
4.1.3	Symmetrieverhalten	145
4.1.4	Monotonie	146
4.1.5	Periodizität	147
4.1.6	Umkehrfunktion (inverse Funktion)	148
4.2	Polynome	152
4.2.1	Festlegung von Polynomen durch gegebene Wertepaare	153
4.2.2	Koeffizientenvergleich	154
4.2.3	Teilbarkeit durch einen Linearfaktor	155
4.2.4	Nullstellenproblem	156
4.2.5	Interpolationspolynome mit dem Newton-Algorithmus	159
4.3	Rationale Funktionen	162
4.3.1	Definitionslücken, Nullstellen, Pole	163
4.3.2	Verhalten im Unendlichen	164
4.3.3	Anwendung: Übertragung bei Wechselstrom-Kreisen	166
4.4	Potenz- und Wurzelfunktionen	167
4.5	Exponential- und Logarithmusfunktion	170
4.5.1	Exponentialfunktion	170
4.5.2	Logarithmusfunktion	172
4.6	Trigonometrische Funktionen	175
4.6.1	Grundbegriffe	175
4.6.2	Sinus- und Kosinusfunktion	175
4.6.3	Tangens- und Kotangensfunktion	180
4.6.4	Zusammenstellung wichtiger trigonometrischer Formeln	181
4.7	Arkusfunktionen	182
4.7.1	Arkussinusfunktion	183
4.7.2	Arkuskosinusfunktion	184
4.7.3	Arkustangensfunktion	184
4.7.4	Arkuskotangensfunktion	185
4.8	Aufgaben zu elementaren Funktionen	188

Zusätzliche Abschnitte auf der CD-Rom

4.9	MAPLE: Elementare Funktionen	cd
-----	------------------------------	----

4 Elementare Funktionen

In diesem Kapitel werden elementare Funktionen, die u.a. zur Beschreibung von physikalischen Vorgängen notwendig sind, angegeben und allgemeine Funktionseigenschaften zur Charakterisierung dieser Funktionen bereitgestellt. Neben den Polynomen werden die gebrochenrationalen Funktionen qualitativ diskutiert. Potenz- und Wurzelfunktionen sowie Exponential- und Logarithmusfunktion werden im Zusammenhang mit wichtigen Anwendungen eingeführt. Zur Beschreibung von Wechselspannungen bzw. harmonischen Schwingungen werden die trigonometrischen Funktionen mit den zugehörigen Umkehrfunktionen benötigt, die im Zusammenhang mit den jeweiligen Anwendungen diskutiert werden.

4.1 Allgemeine Funktionseigenschaften

4.1

In diesem Abschnitt werden grundlegende Begriffe wie z.B. Funktion und Umkehrfunktion, Nullstellen und Symmetrieverhalten als auch Monotonieverhalten wiederholt. Diese Eigenschaften werden in den weiteren Kapiteln zur Charakterisierung der elementaren Funktionen verwendet.

➤ 4.1.1 Grundbegriffe

Für die mathematische Formulierung naturwissenschaftlicher Gesetzmäßigkeiten ist der Begriff der *Funktion* unentbehrlich. Nahezu alle quantitativen Aussagen werden in Form eines funktionalen Zusammenhangs aufgestellt. Bevor der Begriff der Funktion mathematisch präzisiert wird, soll seine zentrale Stellung anhand technischer Beispiele sichtbar gemacht werden.

Beispiel 4.1 (Weg-Zeit-Diagramm): Bewegungsabläufe werden in sog. *Weg-Zeit-Diagrammen* wie in Abb. 4.1 graphisch dargestellt.

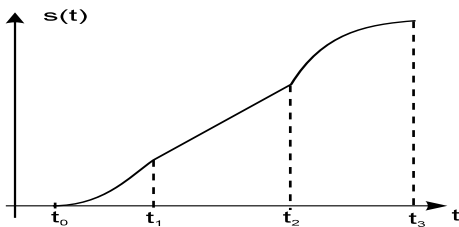


Abb. 4.1. Weg-Zeit-Diagramm

Im Intervall $t_0 \leq t \leq t_1$ findet eine beschleunigte Bewegung, im Intervall $t_1 \leq t \leq t_2$ eine gleichförmige Bewegung und im Intervall $t_2 \leq t \leq t_3$ eine abgebremsete Bewegung statt. Charakteristisch für den dargestellten Vorgang

ist, dass zu jedem Zeitpunkt $t \in [t_0, t_3]$ der Ort $s(t)$ eindeutig bestimmt ist. Durch diese Zuordnung wird eine *Funktion* definiert. Die allgemeine Formulierung lautet daher:

Definition: (Funktion). Sei $ID \subset \mathbb{R}$ eine Teilmenge von $\mathbb{R}$. Unter einer **reellen Funktion** f auf ID versteht man eine Abbildung

$$\begin{aligned} f : ID &\rightarrow \mathbb{R} \\ x &\mapsto f(x). \end{aligned}$$

Die Abbildung f bildet jedes Element aus ID **eindeutig** auf ein Element in $\mathbb{R}$ ab.

Die Menge ID heißt **Definitionsbereich** von f und die Menge $\mathbb{R}$ der **Zielbereich** von f . Die Menge

$$\mathbb{W} := f(ID) := \{f(x) : x \in ID\}$$

heißt der **Wertebereich** von f . Man nennt $f(x)$ den Funktionswert oder Funktionsausdruck an der Stelle x .

Zur graphischen Darstellung einer Funktion verwenden wir das kartesische Koordinatensystem $\mathbb{R} \times \mathbb{R} = \{(x, y) : x \in \mathbb{R}, y \in \mathbb{R}\}$, in dem jeder Punkt P der Ebene durch ein Zahlenpaar (x, y) eindeutig beschrieben wird. Zur geometrischen Veranschaulichung der Funktion zeichnet man die Menge der Punkte $(x, f(x))$ in ein Koordinatensystem und erhält so die zugehörige Kurve.

Definition: Unter dem **Graphen** G_f einer Funktion $f : ID \rightarrow \mathbb{R}$ mit $x \mapsto f(x)$ verstehen wir die Menge aller Paare (x, y) :

$$G_f := \{(x, y) : x \in ID \text{ und } y = f(x)\}.$$

Beispiele für Funktionen aus der Technik:

Beispiel 4.2 (Freier Fall): Das Weg-Zeit-Gesetz beim freien Fall aus der Höhe s_0 mit Anfangsgeschwindigkeit v_0 lautet

$$s(t) = \frac{1}{2} g t^2 + v_0 t + s_0, \quad t \geq 0,$$

wenn $g = 9,81 \frac{m}{s^2}$ die Erdbeschleunigung ist.

Beispiel 4.3 (Plattenkondensator): Wird an einem Plattenkondensator eine Ladung Q angelegt, so ist die induzierte Spannung

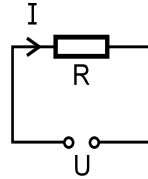
$$U(Q) = \frac{1}{C}Q,$$

wenn C die Kapazität des Kondensators ist.

Beispiel 4.4 (Gleichstromkreis): In einem elektrischen Gleichstromkreis ist die Stromstärke I abhängig von der angelegten Gleichspannung U . Es gilt

$$I(U) = \frac{U}{R},$$

wenn R der Ohmsche Widerstand der Schaltung ist.



Beispiel 4.5 (Wärmeausdehnung von Gasen (Gay-Lussac)): Ist V_0 das Volumen eines idealen Gases bei Temperatur T_0 , so gilt bei konstantem Druck p für das Volumen bei der Temperatur T

$$V(T) = V_0 (1 + \gamma (T - T_0)), \quad T \geq T_0,$$

wenn γ die Wärmeausdehnungskonstante ist.

Beispiel 4.6 (Barometrische Höhenformel): Für den Luftdruck p in Höhe h über dem Erdboden gilt näherungsweise

$$p(h) = p_0 e^{-\frac{h-H}{H}}, \quad h \geq H,$$

wenn $H = \frac{p_0}{\rho_0 \cdot g}$ und p_0, ρ_0 Luftdruck und -dichte am Erdboden ist.

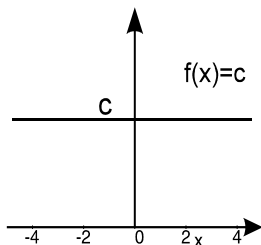
All diesen physikalischen Gesetzmäßigkeiten ist gemeinsam, dass zu einer unabhängigen Variablen t, Q, U, T, h die physikalische Größe $s(t), U(Q), I(U), V(T)$ bzw. $p(h)$ **eindeutig** berechenbar ist. Man spricht daher auch oftmals von der *unabhängigen Variablen* x und der *abhängigen Variablen* (= Funktion) $f(x)$.

Mathematische Beispiele:

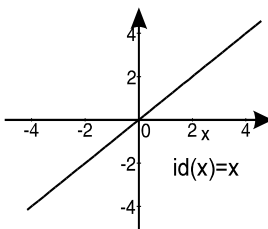
Beispiel 4.7. $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto f(x) = c$ heißt *konstante Funktion*.

Beispiel 4.8. $id_{\mathbb{R}} : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto id_{\mathbb{R}}(x) = x$ heißt *identische Funktion*.

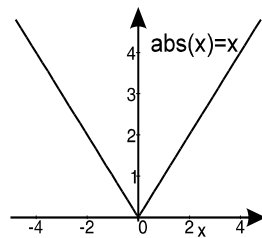
Beispiel 4.9. $abs : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto abs(x) = |x|$ heißt *Betragsfunktion*.



7. Konstante Funktion



8. Identische Funktion

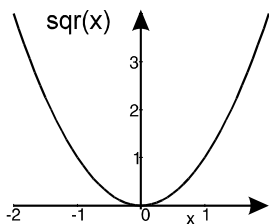


9. Betragsfunktion

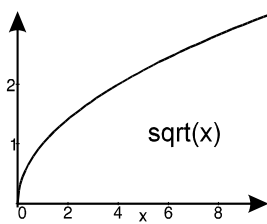
Beispiel 4.10. $\text{sqr} : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto \text{sqr}(x) = x^2$ heißt *Quadratfunktion*.

Beispiel 4.11. $\text{sqr}t : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$ mit $x \mapsto \text{sqr}t(x) = \sqrt{x}$ heißt *Wurzelfunktion*.

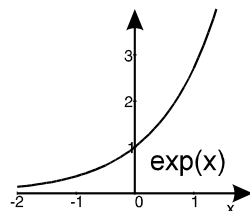
Beispiel 4.12. $\exp : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto \exp(x) = e^x$ heißt *Exponentialfunktion*.



10. Quadratfunktion

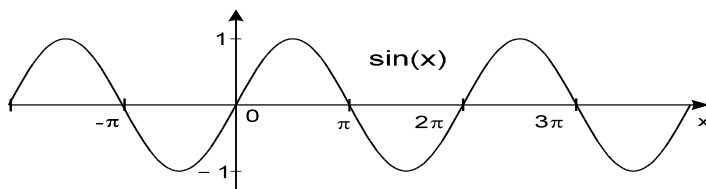


11. Wurzelfunktion



12. Exponentialfunktion

Beispiel 4.13. $\sin : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto \sin(x)$ heißt *Sinusfunktion*.



13. Sinusfunktion

Achtung: Es gibt auch Funktionen, die man nicht zeichnen kann:

Beispiel 4.14. $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto f(x) = \begin{cases} 0 & \text{falls } x \text{ rational} \\ 1 & \text{falls } x \text{ irrational} \end{cases}$.

Bemerkung: Im Hinblick auf die Programmierung bzw. der Anwendung von MAPLE ist zu unterscheiden zwischen einer *Funktion* und einem *Ausdruck* (bzw. Funktionsausdruck): Z.B. $\text{sqr}(x)$ ist ein Ausdruck für $\sqrt{x}$, welcher zu gegebenem x die Wurzel aus x darstellt. Hingegen ist sqr eine Funktion, die an einer Stelle x ausgewertet werden kann. Eine Funktion ist mehr als nur der Funktionsausdruck: Jede Funktion besteht aus Definitionsbereich, Zielbereich und einer Funktionszuweisung.

Anwendung: Logarithmische Darstellung von Funktionen.

Häufig werden in Physik und Technik logarithmische oder doppel-logarithmische Darstellungen der Messergebnisse verwendet, um exponentielle, potenzielle oder logarithmische Gesetzmäßigkeiten aufzuzeigen.

- (1) Liegt ein exponentielles Gesetz vor

$$y = c \cdot a^x,$$

so liefert eine logarithmische Skalierung der y -Achse

$$\log(y) = \log(c) + x \cdot \log(a)$$

als Graphen eine Gerade mit Steigung $\log(a)$ und dem Achsenabschnitt $\log(c)$. Dabei dürfen die Funktionswerte nur positive Werte enthalten, da der Logarithmus nur für positive reelle Zahlen definiert ist. Beispielsweise wird durch die logarithmische Skalierung der y -Achse eine allgemeine Exponentialfunktion $y = 10 \cdot 4^{-x}$ als Gerade gezeichnet.

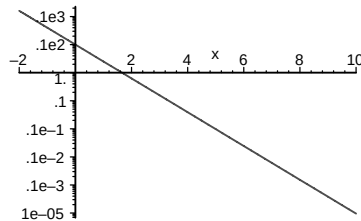


Abb. 4.2. Logarithmische Darstellung der Funktion $y = 10 \cdot 4^{-x}$.

- (2) Liegt eine allgemeine Potenzfunktion vor

$$y = a \cdot x^b,$$

dann liefert eine doppel-logarithmische Auftragung

$$\log(y) = \log(a) + b \cdot \log(x)$$

als Graphen eine Gerade mit Steigung b und dem Achsenabschnitt $\log(a)$. Sowohl der x -Bereich als auch die zugehörigen Funktionswerte dürfen nur positive Werte enthalten, da der Logarithmus nur für positive reelle Zahlen definiert ist. Eine doppel-logarithmische Auftragung der Potenzfunktion $y = 10 x^{\frac{1}{3}}$ liefert die folgende Gerade:

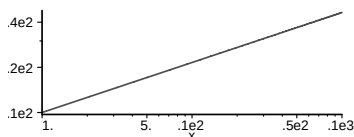


Abb. 4.3. Doppel-logarithmische Darstellung der Funktion $y = 10x^{\frac{1}{3}}$.

- (3) Liegt ein logarithmisches Gesetz der Form $y = a \log(x) + b$ vor, dann liefert die logarithmische Auftragung der x -Achse eine Gerade mit der Steigung a und dem Achsenabschnitt b .



Hinweis: Auf der CD-Rom befindet sich ein [MAPLE-Worksheet](#), mit dem man z.B. Messdaten aus einer Datei einlesen und diese dann logarithmisch auftragen kann, um die physikalische Gesetzmäßigkeit zu erschließen.

Allgemeine Funktionseigenschaften

4.1.2 Nullstellen

Definition: Eine Funktion f besitzt in x_0 eine **Nullstelle**, wenn

$$f(x_0) = 0.$$

Beispiele 4.15:

- ① Die Funktion $f(x) = x - 2$ schneidet die x -Achse an der Stelle $x_0 = 2$.
- ② Die Funktion $f(x) = \sin(x)$ hat unendlich viele Nullstellen: $0, \pm\pi, \pm2\pi, \dots$

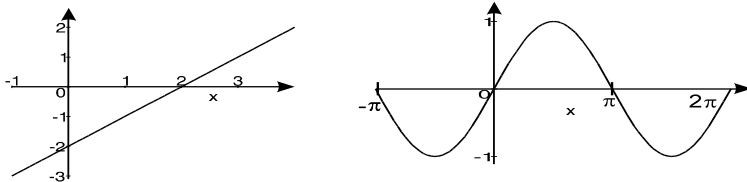


Abb. 4.4. Nullstellen von Funktionen

4.1.3 Symmetrieverhalten

Definition: Eine Funktion f mit symmetrischem Definitionsbereich heißt **gerade**, wenn für alle $x \in \mathbb{D}$ gilt

$$f(-x) = f(x).$$

Bemerkung: Die Funktionskurve einer geraden Funktion ist **spiegelsymmetrisch** zur y -Achse.

Beispiele 4.16:

- ① Die Funktion $f(x) = x^2$ ist gerade.
- ② Die Funktion $f(x) = \cos x$ ist gerade.
- ③ Jede Potenzfunktion $f(x) = x^n$ mit geradem n ist gerade.

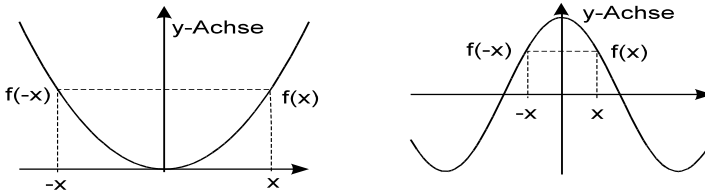


Abb. 4.5. Gerade Funktionen x^2 und $\cos(x)$: Spiegelsymmetrisch zur y -Achse

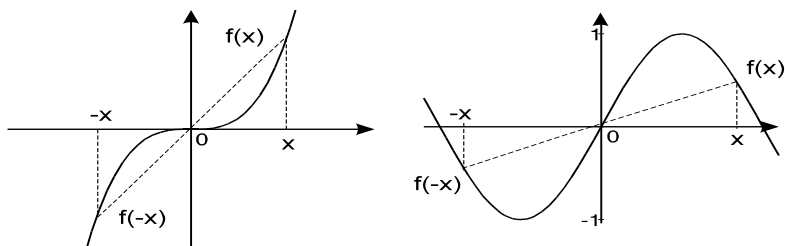
Definition: Eine Funktion f mit symmetrischem Definitionsbereich heißt **ungerade**, wenn für alle $x \in \mathbb{D}$ gilt

$$f(-x) = -f(x).$$

Bemerkung: Die Funktionskurve einer ungeraden Funktion ist **punktsymmetrisch** zum Ursprung.

Beispiele 4.17:

- ① Die Funktion $f(x) = x^3$ ist ungerade.
- ② Die Funktion $f(x) = \sin x$ ist ungerade.
- ③ Jede Potenzfunktion $f(x) = x^n$ mit ungeradem n ist ungerade.

Abb. 4.6. Ungerade Funktionen x^3 und $\sin(x)$: Spiegelsymmetrisch zum Ursprung

4.1.4 Monotonie

Definition: Eine Funktion $f : \mathbb{D} \rightarrow \mathbb{R}$ heißt

monoton wachsend, falls	$f(x_1) \leq f(x_2).$
streng monoton wachsend, falls	$f(x_1) < f(x_2).$
monoton fallend, falls	$f(x_1) \geq f(x_2).$
streng monoton fallend, falls	$f(x_1) > f(x_2).$

für alle $x_1, x_2 \in \mathbb{D}$ mit $x_1 < x_2$ gilt.

Bemerkung: Eine streng monoton wachsende Funktion f (Abb. 4.7 (a)) hat also die Eigenschaft, dass zum kleineren x -Wert der kleinere Funktionswert gehört. Bei einer streng monoton fallenden Funktion (Abb. 4.7 (b)) ist es genau umgekehrt: Zum kleineren x -Wert gehört der größere Funktionswert.

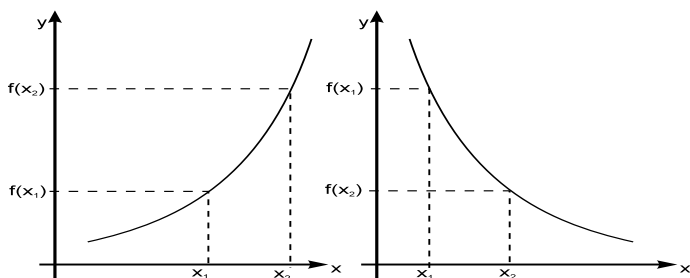


Abb. 4.7. Streng monoton wachsende (a) und streng monoton fallende (b) Funktionen

Beispiele 4.18:

① Streng monoton wachsende Funktionen:

- (i) Jede Gerade mit positiver Steigung.
- (ii) Die Funktion $f(x) = x^3$.
- (iii) Die Funktion $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$ mit $f(t) = y_0(1 - e^{-t/t_0})$.

② Der Ladevorgang eines Kondensators entspricht einer streng monoton wachsenden Funktion: Beim Laden eines Kondensators C über eine Batterie mit Spannung U_0 ist der Zeitverlauf der Spannung am Kondensator gegeben durch

$$U(t) = U_0 \left(1 - e^{-t/RC}\right).$$

③ Streng monoton fallende Funktionen:

- (i) Jede Gerade mit negativer Steigung.
- (ii) Die Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f(t) = y_0 e^{-t/t_0}$.

④ Der radioaktive Zerfall entspricht einer streng monoton fallenden Funktion: Beim radioaktiven Zerfall nimmt die Anzahl der Atomkerne $n(t)$ nach dem Exponentialgesetz $n(t) = n_0 e^{-t/t_0}$ mit der Zeit ab.

⑤ Der Entladevorgang eines Kondensators entspricht einer streng monoton fallenden Funktion: Entladet man einen Kondensator C über einen Ohmschen Widerstand R , so klingt die Spannung am Kondensator exponentiell ab:

$$U(t) = U_0 e^{-t/RC}.$$

⑥ Die quadratische Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto f(x) = x^2$ ist weder monoton fallend noch monoton wachsend. Schränkt man sie jedoch auf ein Teilintervall $\mathbb{R}_{\geq 0}$ oder $\mathbb{R}_{\leq 0}$ ein, so ist sie dort monoton:

- $f_+: \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$ mit $x \mapsto f_+(x) = x^2$ ist streng monoton wachsend.
- $f_-: \mathbb{R}_{\leq 0} \rightarrow \mathbb{R}$ mit $x \mapsto f_-(x) = x^2$ ist streng monoton fallend.

► 4.1.5 Periodizität

Definition: Eine Funktion $f: \mathbb{D} \rightarrow \mathbb{R}$ heißt **periodisch** mit Periode $p > 0$, wenn $f(x+p) = f(x)$ für alle $x \in \mathbb{D}$.

Beispiele 4.19:

- ① Die Sinusfunktion $\sin: \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto \sin(x)$ ist periodisch mit der Periode 2π . Die Kosinusfunktion $\cos: \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto \cos(x)$ ist periodisch mit Periode 2π (siehe §4.6).
- ② Die Tangens- und Kotangensfunktion, die wir in §4.6 noch näher besprechen werden, sind periodisch mit Periode π .

➤ 4.1.6 Umkehrfunktion (inverse Funktion)

Eine Funktion f ordnet jeder Zahl x_0 ihres Definitionsbereiches $\mathbb{D}$ eindeutig einen Wert $y_0 = f(x_0)$ zu. Diese Zuordnung ist in Abb. 4.8 durch den Pfeil von x_0 nach y_0 gekennzeichnet.

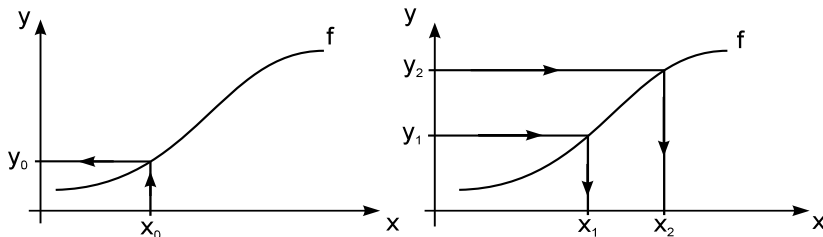


Abb. 4.8. Zuordnung und Umkehrzuordnung

Häufig stellt sich das umgekehrte Problem. Zu gegebenem Funktionswert (y -Wert) ist der zugehörige x -Wert zu bestimmen. In Abb. 4.8 ist diese Umkehrung der Problematik durch die Umkehrung der Pfeile gekennzeichnet. Gegeben sind y_1 und y_2 , gesucht sind die zugehörigen x_1 und x_2 . Zu gegebenem y -Wert aus dem Wertebereich von f kann genau ein x -Wert angegeben werden, wenn die Funktion die Eigenschaft hat, dass aus $x_1 \neq x_2$ stets folgt $f(x_1) \neq f(x_2)$.

Definition:

- (1) Eine Funktion $f : \mathbb{D}_f \rightarrow \mathbb{W}_f$ heißt **umkehrbar**, wenn aus $x_1 \neq x_2$ stets folgt $f(x_1) \neq f(x_2)$.
- (2) Ist die Funktion f umkehrbar, so gibt es zu jedem $y \in \mathbb{W}_f$ genau ein $x \in \mathbb{D}_f$. Diese eindeutige Zuordnung $y \mapsto x$ liefert eine Funktion, die **Umkehrfunktion (inverse Funktion)**

$$f^{-1} : \mathbb{W}_f \rightarrow \mathbb{D}_f.$$

⚠ **Achtung:** Nicht jede Funktion ist umkehrbar, wie Beispiel 4.20 zeigt.

Beispiele 4.20:

- ① Die Funktion $f_1 : \mathbb{R} \rightarrow \mathbb{R}$ mit $f_1(x) = x^2$ ist **nicht** umkehrbar, da für $x_0 \neq 0$ $f(x_0) = x_0^2 = f(-x_0)$ aber $x_0 \neq -x_0$ (Abb. 4.9 (a)).
- ② Die Funktion $f_2 : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$ mit $f_2(x) = x^2$ ist **nicht** umkehrbar. Obwohl man den Definitionsbereich so eingeschränkt hat, dass f_2 eine streng monoton wachsende Funktion ist, so gibt es doch für die negativen y -Werte keine zugehörigen x -Werte.

- ③ Die Funktion $f_3 : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ mit $f_3(x) = x^2$ ist **umkehrbar**. f_3 ist auf dem Definitionsbereich eine streng monoton wachsende Funktion und der Zielbereich fällt mit dem Wertebereich zusammen. Jedem Wert y aus dem Zielbereich kann nun eindeutig ein x zugeordnet werden (Abb. 4.9 (b)).

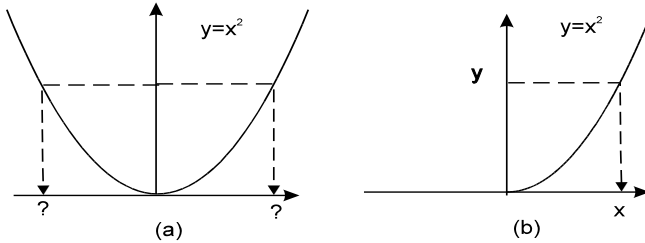


Abb. 4.9. Umkehrung der Funktion $y = x^2$

Bemerkung: Anhand des Beispiels erkennt man, dass eine Funktion nicht nur durch eine Funktionsvorschrift charakterisiert werden kann, sondern Definitionsbereich **und** Zielbereich dazugehören.

In vielen (aber nicht allen) Fällen kann man für umkehrbare Funktionen eine explizite Zuordnungsvorschrift für die Umkehrfunktion wie folgt angeben:

Bestimmung der Funktionsgleichung der Umkehrfunktion

- (1) Man schränkt den Definitionsbereich der Funktion $\mathbb{D}_f$ so ein, dass f in dem eingeschränkten Definitionsbereich streng monoton verläuft.
- (2) Man schränkt den Zielbereich auf den Wertebereich der Funktion $\mathbb{W}_f$ ein. Damit erhält man eine umkehrbare Funktion!
- (3) Man löst die Funktionsgleichung $y = f(x)$ nach der Variablen x auf und erhält die aufgelöste Form $x = g(y)$.
- (4) Durch formales Vertauschen der beiden Variablen x und y gewinnt man die Zuordnungsvorschrift der Umkehrfunktion

$$f^{-1}(x) = y = g(x).$$

Beispiele 4.21:

- ① $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = 2x + 1$.

Wie man aus dem Funktionsgraphen entnimmt, stellt f eine Gerade dar. f ist auf ganz $\mathbb{R}$ streng monoton steigend mit dem Wertebereich $\mathbb{R}$; daher ist f umkehrbar!

Durch Auflösen von

$$y = 2x + 1$$

nach x folgt

$$x = \frac{1}{2}y - \frac{1}{2}$$

und durch formales Vertauschen von x und y erhält man die Umkehrfunktion

$$f^{-1} : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f^{-1}(x) = \frac{1}{2}x - \frac{1}{2}.$$

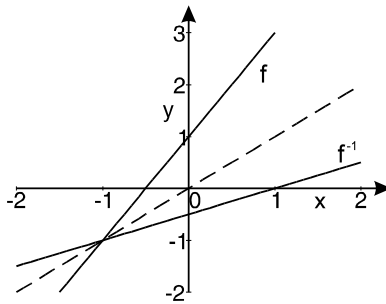


Abb. 4.10. Funktion $f(x) = \frac{1}{2}x + 1$ und Umkehrfunktion $f^{-1}(x) = \frac{1}{2}x - \frac{1}{2}$

② $f : [0, \infty) \rightarrow [1, \infty)$ mit $f(x) = 1 + x^2$.

f ist auf dem Definitionsbereich $\mathbb{D} = [0, \infty)$ streng monoton wachsend und der Zielbereich fällt mit dem Wertebereich $\mathbb{W} = [1, \infty)$ zusammen. Auflösen der Funktionsgleichung

$$y = 1 + x^2$$

nach x liefert

$$x = \pm \sqrt{y - 1}.$$

Da nur positive x -Werte durch den Definitionsbereich zugelassen sind, entfällt das Minuszeichen! Vertauschen von x und y liefert die Umkehrfunktion:

$$f^{-1} : [1, \infty) \rightarrow [0, \infty) \text{ mit } f^{-1}(x) = \sqrt{x - 1}.$$

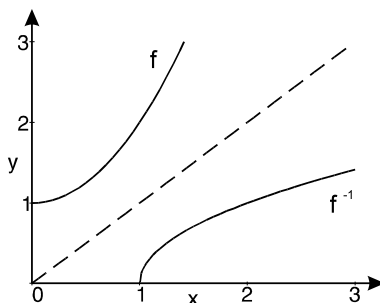


Abb. 4.11. Funktion $f(x) = 1 + x^2$ und Umkehrfunktion $f^{-1}(x) = \sqrt{x - 1}$

Zusammenfassung über die Umkehrung einer Funktion

- (1) Jede streng monoton wachsende oder fallende Funktion $f : \mathbb{D}_f \rightarrow \mathbb{W}_f$ ist umkehrbar.
- (2) Bei der Umkehrung einer Funktion werden Definitions- und Wertebereich miteinander vertauscht.
- (3) Zeichnerisch erhält man das Schaubild der Umkehrfunktion durch Spiegelung von f an der Winkelhalbierenden $y = x$.

Zur Charakterisierung der umkehrbaren Funktionen führen wir noch die folgenden Begriffe ein:

Definition: Sei $f : \mathbb{D}_f \rightarrow \mathbb{W}_f$ eine Funktion. f heißt

- (1) **injektiv**, falls gilt: $x_1 \neq x_2 \Rightarrow f(x_1) \neq f(x_2)$.
- (2) **surjektiv**, falls es zu jedem $y \in \mathbb{W}_f$ ein $x \in \mathbb{D}_f$ gibt mit $y = f(x)$.
- (3) **bijektiv**, falls gilt: f ist injektiv und surjektiv.

Man nennt die Injektivität einer Funktion auch die *Eineindeutigkeit* einer Funktion. Sie bedeutet, dass eine Parallele zur x -Achse den Graphen der Funktion höchstens einmal schneidet. Eine Funktion f ist surjektiv, wenn zu jedem Wert y aus dem Zielbereich mindestens ein x aus dem Definitionsbereich gefunden wird, welches auf y abgebildet wird.

Beispiele 4.22:

- ① Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist weder injektiv, noch surjektiv. Denn sowohl x als auch $-x$ führen zu dem selben Funktionswert. Die Funktion ist auch nicht surjektiv, da zu den negativen y -Werten keine zugehörigen x -Werte gibt.
- ② Die Funktion $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist injektiv, aber nicht surjektiv, da z.B. zu $y = (-1)$ kein x aus dem Definitionsbereich gefunden wird mit $x^2 = -1$.
- ③ Die Funktion $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ ist bijektiv. □

Die bijektiven Funktionen sind genau die umkehrbaren Funktionen:

Satz: Zu einer Funktion f existiert die Umkehrfunktion genau dann, wenn f bijektiv ist.

4.2 Polynome

Im Abschnitt über die Polynome werden grundlegende Techniken, wie z.B. Nullstellenbestimmung, Linearfaktorzerlegung, Horner-Schema und Interpolationspolynom, eingeführt, die wir dann z.B. bei der Diskussion gebrochenrationaler Funktionen anwenden.

Polynome (oder *Polynomfunktionen*) spielen in der Angewandten Mathematik eine wichtige Rolle. Sie sind nicht nur besonders einfach in ihrer Darstellung, sondern sie lassen sich auch auf einfache Weise auswerten, da nur Additionen und Multiplikationen ausgeführt werden müssen. Gerade deshalb sind sie für die Anwendung in der Numerik ideale Funktionen. Im Kapitel über Taylor-Reihen ($\rightarrow$ §9.3) werden wir zeigen, dass man komplizierte Funktionen auf eine Darstellung über Polynome zurückspielt. In diesem Kapitel lernen wir für die Anwendung wichtige Eigenschaft kennen, wie z.B. die Festlegung von Polynomen durch vorgegebene Punkte.

Definition: Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ der Form

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 \quad \text{mit } a_n \neq 0$$

heißt **Polynom** (**Polynomfunktion**, **ganzzrationale Funktion**) vom Grad n . Die reellen Zahlen $a_0, a_1, \dots, a_n$ heißen **Koeffizienten** des Polynoms.

Beispiele 4.23:

- | | | | |
|---|----------------------------|---------------------|-------------------------|
| ① | $p_1(x) = 4$ | Polynom vom Grad 0. | (konstante Funktion) |
| | $p_2(x) = 2x - 3$ | Polynom vom Grad 1. | (lineare Funktion) |
| | $p_3(x) = 2x^2 - 3x + 5$ | Polynom vom Grad 2. | (quadratische Funktion) |
| | $p_4(x) = x^3 - x$ | Polynom vom Grad 3. | (kubische Funktion) |
| | $p_5(x) = 4x^8 - x^5 - 10$ | Polynom vom Grad 8. | |

- ② Die Höhe H einer **Quecksilbersäule** hängt von der Temperatur T ab:

$$H(T) = H_0 (1 + \alpha (T - T_0)).$$

Dabei ist H_0 die Höhe bei der Temperatur T_0 und α der Wärmeausdehnungskoeffizient.

- ③ Ein **strömendes Medium** (Luft, Wasser), das mit einer mittleren Geschwindigkeit v auf einen Körper trifft, übt die Kraft

$$F_w(v) = c_w A \frac{1}{2} \rho v^2$$

auf ihn aus. Dabei ist c_w der Widerstandsbeiwert, A die Querschnittsfläche des Körpers und ρ die Dichte des strömenden Mediums.

- ④ Die **Biegelinie** eines einseitig eingespannten Trägers wird beschrieben durch

$$y(x) = \frac{F}{6 E \cdot I} (3 l x^2 - x^3)$$

$E \cdot I$ ist die Biegefestigkeit, F die Kraft am Balkenende und l die Balkenlänge. $\square$

► 4.2.1 Festlegung von Polynomen durch gegebene Wertepaare

Eine der wichtigsten Eigenschaften von Polynomen ist die Festlegung durch gegebene Wertepaare. Im Falle des Thermometers braucht man 2 Werte, um eine lineare Skala festzulegen. Schwach durchhängende Seile haben in guter Näherung eine Parabelform. Zur Festlegung dieses funktionalen Zusammenhangs müssen 3 Wertepaare angegeben werden. Denn durch 3 Punkte wird genau ein Polynom 2-ten Grades festgelegt. Allgemein gilt der folgende Satz:

Satz: Zu $n + 1$ verschiedenen Wertepaaren $(x_1, y_1), (x_2, y_2), \dots, (x_{n+1}, y_{n+1})$ gibt es genau eine Polynomfunktion f mit $f(x_i) = y_i$, deren Grad nicht größer als n ist.

Eine Möglichkeit das gesuchte Polynom anzugeben, ist gegeben durch die **Lagrange Interpolationsformel**:

$$\begin{aligned} f(x) = & y_1 \cdot \frac{(x-x_2)(x-x_3) \cdots (x-x_{n+1})}{(x_1-x_2)(x_1-x_3) \cdots (x_1-x_{n+1})} + \\ & y_2 \cdot \frac{(x-x_1)(x-x_3) \cdots (x-x_{n+1})}{(x_2-x_1)(x_2-x_3) \cdots (x_2-x_{n+1})} + \\ & \vdots \\ & y_{n+1} \cdot \frac{(x-x_1)(x-x_2) \cdots (x-x_n)}{(x_{n+1}-x_1)(x_{n+1}-x_2) \cdots (x_{n+1}-x_n)}. \end{aligned}$$

Beim Auswerten der Funktion an der Stelle x_i verschwinden alle Terme $k \neq i$, da $(x-x_i)$ als Faktor enthalten ist; nur der Term $k = i$ bleibt erhalten, da bei diesem Summand der Faktor $(x-x_i)$ nicht auftritt:

$$\begin{aligned} f(x_i) &= y_i \frac{(x_i-x_1)(x_i-x_2) \cdots (x_i-x_{i-1})(x_i-x_{i+1}) \cdots (x_i-x_{n+1})}{(x_i-x_1)(x_i-x_2) \cdots (x_i-x_{i-1})(x_i-x_{i+1}) \cdots (x_i-x_{n+1})} \\ &= y_i. \end{aligned}$$

Außerdem erhält man höchstens ein Polynom vom Grade n , da alle Terme nur n Faktoren $(x-x_i)$ enthalten. $\square$

➤ 4.2.2 Koeffizientenvergleich

Oftmals führt man bei Polynomen einen *Koeffizientenvergleich* durch. Diese Methode beruht auf der Tatsache, dass

Satz: (Koeffizientenvergleich). Zwei Polynome f und g mit

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$$

und

$$g(x) = b_m x^m + a_{m-1} x^{m-1} + \dots + b_0$$

sind genau dann gleich, wenn $n = m$ und $a_i = b_i$ für alle i .

Dieser Satz besagt, dass Polynome nur dann gleich sind, wenn sie denselben Grad besitzen und die Koeffizienten identisch sind.

Begründung: Wir nehmen an, dass $m \geq n$. Da $f(x) = g(x)$ für alle $x \in \mathbb{R}$, gilt für $x = 0$

$$\begin{aligned} f(0) = g(0) &\Rightarrow a_0 = b_0 \\ &\Rightarrow x(a_n x^{n-1} + \dots + a_1) = x(b_m x^{m-1} + \dots + b_1) \end{aligned}$$

Damit ist

$$a_n x^{n-1} + \dots + a_2 x + a_1 = b_m x^{m-1} + \dots + b_2 x + b_1$$

für alle $x \in \mathbb{R}$, also insbesondere für

$$x = 0 \Rightarrow a_1 = b_1.$$

Wieder kann man x ausklammern und die Vorgehensweise wiederholen:

$$\Rightarrow a_2 = b_2, a_3 = b_3, \dots, a_{n-1} = b_{n-1},$$

bis man schließlich

$$a_n = b_n + b_{n+1} x + \dots + b_m x^{m-n+1}$$

erhält. Durch Einsetzen von $x = 0$ folgt $a_n = b_n$, so dass für alle $x \in \mathbb{R}$:

$$\Rightarrow b_m x^{m-n} + \dots + b_{n+1} = 0.$$

Ein Polynom kann aber nur dann das Nullpolynom sein, wenn alle Koeffizienten verschwinden, d.h. $b_m = b_{m-1} = \dots = b_{n+1} = 0$. Folglich ist der Grad von f gleich dem Grad von g und die Koeffizienten a_i von f identisch mit den Koeffizienten b_i von g . □

4.2.3 Teilbarkeit durch einen Linearfaktor

Eine der elementaren Aufgaben von Polynomen besteht in der **Auswertung** eines Polynoms an einer Stelle x_0 . Ein sehr einfaches Schema erhält man durch den folgenden Ansatz:

Satz: Für jedes Polynom f und jeden Wert x_0 ist die folgende Umformung möglich:

$$\begin{aligned} f(x) &= a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 \\ &= (x - x_0) (b_n x^{n-1} + \dots + b_2 x + b_1) + r. \end{aligned}$$

Begründung: Um diese Gleichheit zu überprüfen, multiplizieren wir das Produkt aus und führen einen Koeffizientenvergleich durch.

$$\begin{aligned} (x - x_0) (b_n x^{n-1} + \dots + b_2 x + b_1) + r &= \\ &= b_n x^n + b_{n-1} x^{n-1} + b_{n-2} x^{n-2} + \dots + b_1 x \\ &\quad - x_0 b_n x^{n-1} - x_0 b_{n-1} x^{n-2} - \dots - x_0 b_2 x - x_0 b_1 + r \\ &\stackrel{!}{=} a_n x^n + a_{n-1} x^{n-1} + a_{n-2} x^{n-2} + \dots + a_1 x + a_0. \end{aligned}$$

Der Koeffizientenvergleich nach absteigenden Potenzen von x liefert

$$\begin{array}{llll} n & : & & b_n = a_n \\ n-1 & : & b_{n-1} - x_0 b_n = a_{n-1} & \Rightarrow b_{n-1} = a_{n-1} + x_0 b_n \\ n-2 & : & b_{n-2} - x_0 b_{n-1} = a_{n-2} & \Rightarrow b_{n-2} = a_{n-2} + x_0 b_{n-1} \\ \vdots & & & \\ 1 & : & b_1 - x_0 b_2 = a_1 & \Rightarrow b_1 = a_1 + x_0 b_2 \\ 0 & : & -x_0 b_1 + r = a_0 & \Rightarrow r = a_0 + x_0 b_1 \end{array}$$

Mit diesem Vorgehen ist man in der Lage, systematisch $b_n, b_{n-1}, \dots, b_1$ und r zu bestimmen. Außerdem ist

$f(x_0) = r.$

□

Dieses Verfahren ist die Grundlage für die effektive Auswertung von Polynomen. Schon seit Anfang des vorvorigen Jahrhunderts (1819) ist bekannt, dass die Funktionswerte eines Polynoms an einer beliebigen Zwischenstelle x_0 so durch das **Horner-Schema** berechnet werden können:

a_n	a_{n-1}	a_2	a_1	a_0
+	$x_0 \cdot b_n$	$x_0 \cdot b_3$	$x_0 \cdot b_2$	$x_0 \cdot b_1$
b_n	$\nearrow$ b_{n-1}	$\nearrow$ $\dots$	$\nearrow$ b_2	$\nearrow$ b_1
				$r = f(x_0)$

Der Vorteil des *Horner-Schemas* besteht darin, dass für Polynome nicht jeweils die Potenzen $x_0^n, x_0^{n-1}, \dots$ berechnet, sondern Zwischenwerte mit x_0 multipliziert und geeignet aufsummiert werden.

⚠ Achtung: Ist die Potenz x^k ($k < n$) nicht im Polynom vorhanden, muss an der entsprechenden Stelle eine 0 im Horner-Schema eingetragen werden.

Beispiel 4.24. Man berechne den Funktionswert von $f(x) = 2x^4 - 6x^3 - 35x + 10$ an der Stelle $x_0 = 4$:

	2	-6	0	-35	10
+		2·4	2·4	8·4	-3·4
	2	2	8	-3	-2=f(4)

Dieses Schema der Auswertung ist besonders effektiv, da es das Potenzieren vermeidet. Es zeigt sich auch, dass das Horner-Schema für Rundungsfehler unanfällig ist und sich dadurch für den numerischen Einsatz eignet. $\square$

➤ 4.2.4 Nullstellenproblem

Für den Spezialfall, dass x_1 eine Nullstelle des Polynoms $f(x)$ ist, liefert das Horner-Schema sofort die sog. *Produktdarstellung* des Polynoms in der Form $(x - x_1) \cdot \text{Restpolynom}$:

Satz: Wenn x_1 eine Nullstelle des Polynoms n -ten Grades f , dann gilt

$$f(x) = (x - x_1) (b_n x^{n-1} + b_{n-1} x^{n-2} + \dots + b_1).$$

Begründung: Denn wenn x_1 eine Nullstelle, ist das Restglied $r = f(x_1) = 0$ und wir erhalten die obige Behauptung. $\square$

Beispiel 4.25. $x_1 = 1$ ist eine Nullstelle des Polynoms $f(x) = x^3 + 2x^2 - 13x + 10$:

	1	2	-13	10
+		1·1	3·1	-10·1
	1	3	-10	0=f(1)

Das Horner-Schema liefert dann nicht nur den Funktionswert $f(1) = 0$, sondern auch gleich die Koeffizienten b_i des Restpolynoms und somit die Zerlegung des Polynoms f in

$$x^3 + 2x^2 - 13x + 10 = (x - 1) (1x^2 + 3x + (-10)).$$

Das gleiche Ergebnis erhält man, wenn man eine **Polynomdivision** durchführt:

$$\begin{array}{r}
 (x^3 + 2x^2 - 13x + 10) : (x - 1) = 1x^2 + 3x - 10 \\
 \underline{-(x^3 - x^2)} \\
 3x^2 - 13x \\
 \underline{-(3x^2 - 3x)} \\
 -10x + 10 \\
 \underline{-(-10x + 10)} \\
 0
 \end{array}$$

□

Ist also f ein Polynom n -ten Grades und x_1 eine Nullstelle von f , so lässt sich $f(x)$ durch $(x - x_1)$ ohne Rest dividieren und das Resultat ist ein Polynom vom Grade $n - 1$. Das Abspalten eines Linearfaktors von einem Polynom vom Grade n kann man jedoch höchstens n -mal durchführen, solange bis das Restpolynom nur noch den Grad Null hat:

Satz: Jedes Polynom n -ten Grades hat **höchstens** n verschiedene Nullstellen.

Eine schwierige Aufgabe ist die konkrete Bestimmung der Nullstellen von Polynomen. Für $n = 2$ hat man noch die quadratische Formel:

Anwendungsbeispiel 4.26 (Reihen- und Parallelschaltung von Widerständen).

An eine Spannungsquelle $U = 220V$ werden zwei Widerstände R_1 und R_2 einmal in Reihe und einmal in Parallelschaltung angeschlossen. Im ersten Falle ist die Stromstärke $I_1 = 0.9A$ im zweiten Falle $I_2 = 6A$. Wie groß sind R_1 und R_2 ?

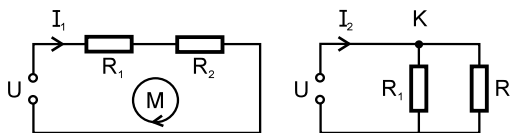


Abb. 4.12. Schaltungsdiagramme

Für die Reihenschaltung gilt der Maschensatz: $U = I_1 R_1 + I_1 R_2$ und für die Parallelschaltung der Knotensatz für den Knoten K : $I_2 = \frac{U}{R_1} + \frac{U}{R_2}$. Löst man

die erste Gleichung nach R_2 auf und setzt dies in die zweite ein, folgt nach Umformungen

$$R_1^2 - \frac{U}{I_1} R_1 + \frac{U^2}{I_1 I_2} = 0.$$

Gesucht sind also die Nullstellen des Polynoms $f(x) = x^2 - \frac{220}{0,9}x + \frac{220^2}{0,9 \cdot 6}$. Mit

$$x_{1/2} = -\frac{p}{2} \pm \sqrt{\left(\frac{p}{2}\right)^2 - q}$$

erhält man als Nullstellen $x_1 = 44.9$ und $x_2 = 199.5$. Wählt man $R_1 = 44.9\Omega$, so ist $R_2 = 199.5\Omega$. Aus Symmetriegründen können R_1 und R_2 auch vertauscht werden. $\square$

Für $n = 3$ und 4 sind die Lösungsformeln wesentlich komplizierter und für $n \geq 5$ gibt es keine geschlossenen Lösungsformeln mehr. Daher ist man bei der Suche nach Nullstellen bei Polynomen höheren Grades auf numerische Verfahren wie z.B. die Intervallhalbierungs-Methode 6.4 oder das Newton-Verfahren 7.9 angewiesen. Wenn man allerdings eine Nullstelle x_0 erraten kann, so liefert das Horner-Schema ein nützliches Verfahren zur Reduktion des Problems, da man durch Abspalten des Linearfaktors $(x - x_0)$ ein Polynom vom Grad $n - 1$ erhält.

Beispiel 4.27. Gesucht sind die Nullstellen des Polynoms

$$f(x) = 3x^3 + 3x^2 - 3x - 3.$$

Durch Erraten findet man eine Nullstelle bei $x_1 = -1$:

$$\begin{array}{r} 3 \quad 3 \quad -3 \quad -3 \\ + \quad -3 \quad 0 \quad 3 \\ \hline 3 \quad 0 \quad -3 \quad 0 = f(-1) \end{array}$$

Hiernach gilt also

$$3x^3 + 3x^2 - 3x - 3 = (x + 1)(3x^2 - 3).$$

Durch Anwenden der quadratischen Formel berechnen sich die beiden weiteren Nullstellen $x_2 = 1$ und $x_3 = -1$. Damit folgt die *Linearfaktorenzerlegung* des Polynoms:

$$3x^3 + 3x^2 - 3x - 3 = 3(x + 1)(x + 1)(x - 1). \quad \square$$

Wir fassen das Ergebnis in folgendem allgemein gültigen Satz zusammen:

Satz: Besitzt ein Polynom n -ten Grades n Nullstellen $x_1, x_2, \dots, x_n$, dann lässt es sich als Produkt aus n **Linearfaktoren** darstellen:

$$\begin{aligned} f(x) &= a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 \\ &= a_n (x - x_1) (x - x_2) \cdot \dots \cdot (x - x_n). \end{aligned}$$

Bemerkungen:

- (1) **⚠ Achtung:** Man darf den Koeffizient a_n in der Produktdarstellung nicht vergessen!
- (2) **⚠ Achtung:** Bei doppelten Nullstellen tritt der zugehörige Linearfaktor doppelt, bei dreifachen Nullstellen dreifach usw. auf!

➤ 4.2.5 Interpolationspolynome mit dem Newton-Algorithmus

In naturwissenschaftlichen Anwendungen stellt sich oftmals das Problem: Von einem unbekannten funktionalen Zusammenhang sind $(n + 1)$ Messpunkte gegeben: $P_1(x_1, y_1); P_2(x_2, y_2); \dots; P_{n+1}(x_{n+1}, y_{n+1})$. Gesucht sind die Funktionswerte an Zwischenstellen. Wie bereits nach dem ersten Satz bekannt ist, gibt es genau ein Polynom vom Grad n , welches durch diese $(n + 1)$ Messwerte geht. Diese Polynome werden auch als *Interpolationspolynome* bezeichnet, da sich mit ihnen näherungsweise beliebige Zwischenwerte der unbekannten Funktion im Bereich $[x_1, x_{n+1}]$ berechnen (=interpolieren) lassen.

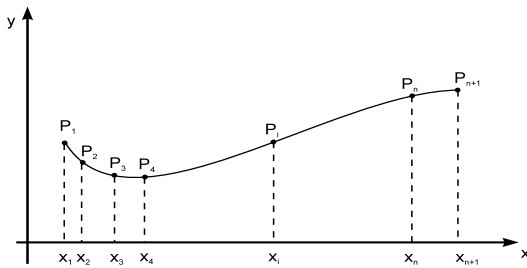


Abb. 4.13. Interpolationspolynom

Durch das Lagrange Interpolationspolynom erhält man das eindeutig bestimmte Polynom vom Grade höchstens n , welches diese Eigenschaft hat. Allerdings ist der Rechenaufwand zur Berechnung der Koeffizienten erheblich. Ein einfacheres Rechenschema zur Bestimmung des Interpolationspolynoms geht auf einen Ansatz von Newton zurück:

$$\begin{aligned} f(x) &= a_0 + a_1(x - x_1) + a_2(x - x_1)(x - x_2) + \dots \\ &\quad + a_n(x - x_1)(x - x_2) \cdot \dots \cdot (x - x_n). \end{aligned}$$

Die Koeffizienten $a_0, a_1, \dots, a_n$ können durch diesen Ansatz *iterativ* bestimmt werden:

$$y_1 = f(x_1) = a_0 \hookrightarrow \boxed{a_0 = y_1.}$$

$$y_2 = f(x_2) = a_0 + a_1(x_2 - x_1) \hookrightarrow \boxed{a_1 = \frac{y_2 - y_1}{x_2 - x_1}.}$$

$$y_3 = f(x_3) = a_0 + a_1(x_3 - x_1) + a_2(x_3 - x_1)(x_3 - x_2)$$

$$\begin{aligned} a_2 &= \frac{y_3 - a_0 - a_1(x_3 - x_1)}{(x_3 - x_1)(x_3 - x_2)} = \frac{y_3 - y_1 - \frac{y_2 - y_1}{x_2 - x_1}(x_3 - x_1)}{(x_3 - x_1)(x_3 - x_2)} \\ &= \frac{1}{x_3 - x_1} \left\{ \frac{y_3 - y_1}{x_3 - x_2} - \frac{y_2 - y_1}{x_2 - x_1} \frac{x_3 - x_1}{x_3 - x_2} \right\} \\ &= \frac{1}{x_3 - x_1} \left\{ \frac{y_3 - y_2}{x_3 - x_2} + \frac{y_2 - y_1}{x_3 - x_2} - \frac{y_2 - y_1}{x_2 - x_1} \frac{x_3 - x_1}{x_3 - x_2} \right\} \\ &= \frac{1}{x_3 - x_1} \left\{ \frac{y_3 - y_2}{x_3 - x_2} - \frac{y_2 - y_1}{x_2 - x_1} \right\}. \end{aligned}$$

$$\text{Setzt man } D_{2,1} = \frac{y_2 - y_1}{x_2 - x_1}, D_{3,2} = \frac{y_3 - y_2}{x_3 - x_2} \hookrightarrow \boxed{a_2 = \frac{D_{3,2} - D_{2,1}}{x_3 - x_1}.}$$

Mit vollständiger Induktion zeigt man, dass mit den Abkürzungen

$$D_{4,3,2} = \frac{D_{4,3} - D_{3,2}}{x_4 - x_2}; D_{3,2,1} = \frac{D_{3,2} - D_{2,1}}{x_3 - x_1} \quad \text{usw. gilt}$$

$$\boxed{a_{k-1} := D_{k,\dots,1} \text{ für } k \geq 1.}$$

Die Ausdrücke $D_{k,\dots,1}$ heißen *dividierte Differenzen* und die Berechnung erfolgt mittels des folgenden Schemas:

k	x_k	y_k				
1	x_1	y_1				
2	x_2	y_2	$\searrow$	$\boxed{D_{2,1} = \frac{y_2 - y_1}{x_2 - x_1}}$		
3	x_3	y_3	$\rightarrow$	$D_{3,2} = \frac{y_3 - y_2}{x_3 - x_2}$	$\searrow$	$\boxed{D_{3,2,1} = \frac{D_{3,2} - D_{2,1}}{x_3 - x_1}}$
4	x_4	y_4	$\rightarrow$	$D_{4,3} = \frac{y_4 - y_3}{x_4 - x_3}$	$\searrow$	$D_{4,3,2} = \frac{D_{4,3} - D_{3,2}}{x_4 - x_2}$
$\vdots$						$\rightarrow \dots$
$\vdots$			$\searrow$			
$n+1$	x_{n+1}	y_{n+1}	$\rightarrow$	$D_{n+1,n} = \frac{y_{n+1} - y_n}{x_{n+1} - x_n}$	$\rightarrow$	$\dots \rightarrow \dots$

Die Zahlen $y_1 = a_0$, $D_{2,1} = a_1$, $D_{3,2,1} = a_2$, $D_{4,3,2,1} = a_3, \dots$, $D_{n+1,\dots,1} = a_n$ sind dann die gesuchten Koeffizienten des **Newtonschen Interpolationspolynoms**.

Bemerkung: Der große **Vorteil** des Newton-Verfahrens gegenüber dem Lagrangeschen Interpolationspolynom besteht darin, dass bei der Hinzunahme noch weiterer Messpunkte nur weitere Zeilen in die Tabelle aufgenommen werden müssen. Die bereits berechneten Koeffizienten bleiben gültig!

Beispiel 4.28 (Mit MAPLE-Worksheet).

- (i) Das Ergebnis einer Messreihe liefert die Werte $P_1(0, -12)$, $P_2(2, 16)$ und $P_3(5, 28)$. Gesucht ist das Interpolationspolynom durch diese 3 Punkte.

Ansatz: $f(x) = a_0 + a_1(x - x_1) + a_2(x - x_1)(x - x_2)$.

Bestimmung der Koeffizienten:

k	x_k	y_k	
1	0	-12	
2	2	16	$\frac{16+12}{2-0} = $ 14
3	5	28	$\frac{28-16}{5-2} = 4$ $\frac{4-14}{5-0} = $ -2

Damit erhält man das quadratische Polynom

$$f(x) = -12 + 14x - 2x(x - 2) = -2x^2 + 18x - 12.$$

- (ii) Durch die Hinzunahme eines weiteren Messpunktes $P_4(7, -54)$ kann man das bestehende Schema erweitern

k	x_k	y_k	
1	0	-12	
2	2	16	14
3	5	28	4 -2
4	7	-54	$\frac{-54-28}{7-5} = -41$ $\frac{-41-4}{7-2} = -9$ $\frac{-9+2}{7-0} = $ -1

und man findet als Interpolationspolynom vom Grade 3

$$\begin{aligned}
 f(x) &= a_0 + a_1(x - x_1) + a_2(x - x_1)(x - x_2) \\
 &\quad + a_3(x - x_1)(x - x_2)(x - x_3) \\
 &= -12 + 14(x - 0) - 2(x)(x - 2) - 1(x)(x - 2)(x - 5) \\
 &= -x^3 + 5x^2 + 8x - 12.
 \end{aligned}$$

- (iii) Im **MAPLE-Worksheet** wird bei gegebenen Wertepaaren das Interpolationspolynom bestimmt und mit den Werten graphisch dargestellt. □

4.3 Rationale Funktionen

Ziel dieses Abschnitts ist, rationale Funktionen qualitativ durch die Bestimmung der Null- und Polstellen sowie dem asymptotischen Verhalten zu charakterisieren.

In Physik und Technik werden viele Vorgänge von Funktionen beschrieben, die sich als Quotient zweier Polynome darstellen. So ergibt sich z.B. bei einer Sammellinse mit Brennweite f die Bildweite b als $b(x) = \frac{f x}{x-f}$, wenn x die Gegenstandsweite ist.

Definition: Unter einer **rationalen Funktion (gebrochenrationalen Funktion)** versteht man eine Funktion f , die sich als Quotient zweier Polynomfunktionen $g(x)$ und $h(x)$ darstellen lässt

$$f : \mathbb{R} \setminus \{x \in \mathbb{R} : h(x) = 0\} \rightarrow \mathbb{R}$$

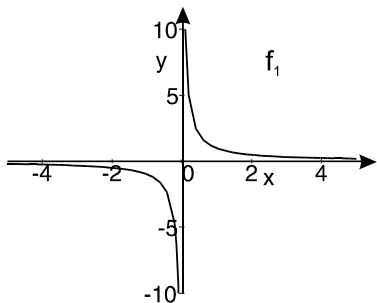
mit

$$x \mapsto f(x) = \frac{g(x)}{h(x)} = \frac{a_m x^m + \dots + a_1 x + a_0}{b_n x^n + \dots + b_1 x + b_0} = \frac{\sum_{i=0}^m a_i x^i}{\sum_{j=0}^n b_j x^j}.$$

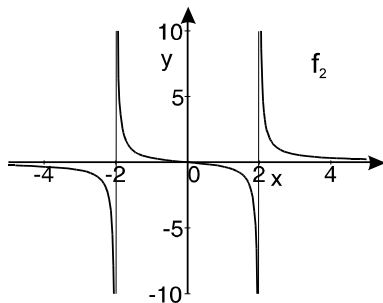
Dabei unterscheidet man analog zu Brüchen zwischen *ganzrationalen Funktionen* ($n = 0$) (= Polynome), *echt gebrochenrationalen Funktionen* ($m < n$) und *unecht gebrochenrationalen Funktionen* ($m \geq n$).

Beispiele 4.29:

- ① $f_1 : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$ mit $x \mapsto \frac{1}{x}$.
 ② $f_2 : \mathbb{R} \setminus \{-2, 2\} \rightarrow \mathbb{R}$ mit $x \mapsto \frac{x}{x^2 - 4}$.



① Graph von $f_1(x) = \frac{1}{x}$



② Graph von $f_2(x) = \frac{x}{x^2 - 4}$

► 4.3.1 Definitionslücken, Nullstellen, Pole

Sei $f(x) = \frac{g(x)}{h(x)}$ eine gebrochenrationale Funktion. x_0 ist **Nullstelle** von f , wenn $g(x_0) = 0$ und $h(x_0) \neq 0$. In den Nullstellen des Nenners ist die Funktion f nicht definiert. Man spricht daher von einer **Definitionslücke**, wenn $h(x_0) = 0$. Aber nicht in allen Definitionslücken strebt die Funktion gegen Unendlich. Man unterscheidet zwischen *hebbaren Definitionslücken* und **Polen**: x_0 ist eine *Polstelle* (*Pol*), wenn in der unmittelbaren Umgebung von x_0 die Funktionswerte betragsmäßig über alle Grenzen anwachsen. Der Funktionsgraph schmiegt sich dabei asymptotisch an die in der Polstelle errichtete Parallele zur y -Achse an.

Falls Zähler und Nennerpolynom eine gemeinsame Nullstelle x_0 besitzen, so enthalten beide Polynome $(x - x_0)$ als Linearfaktor. Gemeinsamen Faktoren werden gekürzt. Definitionslücken können so gegebenenfalls durch Kürzen behoben und der Definitionsbereich erweitert werden.

Bestimmung der Null- und Polstellen

- (1) Man zerlegt Zähler und Nennerpolynom soweit möglich in Linearfaktoren und kürzt gemeinsame Faktoren.
- (2) Die im Zähler verbleibenden Linearfaktoren ergeben die Nullstellen der Funktion, die im Nenner verbleibenden Linearfaktoren liefern die Polstellen.

Beispiel 4.30. Gesucht sind die Definitionslücken, Nullstellen und Polstellen der Funktion $f(x) = \frac{2x^3 + 2x^2 - 32x + 40}{x^3 + 2x^2 - 13x + 10}$.

Zur Bestimmung dieser Stellen zerlegen wir sowohl das Nenner- als auch das Zählerpolynom in Linearfaktoren:

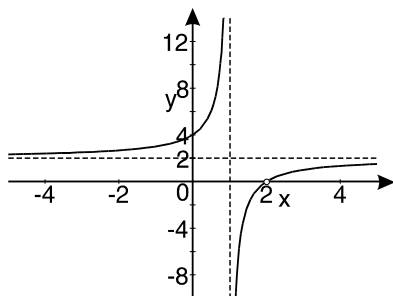
$$\begin{aligned} 2x^3 + 2x^2 - 32x + 40 &= 2(x-2)^2(x+5) \\ x^3 + 2x^2 - 13x + 10 &= (x-1)(x-2)(x+5). \end{aligned}$$

Damit lässt sich die gebrochenrationale Funktion in Faktoren schreiben als

$$\Rightarrow f(x) = \frac{2x^3 + 2x^2 - 32x + 40}{x^3 + 2x^2 - 13x + 10} = \frac{2(x-2)^2(x+5)}{(x-1)(x-2)(x+5)} = \frac{2(x-2)}{x-1}.$$

Alle Nullstellen des Nenner im *ungekürzten* Zustand liefern die Definitionslücken; die verbleibenden Nullstellen des Nenners sind die Polstellen der Funktion. Die Nullstellen des Zählers liefern, sofern sie nicht gleichzeitig eine Definitionslücke sind, die Nullstellen der Funktion. Zusammenfassend erhalten wir den qualitativen Verlauf der Funktion:

Definitionslücken	: 1, 2, -5
Nullstellen	: (2)
Polstellen	: 1
Hebbare Definitionslücken	: 2, -5



4.3.2 Verhalten im Unendlichen

Gerade für die Diskussion der Anwendungsbeispiele ist die Bestimmung des Verhaltens der Funktionen für $x \rightarrow \pm\infty$ von Interesse. Wir müssen dabei drei verschiedene Fälle betrachten. Entscheidend dabei ist, wie sich der Grad des Zählerpolynoms zum Grad des Nennerpolynoms verhält, wie die folgenden Beispiele zeigen:

Beispiele 4.31.

$$(1.) f_1(x) = \frac{5x^4 + x^3 + 4}{2x^5 + x + 2}$$

$$(2.) f_2(x) = \frac{5x^4 + x^3 + 4}{2x^4 + x + 2}$$

$$(3.) f_3(x) = \frac{5x^4 + x^3 + 4}{2x^3 + x + 2}$$

In allen 3 Fällen ist $f(x) \xrightarrow{x \rightarrow \infty} \infty$. Die folgende Diskussion zeigt aber, dass man jeweils unterschiedliche Ergebnisse erhält. Dazu erweitert man den entsprechenden Funktionsausdruck mit $\frac{1}{x^k}$, wenn k der höchste auftretende Exponent:

$$\lim_{x \rightarrow \infty} f_1(x) = \lim_{x \rightarrow \infty} \frac{5x^4 + x^3 + 4}{2x^5 + x + 2} \cdot \frac{\frac{1}{x^5}}{\frac{1}{x^5}} = \lim_{x \rightarrow \infty} \frac{\frac{5}{x} + \frac{1}{x^2} + \frac{4}{x^5}}{2 + \frac{1}{x^4} + \frac{2}{x^5}} = \frac{0}{2} = 0.$$

$$\lim_{x \rightarrow \infty} f_2(x) = \lim_{x \rightarrow \infty} \frac{5x^4 + x^3 + 4}{2x^4 + x + 2} \cdot \frac{\frac{1}{x^4}}{\frac{1}{x^4}} = \lim_{x \rightarrow \infty} \frac{5 + \frac{1}{x} + \frac{4}{x^4}}{2 + \frac{1}{x^3} + \frac{2}{x^4}} = \frac{5}{2}.$$

$$\lim_{x \rightarrow \infty} f_3(x) = \lim_{x \rightarrow \infty} \frac{5x^4 + x^3 + 4}{2x^3 + x + 2} \cdot \frac{\frac{1}{x^4}}{\frac{1}{x^4}} = \lim_{x \rightarrow \infty} \frac{5 + \frac{1}{x} + \frac{4}{x^4}}{\frac{2}{x} + \frac{1}{x^3} + \frac{2}{x^4}} = \frac{5}{0} = \infty. \quad \square$$

Diese Unterscheidung lässt sich auf jede rationale Funktion anwenden:

Verhalten einer rationalen Funktion $f(x) = \frac{g(x)}{h(x)}$ im Unendlichen

$$(1) \text{ grad } g < \text{grad } h \Rightarrow f(x) \rightarrow 0 \quad \text{für } x \rightarrow \pm\infty.$$

$$(2) \text{ grad } g = \text{grad } h \Rightarrow f(x) \rightarrow \frac{a_n}{b_n} \quad \text{für } x \rightarrow \pm\infty.$$

$$(3) \text{ grad } g > \text{grad } h \Rightarrow f(x) \rightarrow \pm\infty \quad \text{für } x \rightarrow \pm\infty.$$

Im letzten Fall zerlegt man die unecht gebrochenrationale Funktion durch Polynomdivision in eine Polynomfunktion $p(x)$ und eine echt gebrochenrationale Funktion $r(x)$: $f(x) = p(x) + r(x)$ mit $r(x) \rightarrow 0$ für $x \rightarrow \infty$. Die Funktion $f(x)$ nähert sich für $x \rightarrow \infty$ an die Funktion $p(x)$ an, da der Rest $r(x)$ gegen 0 geht! Man nennt $p(x)$ die **Asymptote** von f .

Beispiel 4.32.
$$f(x) = \frac{\frac{1}{2}x^3 - \frac{3}{2}x + 1}{x^2 + 3x + 2}.$$

Wir zerlegen Zähler- und Nennerpolynom in Linearfaktoren:

Die Nullstellen des Zählers sind 1 (doppelt) und -2.

Die Nullstellen des Nenners sind -1 und -2.

$$\Rightarrow f(x) = \frac{\frac{1}{2}(x-1)^2(x+2)}{(x+1)(x+2)} = \frac{\frac{1}{2}(x-1)^2}{x+1}.$$

Im gekürzten Ausdruck lassen sich nun die Nullstellen und Polstellen der Funktion identifizieren:

Nullstelle : $x = 1$ (doppelt)

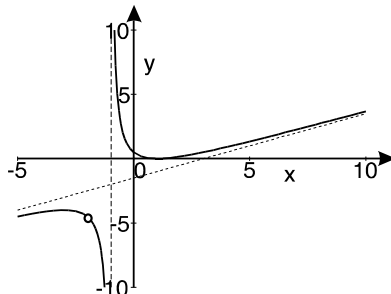
Polstelle : $x = -1$ (einfach).

Um das Verhalten im Unendlichen zu bestimmen, zerlegen wir durch Polynomdivision die Funktion f in eine Polynomfunktion $p(x)$ und eine echt gebrochenrationale Funktion $r(x)$:

$$\begin{array}{r} \left(\frac{1}{2}x^2 - 1x + \frac{1}{2}\right) : (x+1) = \frac{1}{2}x - \frac{3}{2} + \frac{2}{x+1} \\ \underline{-(\frac{1}{2}x^2 + \frac{1}{2}x)} \\ -\frac{3}{2}x + \frac{1}{2} \\ \underline{-(-\frac{3}{2}x - \frac{3}{2})} \\ 2 \end{array}$$

Damit ist $f(x) = p(x) + r(x)$ mit der Asymptote $p(x) = \frac{1}{2}x - \frac{3}{2}$ und dem

Rest $r(x) = \frac{2}{x+1} \xrightarrow{x \rightarrow \infty} 0$.



Graph der Funktion $\frac{1/2 x^3 - 3/2 x + 1}{x^2 + 3x + 2}$

➤ 4.3.3 Anwendung: Übertragung bei Wechselstrom-Kreisen

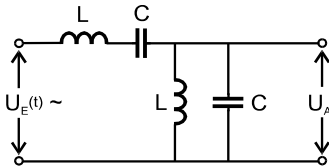


Abb. 4.14. LC-Parallel-Kreis

Die nebenstehende LC-Schaltung hat wie jede RCL-Schaltung die folgende Eigenschaft: Ist die Eingangsspannung $U_E(t)$ eine Wechselspannung mit Frequenz ω , so ist auch die Ausgangsspannung eine Wechselspannung mit Frequenz ω , aber phasenverschoben und mit anderer Amplitude. Diese Amplitude hängt von der Frequenz der Eingangsspannung ab. Das Amplitudenverhältnis $H(\omega)$ ist gegeben durch (vgl. CD-Kapitel [Übertragungsfunktion für RCL-Filterschaltungen](#)):

$$\left| \frac{U_A(t)}{U_E(t)} \right| = |H(\omega)| \text{ mit}$$

$$H(\omega) = \frac{-\omega^2 LC}{\omega^4 L^2 C^2 - 3\omega^2 LC + 1}.$$

$H(\omega)$ ist eine echt gebrochenrationale Funktion in ω . Die Nullstelle liegt bei $\omega = 0$ und um die Polstellen der Funktion zu bestimmen, setzen wir

$$\omega^4 L^2 C^2 - 3\omega^2 LC + 1 = 0:$$

$$Z = \omega^2$$

$$Z^2 (LC)^2 - 3ZLC + 1 = 0$$

$$Z_{1/2} = \frac{3}{2} \frac{1}{LC} \pm \sqrt{\frac{9}{4} \left(\frac{1}{LC}\right)^2 - \left(\frac{1}{LC}\right)^2} = \left(\frac{3}{2} \pm \frac{\sqrt{5}}{2}\right) \frac{1}{LC}$$

$$Z_1 = 2,62 \frac{1}{LC} \Rightarrow \omega_{1/2} = \pm 1,61 \sqrt{\frac{1}{LC}}$$

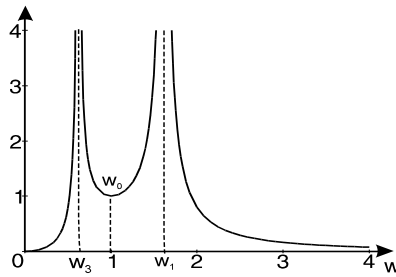
$$Z_2 = 0,76 \frac{1}{LC} \Rightarrow \omega_{3/4} = \pm 0,87 \sqrt{\frac{1}{LC}}.$$

$H(\omega)$ ist achsensymmetrisch zur y -Achse, da

$$H(-\omega) = \frac{-(-\omega)^2 LC}{(-\omega)^4 L^2 C^2 - 3(-\omega)^2 LC + 1} = \frac{-\omega^2 LC}{\omega^4 L^2 C^2 - 3\omega^2 LC + 1} = H(\omega).$$

Der Grad des Zählerpolynoms ist kleiner als der Grad des Nennerpolynoms. Daher gilt: $H(\omega) \rightarrow 0$ für $\omega \rightarrow \infty$. Außerdem ist $H(\omega) \rightarrow 0$ für $\omega \rightarrow 0$. Qualitativ erhalten wir von $|H(\omega)|$ für positive Frequenzen ($L = C = 1$) den graphischen Verlauf, wie er in Abb. 4.15 dargestellt ist.

Diskussion (Mit MAPLE-Worksheet): Für tiefe Frequenzen (ω klein) ist das Amplitudenverhältnis von Ausgangsspannung zu Eingangsspannung klein: Tiefe Frequenzen werden nicht gut übertragen. Für hohe Frequenzen (ω groß) ist $|H(\omega)|$ ebenfalls klein: Hohe Frequenzen werden ebenfalls stark gedämpft. Frequenzen zwischen ω_1 und ω_2 werden etwa mit dem Faktor 1 übertragen.

Abb. 4.15. Übertragungsfunktion $|H(\omega)|$

Dies ist das typische Verhalten eines Bandpasses, der tiefe und hohe Frequenzen dämpft und Frequenzen in einem Band zwischen ω_1 und ω_2 überträgt.

Bemerkung: Die Modellierung von $H(\omega)$ erfolgte unter der Voraussetzung, dass der Ohmsche Widerstand $R = 0$ ist. Dadurch kommt es zu nichtphysikalischen Polstellen bei $\omega = \omega_1$ und $\omega = \omega_3$ bzw. zum Effekt, dass nahe den Polstellen die Amplitude des Ausgangssignals größer als die des Eingangssignals ist. Eine vollständige Diskussion (mit Ohmschem Widerstand) zeigt, dass $|H(\omega)| \leq 1$ und die Polstellen entfallen. $\square$

4.4 Potenz- und Wurzelfunktionen

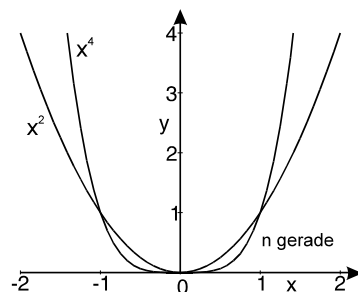
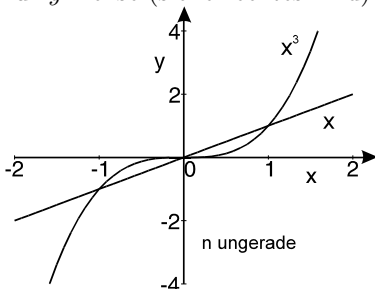
4.4

Definition: Polynomfunktionen der Form

$$p: \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto x^n \quad (n \in \mathbb{N})$$

nennt man auch **Potenzfunktionen**.

Bei den Potenzfunktionen lassen sich qualitativ zwei Fälle unterscheiden, nämlich n gerade und n ungerade. Für ungerades n ist die Potenzfunktion streng monoton wachsend und punktsymmetrisch zum Ursprung (siehe linkes Bild). Für gerades n liegt keine Monotonie vor, die Potenzfunktion ist achsensymmetrisch zur y -Achse (siehe rechtes Bild).



Schränken wir den Definitionsbereich der Potenzfunktionen auf die positiven, reellen Zahlen einschließlich der Null ein, so ist

$$\begin{array}{ccc} p : \mathbb{R}_{\geq 0} & \rightarrow & \mathbb{R}_{\geq 0} \\ x & \mapsto & x^n \end{array}$$

für alle $x \in \mathbb{R}_{\geq 0}$ eine streng monoton wachsende Funktion mit dem Wertebereich $\mathbb{R}_{\geq 0}$. Daher ist diese Funktion umkehrbar. Mit Hilfe der Umformung

$$y = x^n \rightarrow x = \sqrt[n]{y} \rightarrow y = \sqrt[n]{x}$$

erhalten wir die Umkehrfunktion

$$p^{-1} : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0} \quad \text{mit} \quad x \mapsto \sqrt[n]{x}.$$

Definition: Die Funktion

$$\omega : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0} \quad \text{mit} \quad x \mapsto \sqrt[n]{x}$$

heißt **n-te Wurzelfunktion** ($n \in \mathbb{N}$).

Spezialfall: Potenzfunktionen mit ungeradem Exponent $p(x) = x^{2m+1}$, $m \in \mathbb{N}$, sind auf ganz $\mathbb{R}$ streng monoton wachsend und haben als Wertebereich ebenfalls $\mathbb{R}$. Daher existiert die Umkehrfunktion auf ganz $\mathbb{R}$

$$\omega : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \sqrt[2m+1]{x}.$$



Visualisierung mit MAPLE: Auf der CD-Rom befinden sich zwei

Animationen in denen sowohl die Potenz- als auch Wurzelfunktionen mit wachsendem n gezeigt werden. Man erkennt, dass die Potenzfunktionen mit wachsendem n immer steiler anwachsen, während die Wurzelfunktionen immer flacher werden.

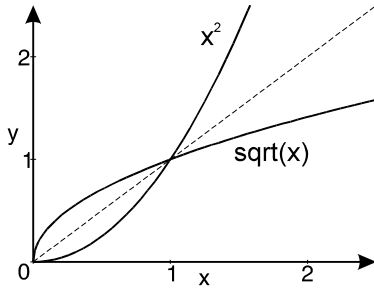
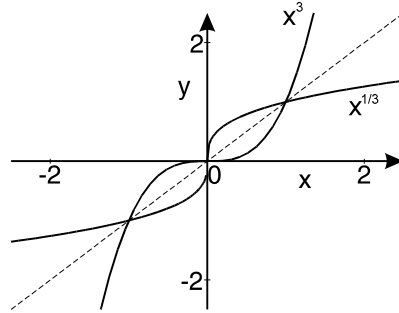
Beispiele 4.33:

① $p_2 : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$ mit $x \mapsto x^2$ hat als Umkehrfunktion

$$\omega_2 : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0} \quad \text{mit} \quad x \mapsto \sqrt{x} = x^{\frac{1}{2}}.$$

② $p_3 : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto x^3$ hat als Umkehrfunktion

$$\omega_3 : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \sqrt[3]{x} = x^{\frac{1}{3}}.$$

① Graph von x^2 und $\sqrt{x}$ ② Graph von x^3 und $x^{\frac{1}{3}}$ **Anwendungsbeispiel 4.34 (Fallgeschwindigkeit eines Körpers).**

Ein Körper der Masse m fällt frei aus der Höhe h_0 mit der Anfangsgeschwindigkeit $v_0 = 0$. Zu jedem Zeitpunkt t gilt für die Bewegung, dass die Gesamtenergie (Summe aus kinetischer und potentieller Energie) konstant bleibt. Es gilt:

$$\left. \begin{aligned} E(t=0) &= m \cdot g \cdot h_0 + \frac{1}{2} m v_0^2 = m \cdot g \cdot h_0 \\ E(t>0) &= m \cdot g \cdot h + \frac{1}{2} m v^2 \end{aligned} \right\} E(t=0) = E(t>0).$$

$$\Rightarrow m \cdot g \cdot h_0 = m \cdot g \cdot h + \frac{1}{2} m v^2.$$

Die Geschwindigkeit v ergibt sich dann bei der Höhe h als Wurzelfunktion

$$v = \sqrt{2g(h_0 - h)}.$$

□

Potenzfunktion mit rationalem Exponenten

Definition: Eine Funktion

$$f: \mathbb{R}_{>0} \rightarrow \mathbb{R} \quad \text{mit} \quad f(x) = \sqrt[n]{x^m} = x^{\frac{m}{n}} \quad m \in \mathbb{Z}, n \in \mathbb{N}$$

heißt *Potenzfunktion mit rationalem Exponent*.

Diesen Begriff der Potenzfunktion werden wir in §4.5 auf beliebige, reelle Exponenten mit Hilfe der Exponential- und Logarithmusfunktion erweitern. Als Spezialfall sind in der Klasse der Potenzfunktionen mit rationalem Exponenten die Funktionen x^{-1} , x^{-2} usw. enthalten.

4.5 Exponential- und Logarithmusfunktion

Eine der wichtigsten wenn nicht sogar die wichtigste Funktion in der Physik ist die Exponentialfunktion. Wir werden die Exponentialfunktion zusammen mit ihrer Umkehrfunktion, der Logarithmusfunktion, vorstellen und deren wichtigsten Eigenschaften diskutieren.

4.5.1 Exponentialfunktion

Die zur Beschreibung naturwissenschaftlicher Phänomene wichtigste Funktion ist die Exponentialfunktion:

Definition: Die Funktion

$$\exp : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto e^x$$

heißt **Exponentialfunktion**. $e \approx 2.718281828$ ist die Eulersche Zahl.

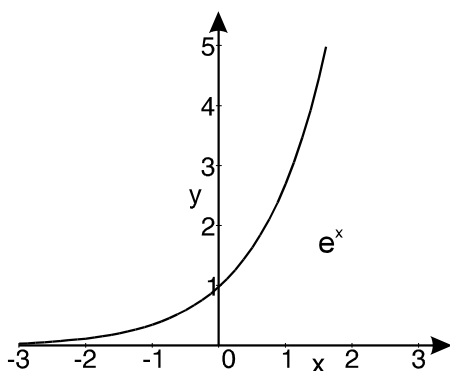


Abb. 4.16. Graph der Exponentialfunktion e^x

Eigenschaften der Exponentialfunktion sind:

	e^x
Definitionsbereich	$\mathbb{R}$
Wertebereich	$\mathbb{R}_{>0}$
Monotonie	streng monoton wachsend
Asymptote	$y = 0$ für $x \rightarrow -\infty$

Für die Exponentialfunktion gelten die **Regeln**:

$$(1) \quad e^0 = 1$$

$$(2) \quad e^{x+y} = e^x \cdot e^y$$

$$(3) \quad e^{-x} = (e^x)^{-1}, \quad e^{n \cdot x} = (e^x)^n$$

Beispiele 4.35:

- ① Die Funktionen $f_a(x) = e^{ax}$ verhalten sich für $a > 0$ qualitativ wie die Exponentialfunktion e^x : Für $x \rightarrow \infty$ gehen sie gegen Unendlich und für $x \rightarrow -\infty$ gegen Null.
- ② Die Funktionen $f_a(x) = e^{-ax}$ verhalten sich für $a > 0$ qualitativ wie die Exponentialfunktion e^{-x} : Für $x \rightarrow \infty$ gehen sie gegen Null und für $x \rightarrow -\infty$ gegen Unendlich.

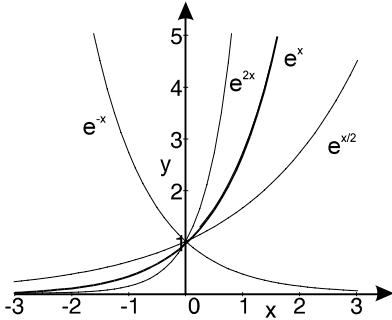


Abb. 4.17. Graphen der Exponentialfunktionen

Anwendungsbeispiel 4.36 (Auftreten der Exponentialfunktion).

- ① **Radioaktiver Zerfall:** Beim Zerfall radioaktiver Atomkerne wird die Zahl $n(t)$ der zur Zeit t noch nicht zerfallenen Kerne durch das Zerfallsgesetz

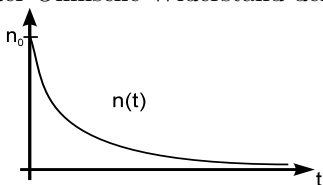
$$n(t) = n_0 e^{-\lambda t}$$

beschrieben. Dabei ist n_0 die Anzahl der zu Beginn ($t = 0$) vorhandenen Atomkerne und $\lambda > 0$ die für den Zerfall typische Zerfallskonstante.

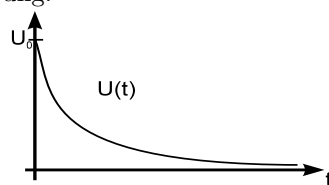
- ② **Entladung eines Plattenkondensators:** Beim Entladen eines Plattenkondensators ist die Spannung am Kondensator $U(t)$ zum Zeitpunkt t gegeben durch

$$U(t) = U_0 e^{-\frac{1}{RC}t}.$$

Dabei ist U_0 die Kondensatorspannung zur Zeit $t = 0$ und C die Kapazität, R der Ohmsche Widerstand der Schaltung.



① Anzahl der Atomkerne $n(t)$



② Spannung am Kondensator $U(t)$

➤ 4.5.2 Logarithmusfunktion

Die Exponentialfunktion $\exp : \mathbb{R} \rightarrow \mathbb{R}_{>0}$ mit $x \mapsto e^x$ ist auf dem gesamten Definitionsbereich streng monoton wachsend. Folglich existiert auf dem Wertebereich $\mathbb{R}_{>0}$ die Umkehrfunktion.

Definition: Die Umkehrfunktion zur Exponentialfunktion wird **natürlicher Logarithmus** genannt:

$$\ln : \mathbb{R}_{>0} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \ln x.$$

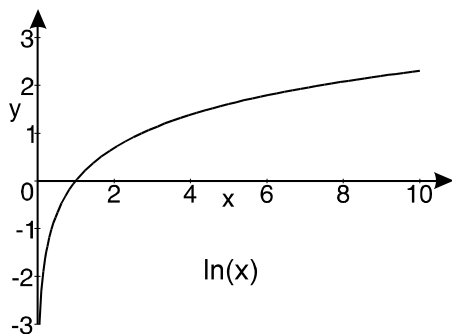


Abb. 4.18. Graph der Logarithmusfunktion $\ln(x)$

Eigenschaften der Logarithmusfunktion sind:

	$\ln(x)$
Definitionsbereich	$\mathbb{R}_{>0}$
Wertebereich	$\mathbb{R}$
Nullstellen	$x_0 = 1$
Monotonie	streng monoton wachsend
Asymptoten	$x = 0$

Rechenregeln für die Logarithmusfunktion. Die Rechenregeln ergeben sich direkt aus den Regeln der Exponentialfunktion. Für $x, y \in \mathbb{R}_{>0}$ gilt

- (1) $\ln(1) = 0$
- (2) $\ln(x \cdot y) = \ln x + \ln y$
- (3) $\ln(x^n) = n \ln x$
- (4) $\ln(e^x) = x$ bzw. $e^{\ln x} = x$

Anwendungsbeispiel 4.37 (Auftreten der Logarithmusfunktion).

- ① **Halbwertszeit τ einer radioaktiven Substanz:** Unter der Halbwertszeit τ einer radioaktiven Substanz versteht man die Zeit, nach der die Hälfte der radioaktiven Kerne zerfallen ist: $n(\tau) = \frac{1}{2}n_0$. Nach Beispiel 4.36 ① ist $n(t) = n_0 e^{-\lambda t}$, also gilt für $t = \tau$:

$$\frac{1}{2}n_0 = n_0 e^{-\lambda \tau}$$

Wir dividieren durch n_0 , wenden den Logarithmus an und lösen nach τ auf

$$\ln \frac{1}{2} = \ln(e^{-\lambda \tau}) = -\lambda \tau$$

$$\Leftrightarrow \tau = -\frac{1}{\lambda} \ln \frac{1}{2} = -\frac{1}{\lambda} (\ln 1 - \ln 2).$$

Die Halbwertszeit τ ergibt sich somit zu

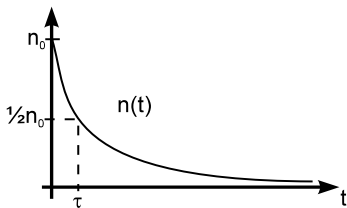
$$\tau = \frac{1}{\lambda} \ln 2.$$

- ② **Abklingzeit eines Kondensators:** Unter der Abklingzeit τ_a eines Kondensators versteht man die Zeit, nach der die Spannung am Kondensator auf z.B. $\frac{1}{e}$ -tel des maximalen Spannungswertes abgefallen ist: $U(\tau_a) = \frac{1}{e}U_0$. Nach Beispiel 4.36 ② ist $U(t) = U_0 e^{-\frac{1}{RC}t}$, also gilt für $t = \tau_a$:

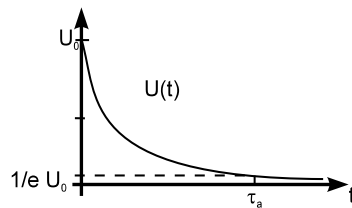
$$\frac{1}{e}U_0 = U_0 e^{-\frac{1}{RC}\tau_a} \Leftrightarrow \ln \frac{1}{e} = \ln(e^{-\frac{1}{RC}\tau_a}) = -\frac{1}{RC}\tau_a$$

Die Abklingzeit am Kondensator τ_a ergibt sich somit zu

$$\tau_a = -RC \ln\left(\frac{1}{e}\right) = RC.$$



① Halbwertszeit einer radioaktiven Substanz



② Abklingzeit am Kondensator

Beispiel 4.38. Wie lautet die Umkehrfunktion von

$$f : \mathbb{R} \rightarrow \mathbb{R}_{>0} \quad \text{mit} \quad x \mapsto f(x) = 3e^{2x-1}?$$

Die Funktion f ist auf ihrem gesamten Definitionsbereich streng monoton wachsend, also existiert auf dem Wertebereich $\mathbb{R}_{>0}$ die Umkehrfunktion. Wir setzen

$$y = 3e^{2x-1}$$

und lösen nach x auf:

$$\frac{1}{3}y = e^{2x-1} \hookrightarrow \ln \frac{1}{3}y = 2x - 1 \hookrightarrow x = \frac{1}{2}(\ln \frac{1}{3}y + 1).$$

Nach Vertauschen der Variablen erhalten wir die Umkehrfunktion

$$g : \mathbb{R}_{>0} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto g(x) = \frac{1}{2}(\ln \frac{1}{3}x + 1). \quad \square$$

Allgemeine Potenz- und Exponentialfunktion. Mit der Exponential- und Logarithmusfunktion ist man in der Lage, die allgemeine Potenz- und Exponentialfunktion zu definieren.

Definition:

(1) *Die Funktion*

$$f : \mathbb{R}_{>0} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto f(x) = x^\alpha := e^{\alpha \ln x}$$

heißt **allgemeine Potenzfunktion**.

(2) *Die Funktion*

$$f : \mathbb{R} \rightarrow \mathbb{R}_{>0} \quad \text{mit} \quad x \mapsto f(x) = a^x := e^{x \ln a} \quad (a > 0)$$

heißt **allgemeine Exponentialfunktion**.

4.6 Trigonometrische Funktionen

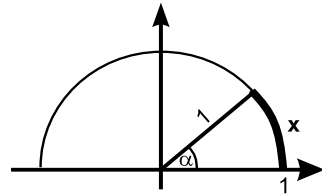
Zur Beschreibung periodischer Vorgänge benötigt man die Sinus- und Kosinusfunktionen. Schon um z.B. eine Wechselspannung $u(t) = u_0 \sin(\omega t + \varphi)$ anzugeben, verwendet man die allgemeine Sinusfunktion. Wir werden in diesem Abschnitt die trigonometrischen Funktionen zusammen mit ihren Umkehrfunktionen einführen, wichtige Eigenschaften diskutieren und durch Anwendungsbeispiele den Einsatz der Funktionen verdeutlichen.

4.6.1 Grundbegriffe

In den Naturwissenschaften und in der Technik spielen *periodische Vorgänge* eine wichtige Rolle. Sie sind dadurch gekennzeichnet, dass sich ein bestimmter Zustand regelmäßig wiederholt, z.B. bei akustischen und elektromagnetischen Schwingungen; Schwingungen einer Saite oder Feder; Umlaufbahnen von Satelliten. Periodische Funktionen von besonderer Bedeutung sind die **trigonometrischen Funktionen** bzw. *Winkelfunktionen*.

Zur Winkelmessung werden verschiedene Einheiten zugrunde gelegt:

- *Gradmaß* α (360° für den Vollkreis)
- *Bogenmaß* x (2π für den Vollkreis)
= Länge der Strecke auf dem Einheitskreis,
die der Winkel α herauschneidet.



Die im Bogenmaß gemessene Winkelgröße bezeichnen wir mit x . x ist positiv, falls der Winkel im Gegenuhrzeigersinn gemessen, und negativ, wenn der Winkel im Uhrzeigersinn gemessen wird. Die Einheit des Bogenmaßes heißt *Radian* (*rad*). Zwischen Winkelgröße α im Gradmaß und x im Bogenmaß besteht der Zusammenhang

$$\frac{\alpha}{360^\circ} = \frac{x}{2\pi}.$$

4.6.2 Sinus- und Kosinusfunktion

Beschreibt man einen Winkel im Bogenmaß, so können der Größe x Funktionswerte gemäß folgender Festlegung zugeschrieben werden:

Definition: Unter dem **Sinus** bzw. **Kosinus** eines Winkels x ($\sin x$ bzw. $\cos x$) versteht man die Ordinate bzw. Abszisse des Schnittpunktes des freien Schenkels des Winkels x mit dem Einheitskreis.

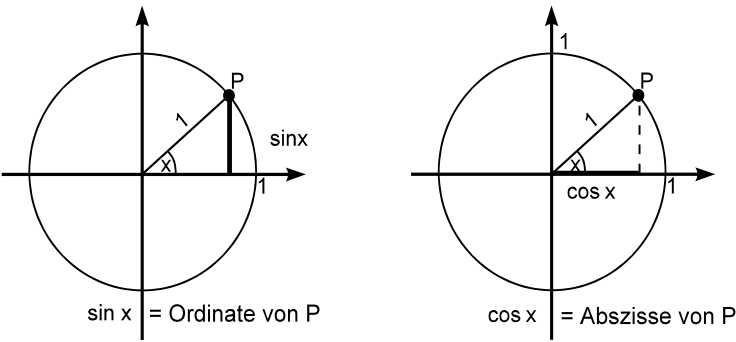


Abb. 4.19. Sinus und Kosinus im Einheitskreis

Mit dieser Definition im Einheitskreis erhalten wir die trigonometrischen Funktionen $\sin$ und $\cos$ als Funktionen auf $\mathbb{R}$:

$\sin : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \sin(x)$
 $\cos : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \cos(x).$

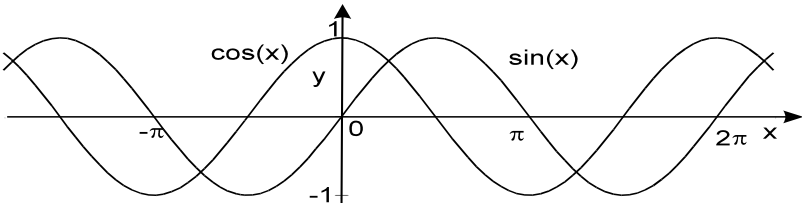


Abb. 4.20. Sinus- und Kosinusfunktion

 **Visualisierung mit MAPLE:** Auf der CD-Rom befindet sich eine Animation, in der gezeigt wird, wie durch Projektion des rotierenden Punktes P auf die x -Achse bzw. y -Achse jeweils die Sinus- bzw. Kosinusfunktion entstehen.

Eigenschaften der Sinus- und Kosinusfunktion:

	$f(x) = \sin x$	$f(x) = \cos x$
Definitionsbereich	$\mathbb{R}$	$\mathbb{R}$
Wertebereich	$[-1, 1]$	$[-1, 1]$
Periode	2π	2π
Symmetrie	ungerade	gerade
Nullstellen	$x_n = n \cdot \pi$	$x_n = \frac{\pi}{2} + n \pi$
relative Maxima	$x_k = \frac{\pi}{2} + k \cdot 2\pi$	$x_k = k \cdot 2\pi$
relative Minima	$x_k = \frac{3}{2}\pi + k \cdot 2\pi$	$x_k = \pi + k \cdot 2\pi$

$n, k \in \mathbb{Z}$

Die allgemeine Sinusfunktion. In den Anwendungen kommen Sinus und Kosinus nicht nur mit dem Argument x , sondern in der allgemeineren Form

$$y(x) = a \sin(bx + c) + d \quad \text{bzw.} \quad y(x) = a \cos(bx + c) + d$$

vor. Im Folgenden diskutieren wir die Bedeutung jedes der Parameter einzeln:

(1) **Bedeutung von a :** Der Faktor a in

$$y(x) = a \cdot \sin x$$

gibt die **maximale Amplitude** der Funktion an. Der Wertebereich dieser Funktion ist $\mathbb{W} = [-a, a]$.

Beispiel 4.39. $y(x) = 2 \cdot \sin(x) \Rightarrow$ Amplitude 2.

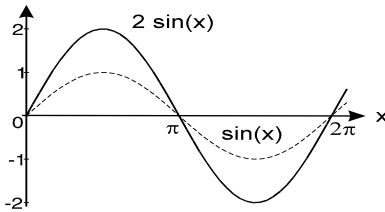


Abb. 4.21. Amplitude a

(2) **Bedeutung von b :** Der Faktor b in

$$y(x) = \sin(bx)$$

bewirkt eine Veränderung der Periode gegenüber der reinen Sinusfunktion. Die Periode p von $\sin(bx)$ erhält man, wenn das Argument des Sinus die dritte Nullstelle liefert, also für

$$bp = 2\pi \Rightarrow \boxed{p = \frac{2\pi}{b}} \text{ Periode.}$$

⚠ Damit verkleinert sich die Periode für $b > 1$ und vergrößert sich für $b < 1$.

Beispiel 4.40. $y(x) = \sin(2x) \Rightarrow$ Periode: $p = \frac{2\pi}{2} = \pi$.

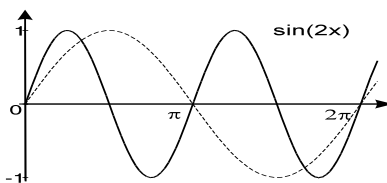


Abb. 4.22. Periode $p = \frac{2\pi}{b}$

(3) **Bedeutung der Konstanten c :** Die Konstante c in

$$y(x) = \sin(x + c)$$

bewirkt eine **Verschiebung** der reinen Sinusfunktion **entlang der x -Achse**. Man bezeichnet c auch als *Phase* oder *Nullphase*. Die erste Nullstelle von $\sin(x + c)$ findet man, wenn das Argument $x + c$ Null wird:

$$\sin(x + c) = 0 \quad \Rightarrow \quad x_0 + c = 0 \quad \Rightarrow \quad x_0 = -c.$$

△ Für $c > 0$ wird die Kurve um x_0 nach links verschoben.

△ Für $c < 0$ wird die Kurve um x_0 nach rechts verschoben.

Beispiel 4.41. $y = \sin(x + \pi/2) \Rightarrow$ Kurve um $\pi/2$ nach links verschoben.

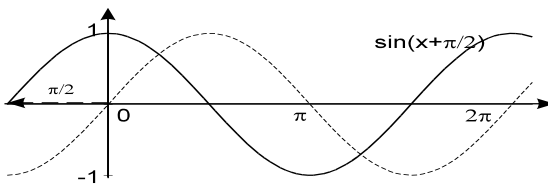


Abb. 4.23. Verschiebung entlang der x -Achse

(4) **Bedeutung der Konstanten d :** Die Konstante d in

$$y(x) = \sin(x) + d$$

bewirkt eine **Verschiebung** der reinen Sinusfunktion **entlang der y -Achse** um den Wert d .

Beispiel 4.42. $y(x) = \sin(x) + 2 \Rightarrow$ Verschiebung um 2 in y -Richtung:

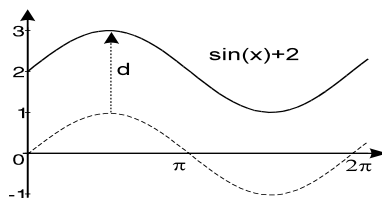


Abb. 4.24. Verschiebung entlang der y -Achse

Damit haben wir die Parameter, die in der allgemeinen Sinusfunktion auftreten, separat behandelt und deren Bedeutung beschrieben. Abschließend fehlt noch die Diskussion, wenn alle Parameter gleichzeitig auftreten:

(5) Zusammenfassende Diskussion der allgemeinen Sinusfunktion:

$$y(x) = a \sin(bx + c) = a \sin\left(b\left(x + \frac{c}{b}\right)\right)$$

Aus der Funktionsgleichung der allgemeinen Sinusfunktion liest man direkt die Amplitude a ab. Die Periode p berechnet sich über b durch $p = \frac{2\pi}{b}$. Die Nullphase c liest man an der ersten Darstellung ab, während die Phasenverschiebung $-\frac{c}{b}$ aus der zweiten Darstellung ersichtlich ist. Zusammenfassend erhalten wir die folgende Tabelle

Periode:	$p = \frac{2\pi}{b}$
1. Nullstelle:	$x_0 = -\frac{c}{b}$
Wertebereich:	$-a \leq y(x) \leq a$
(Null-)Phase:	c
Verschiebung:	$-\frac{c}{b}$

zusammen mit dem Graphen der allgemeinen Sinusfunktion:

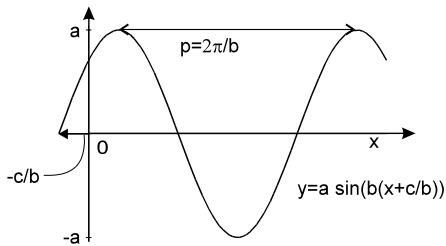
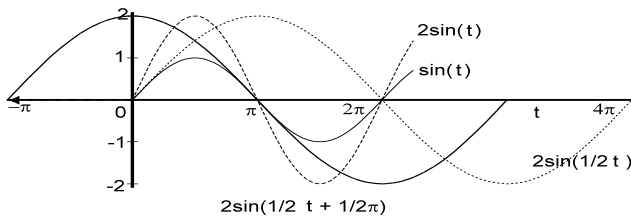


Abb. 4.25. Allgemeine Sinusfunktion

Beispiel 4.43. $u(t) = 2 \sin\left(\frac{1}{2}t + \frac{1}{2}\pi\right)$. Ausgehend von der reinen Sinusfunktion $\sin(t)$ gehen wir zur doppelten Amplitude über: $2 \sin(t)$. Anschließend modifizieren wir die Periode zu $p = \frac{2\pi}{\frac{1}{2}} = 4\pi$ und erhalten $2 \sin\left(\frac{1}{2}t\right)$. Danach berücksichtigen wir die Phasenverschiebung von $-\pi$, da $2 \sin\left(\frac{1}{2}t + \frac{1}{2}\pi\right) = 2 \sin\frac{1}{2}(t + \pi)$.



➤ 4.6.3 Tangens- und Kotangensfunktion

Ausgehend von der Sinus- und Kosinusfunktion können wir analog der geometrischen Interpretation die *Tangens-* und *Kotangensfunktion* als Quotient von Sinus und Kosinus bzw. von Kosinus und Sinus definieren. Dabei ist allerdings zu beachten, dass die Nullstellen des jeweiligen Nenners aus dem Definitionsbereich auszuschließen sind.

Definition:

$\tan : \mathbb{R} \setminus \{\frac{\pi}{2} + k \cdot \pi, \quad k \in \mathbb{Z}\} \rightarrow \mathbb{R}$ mit $x \mapsto \tan x := \frac{\sin x}{\cos x}$
heißt **Tangensfunktion**.

$\cot : \mathbb{R} \setminus \{k \cdot \pi, \quad k \in \mathbb{Z}\} \rightarrow \mathbb{R}$ mit $x \mapsto \cot x := \frac{\cos x}{\sin x}$
heißt **Kotangensfunktion**.

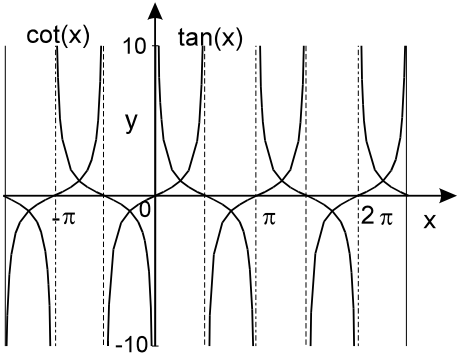


Abb. 4.26. Tangens- und Kotangensfunktion

Eigenschaften der Tangens- und Kotangensfunktion:

	$f(x) = \tan x$	$f(x) = \cot x$
Definitionsbereich	$\mathbb{R} \setminus \{\frac{\pi}{2} + k \cdot \pi, \quad k \in \mathbb{Z}\}$	$\mathbb{R} \setminus \{k \cdot \pi, \quad k \in \mathbb{Z}\}$
Wertebereich	$\mathbb{R}$	$\mathbb{R}$
Periode	π	π
Symmetrie	ungerade	ungerade
Nullstellen	$x_n = n \cdot \pi$	$x_n = \frac{\pi}{2} + n \cdot \pi$
Pole	$x_k = \frac{\pi}{2} + k \cdot \pi$	$x_k = k \cdot \pi$
Asymptoten	$x = \frac{\pi}{2} + k \cdot \pi$	$x = k \cdot \pi$

$k, n \in \mathbb{Z}$

► 4.6.4 Zusammenstellung wichtiger trigonometrischer Formeln

Für die trigonometrischen Funktionen gelten für alle $x, x_1, x_2 \in \mathbb{R}$ die folgenden wichtigen Beziehungen:

(1) Symmetrieverhalten

$\sin(-x) = -\sin(x)$	$\cos(-x) = \cos(x)$
$\tan(-x) = -\tan(x)$	$\cot(-x) = -\cot(x)$

(2) Verschiebungsidentitäten

$\sin x = \cos\left(x - \frac{\pi}{2}\right)$	$\cos x = \sin\left(x + \frac{\pi}{2}\right)$
---	---

(3) Nach dem Satz von Pythagoras

$$\sin^2 x + \cos^2 x = 1$$

(4) Es gelten die Additionstheoreme

$$\sin(x_1 \pm x_2) = \sin x_1 \cos x_2 \pm \cos x_1 \sin x_2$$

$$\cos(x_1 \pm x_2) = \cos x_1 \cos x_2 \mp \sin x_1 \sin x_2$$

$$\tan(x_1 \pm x_2) = \frac{\tan x_1 \pm \tan x_2}{1 \mp \tan x_1 \tan x_2}$$

(5) Aus den Additionstheoremen ergeben sich weitere, oft verwendete Formeln

$\sin(2x) = 2 \cdot \sin x \cos x$	$\sin(3x) = 3 \sin(x) - 4 \sin^3(x)$
$\cos(2x) = \cos^2(x) - \sin^2(x)$	$\cos(3x) = -3 \cos(x) + 4 \cos^3(x)$

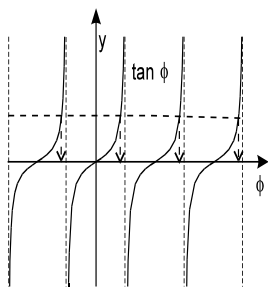
(6) und

$\sin^2 x = \frac{1}{2} (1 - \cos(2x)) = \frac{1}{1 + \tan^2 x}$
$\cos^2 x = \frac{1}{2} (1 + \cos(2x)) = \frac{\tan^2 x}{1 + \tan^2 x}$

Grundlegend sind die Formeln (4). Alle anderen Beziehungen können auf diese Additionstheoreme zurückgespielt werden.

4.7 Arkusfunktionen

Die trigonometrischen Funktionen ordnen jedem x genau einen Funktionswert y zu. Oftmals stellt sich aber das umgekehrte Problem: Gegeben ist der Funktionswert einer trigonometrischen Funktion und gesucht ist das zugehörige Argument.



Beispiel 4.44. In einem RL-Kreis ist die Phase φ zwischen der angelegten Eingangsspannung

$$U_E(t) = U_0 \sin(\omega t)$$

und der am Ohmschen Widerstand abgegriffenen Ausgangsspannung

$$U_A(t) = \frac{U_0}{\sqrt{R^2 + (\omega L)^2}} \sin(\omega t + \varphi)$$

bestimmt durch

$$\tan \varphi = \frac{\omega L}{R}.$$

Gesucht ist die Phase φ . Die Umkehrung der Tangensfunktion ist nicht eindeutig, da der Tangens wie alle anderen trigonometrischen Funktionen periodisch ist. Man schränkt daher den Definitionsbereich so ein, dass auf dem eingeschränkten Definitionsbereich eine monotone Funktion entsteht. Die Umkehrfunktionen der trigonometrischen Funktionen sind die **Arkusfunktionen**.

⚠ Merke: Grundsätzlich lassen sich die trigonometrischen Funktionen in $\mathbb{R}$ nicht umkehren. Schränkt man den Definitionsbereich jedoch auf ein Intervall ein, in dem die Funktionen sich streng monoton verhalten, so sind auf diesem eingeschränkten Intervall die trigonometrischen Funktionen umkehrbar. Die unten gewählten Einschränkungen des Definitionsbereichs liefern jeweils den sog. *Hauptwert* der Funktionen.

Bemerkung: Da die Arkusfunktionen die Umkehrfunktionen der trigonometrischen Funktionen sind, erhält man den Graphen dieser Arkusfunktionen durch **Spiegelung der jeweiligen trigonometrischen Funktion an der Winkelhalbierenden**. Definitionsbereich und Wertebereich der Arkusfunktionen ergeben sich aus dem Wertebereich und Definitionsbereich der zugehörigen trigonometrischen Funktionen.

4.7.1 Arkussinusfunktion

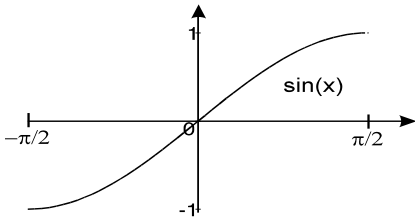
Die eingeschränkte Sinusfunktion

$$\sin : \left[-\frac{\pi}{2}, \frac{\pi}{2}\right] \rightarrow [-1, 1] \quad \text{mit} \quad x \mapsto \sin x$$

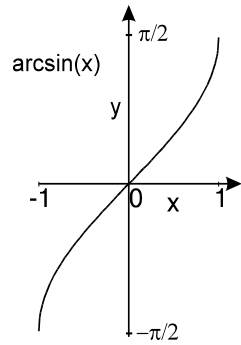
ist im Intervall $\left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$ streng monoton wachsend. Folglich existiert die Umkehrfunktion **Arkussinus**

$$\arcsin : [-1, 1] \rightarrow \left[-\frac{\pi}{2}, \frac{\pi}{2}\right] \quad \text{mit} \quad x \mapsto \arcsin(x).$$

Man beachte, dass damit $\arcsin(\sin(x)) = x$ für alle $x \in \left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$ und $\sin(\arcsin(y)) = y$ für alle $y \in [-1, 1]$.



Sinusfunktion



Arkussinusfunktion

Eigenschaften der Arkussinusfunktion:

	$\sin(x)$	$\arcsin(x)$
Definitionsbereich	$\left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$	$[-1, 1]$
Wertebereich	$[-1, 1]$	$\left[-\frac{\pi}{2}, \frac{\pi}{2}\right]$
Nullstelle	$x_0 = 0$	$x_0 = 0$
Symmetrie	ungerade	ungerade
Monotonie	streng monoton wachsend	streng monoton wachsend

Beispiele 4.45:

① $\arcsin 0 = 0, \quad \arcsin\left(\frac{1}{2}\right) = \frac{\pi}{6}, \quad \arcsin\left(\frac{1}{2}\sqrt{2}\right) = \frac{\pi}{4}.$

② Gesucht sind alle Lösungen der Gleichung $\sin x = 0,5$.

Wegen $\sin x = 0,5$ folgt $x = \arcsin(0,5) = \frac{\pi}{6}$. Wegen der Periodizität von $\sin x = \sin(x + 2\pi)$ und $\sin(\pi - x) = \sin x$ folgt insgesamt

$$\mathbb{L} = \{x \in \mathbb{R} : x = \frac{\pi}{6} + k \cdot 2\pi \text{ oder } x = \frac{5}{6}\pi + k \cdot 2\pi \quad \text{mit } k \in \mathbb{Z}\}.$$

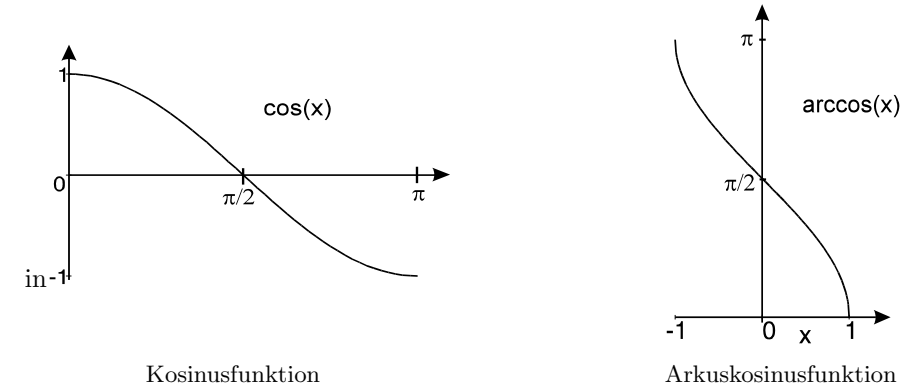
➤ **4.7.2 Arkuskosinusfunktion**

Die eingeschränkte Kosinusfunktion

$$\cos : [0, \pi] \rightarrow [-1, 1] \quad \text{mit} \quad x \mapsto \cos x$$

ist im Intervall $[0, \pi]$ streng monoton fallend. Daher existiert die Umkehrfunktion **Arkuskosinus**

$$\arccos : [-1, 1] \rightarrow [0, \pi] \quad \text{mit} \quad x \mapsto \arccos(x).$$



Eigenschaften der Arkuskosinusfunktion:

	$\cos(x)$	$\arccos(x)$
Definitionsbereich	$[0, \pi]$	$[-1, 1]$
Wertebereich	$[-1, 1]$	$[0, \pi]$
Nullstelle	$\frac{\pi}{2}$	1
Monotonie	streng monoton fallend	streng monoton fallend

Beispiel 4.46. $\arccos \frac{\pi}{2} = 0$, $\arccos \frac{1}{2} = \frac{\pi}{3}$, $\arccos(-0,237) = 1,8101$.

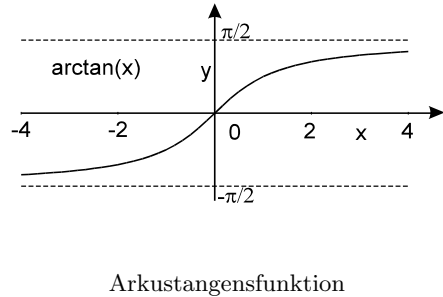
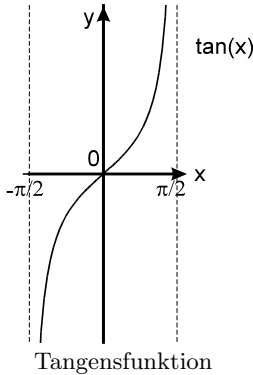
➤ **4.7.3 Arkustangensfunktion**

Die eingeschränkte Tangensfunktion

$$\tan : \left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \tan x$$

ist im Intervall $\left(-\frac{\pi}{2}, \frac{\pi}{2}\right)$ streng monoton wachsend und hat als Wertebereich ganz $\mathbb{R}$. Folglich existiert die Umkehrfunktion **Arkustangens**

$$\arctan : \mathbb{R} \rightarrow \left(-\frac{\pi}{2}, \frac{\pi}{2}\right) \quad \text{mit} \quad x \mapsto \arctan(x).$$



Durch eine Spiegelung des Graphen der Tangensfunktion an der Winkelhalbierenden erhält man den Graphen der Arkustangensfunktion. Damit werden die Polstellen bei $x = \pm \frac{\pi}{2}$ von Tangens zu Asymptoten bei $y = -\frac{\pi}{2}$ für $x \rightarrow -\infty$ und $y = \frac{\pi}{2}$ für $x \rightarrow \infty$ des Arkustangens.

Eigenschaften der Arkustangensfunktion:

	$\tan(x)$	$\arctan(x)$
Definitionsbereich	$(-\frac{\pi}{2}, \frac{\pi}{2})$	$\mathbb{R}$
Wertebereich	$\mathbb{R}$	$(-\frac{\pi}{2}, \frac{\pi}{2})$
Nullstelle	$x_0 = 0$	$x_0 = 0$
Symmetrie	ungerade	ungerade
Monotonie	streng monoton wachsend	streng monoton wachsend
Asymptoten	$x = \pm \frac{\pi}{2}$	$y = \pm \frac{\pi}{2}$

Beispiele 4.47:

- ① $\arctan 1 = \frac{\pi}{4}$, $\arctan(-3\pi) = -1,4651$.
- ② Gesucht sind alle Lösungen von $\tan x = \sqrt{3}$. Mit $x = \arctan \sqrt{3} = \frac{\pi}{3}$ folgt aufgrund der Periodizität des Tangens: $\mathbb{L} = \{x \in \mathbb{R} : x = \frac{\pi}{3} + k \cdot \pi, \quad k \in \mathbb{Z}\}$.

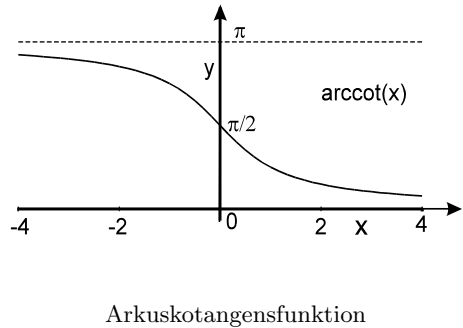
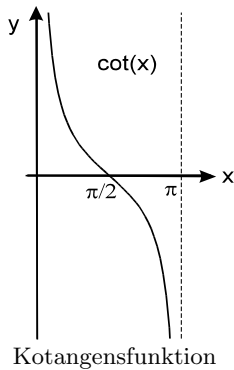
➤ 4.7.4 Arkuskotangensfunktion

Die eingeschränkte Kotangensfunktion

$$\cot : (0, \pi) \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto \cot(x)$$

ist im Intervall $(0, \pi)$ streng monoton fallend und hat als Wertebereich ganz $\mathbb{R}$. Folglich existiert die Umkehrfunktion **Arkuskotangens**

$$\operatorname{arccot} : \mathbb{R} \rightarrow (0, \pi) \quad \text{mit} \quad x \mapsto \operatorname{arccot}(x).$$



Eigenschaften der Arkuskotangensfunktion:

	$\cot(x)$	$\operatorname{arccot}(x)$
Definitionsbereich	$(0, \pi)$	$\mathbb{R}$
Wertebereich	$\mathbb{R}$	$(0, \pi)$
Nullstelle	$x_0 = \frac{\pi}{2}$	–
Monotonie	streng monoton fallend	streng monoton fallend
Asymptoten	$x = 0$ $x = \pi$	$y = 0$ $y = \pi$

Beispiel 4.48. $\operatorname{arccot}(0) = \frac{\pi}{2}$, $\operatorname{arccot} 1 = \frac{\pi}{4}$.

Bemerkungen:

- (1) Auf Taschenrechnern kommen die elementaren Funktionen $\arcsin$, $\arccos$, $\arctan$, arccot nicht als Tasten vor, sondern man verwendet Symbole wie z.B. **INV** und **SIN** für den $\arcsin$ bzw. entsprechende Tastenkombinationen für die anderen inversen trigonometrischen Funktionen.
- (2) Der Arkuskotangens spielt in der Praxis keine Rolle. Er fehlt daher auch auf den meisten Taschenrechnern. Wegen der Beziehung

$$\operatorname{arccot}(x) = \frac{\pi}{2} - \arctan(x)$$

kann man ihn aber dennoch berechnen.

- (3) Denn es zeigt sich, dass ähnliche Formeln wie die für den arccot auch für die anderen trigonometrischen Umkehrfunktionen existieren. Es gilt

$$\begin{aligned}\arccos(x) &= \frac{\pi}{2} - \arcsin(x) \\ \arcsin(x) &= \arctan \frac{x}{\sqrt{1-x^2}} \quad \text{für } -1 < x < 1\end{aligned}$$

- (4) Zu den Funktionen $\sin, \cos, \tan$ und $\cot$ gibt es unendlich viele Intervalle, in denen die betreffenden Funktionen streng monoton wachsend sind. Aus diesem Grund lassen sich beliebig viele Umkehrfunktionen angeben. Die oben genannten stellen jeweils den sog. *Hauptwert* dar. Die Umkehrfunktionen der trigonometrischen Funktionen werden auch *zyklometrische Funktionen* genannt, weil sie für die Kreisberechnung von Bedeutung sind.
- (5) Es gelten die folgenden Beziehungen zwischen den Areafunktionen:

$\arcsin(x)$	$=$	$-\arcsin(-x)$	$=$	$\frac{\pi}{2} - \arccos(x)$	$=$	$\arctan \frac{x}{\sqrt{1-x^2}}$
$\arccos(x)$	$=$	$\pi - \arccos(-x)$	$=$	$\frac{\pi}{2} - \arcsin(x)$	$=$	$\operatorname{arccot} \frac{x}{\sqrt{1-x^2}}$
$\arctan(x)$	$=$	$-\arctan(-x)$	$=$	$\frac{\pi}{2} - \operatorname{arccot}(x)$	$=$	$\arcsin \frac{x}{\sqrt{1+x^2}}$
$\operatorname{arccot}(x)$	$=$	$\pi - \operatorname{arccot}(-x)$	$=$	$\frac{\pi}{2} - \arctan(x)$	$=$	$\arccos \frac{x}{\sqrt{1+x^2}}$

MAPLE-Worksheets zu Kapitel 4



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 4 mit MAPLE zur Verfügung.

- [Grundbegriffe und allgemeine Funktionseigenschaften](#)
- [Polynome](#)
- [Rationale Funktionen](#)
- [Potenz- und Wurzelfunktionen](#)
- [Exponential- und Logarithmusfunktion](#)
- [Hyperbolische Funktionen](#)
- [Trigonometrische Funktionen](#)
- [MAPLE-Lösungen zu den Aufgaben](#)

4.8 Aufgaben zu elementaren Funktionen

- 4.1 Bestimmen Sie den größtmöglichen Definitionsbereich sowie den Wertebereich der folgenden Funktionen
- a) $f(x) = \sqrt{x^2 - 1}$ b) $y = \ln|x|$ c) $f(x) = \frac{x^2}{4x^2 - 16}$
d) $f(x) = \frac{x-1}{x+1}$ e) $y = e^{|x|}$ f) $f(x) = \frac{x}{x^2 + 1}$
- 4.2 Bestimmen Sie das Symmetrieverhalten und den maximalen Definitionsbereich
- a) $f(x) = 4x^2 - 16$ b) $f(x) = \frac{x^3}{x^2 + 1}$ c) $f(x) = \sin x \cdot \cos x$
d) $f(x) = |x^2 - 16|$ e) $f(x) = \frac{x^2 - 1}{1 + x^2}$ f) $f(x) = \frac{1}{x - 1}$
- 4.3 Man untersuche auf Monotonie
- a) $y = x^4$ b) $y = \sqrt{x - 1}$ für $x \geq 1$ c) $y = x^3 + 2x$ d) $y = e^{2x}$
- 4.4 Wie lautet die Umkehrfunktion von
- a) $f: \mathbb{R}_{>0} \rightarrow ?$
 $x \mapsto y = \frac{1}{(2x)}$ b) $f: \mathbb{R}_{\geq 0} \rightarrow ?$
 $x \mapsto y = \sqrt{3x}$
c) $f: \mathbb{R} \rightarrow ?$ d) $f: \mathbb{R}_{>-1} \rightarrow ?$
 $x \mapsto y = 2e^{x-\frac{1}{2}}$ $x \mapsto y = \frac{x-1}{x+1}$
- 4.5 Bestimmen Sie die Polynomfunktion kleinsten Grades, die durch die folgenden Punkte geht: $(-3, 11)$; $(-1, 7)$; $(0, 5)$; $(4, -3)$.
- 4.6 Bestimmen Sie die Nullstellen folgender Funktionen:
- a) $f(x) = x^3 + 2x^2 - 13x + 10$
b) $f(x) = x^3 - x^2 + 2$
c) $f(x) = x^4 - 2x^3 - 25x^2 + 50x$
- 4.7 Man berechne mit dem Horner-Schema den Funktionswert der Funktion $f(x)$ an der Stelle x_0 für
- a) $f(x) = x^3 - 2x^2 - 3x + 1$; $x_0 = 2$
b) $f(x) = 0.1x^4 + x^3 + 2x^2 - 4$; $x_0 = 3$.
- 4.8 Gibt es Polynome, die keine Nullstellen besitzen?
- 4.9 Man berechne mit dem Newton-Schema die Koeffizienten der ganzrationalen Funktion vom Grade ≤ 3 , welche durch die Wertepaare $(0, 1)$; $(1, 0)$; $(2, 5)$; $(-1, 2)$ geht.
- 4.10 Bestimmen Sie die Nullstellen der Funktion
- a) $f(x) = 3x^3 + 3x^2 - 3x - 3$
b) $f(x) = x^4 - 13x^2 + 36$
- 4.11 Welches Polynom kleinsten Grades geht durch die Wertepaare $(-1, 0)$; $(0, 1)$; $(1, 2)$; $(2, 6)$?
- 4.12 Faktorisieren Sie mit MAPLE das Polynom
 $2x^6 + 3x^5 - 63x^4 - 55x^3 + 657x^2 + 216x - 2160$
- 4.13 Bestimmen Sie mit MAPLE alle Nullstellen von
- a) $7x^4 - 59x^3 + 19x^2 + 166x - 1008$
b) $x^3 - x^2 - 100x + 310$

4.14 Werten Sie mit MAPLE die Polynome aus Aufgabe 4.13 an der Stelle $x_0 = 4$ aus, indem Sie entweder mit **unapply** den jeweiligen Ausdruck in eine Funktion konvertieren oder indem Sie mit dem **subs**-Befehl diese Stelle in das Polynom einsetzen.

4.15 Zeichnen Sie mit MAPLE die Funktion $x^3 - x^2 - 100x + 310$ und lesen Sie aus dem Schaubild die Extremwerte ab.

4.16 Faktorisieren Sie mit MAPLE das Polynom soweit möglich. Erstellen Sie das Horner-Schema zu diesem Polynom. Bestimmen Sie den Grad des Polynoms.
 $-17x^6 + 11x^4 - 20x^3 + 13x^2 - 3x + 56x^7 + 4x^5 - 15x^8 + 35x^9$

4.17 Wo besitzen die folgenden Funktionen Nullstellen, wo Pole?

$$\begin{array}{ll} \text{a) } y = \frac{x^2 + x - 2}{x - 2} & \text{b) } y = \frac{x^3 - 5x^2 - 2x + 24}{x^3 + 3x^2 + 2x} \\ \text{c) } y = \frac{x^2 - 2x + 1}{x^2 - 1} & \text{d) } y = \frac{(x - 1)}{(x - 1)^2 (x + 1)} \end{array}$$

4.18 Bestimmen Sie für die folgenden gebrochenrationalen Funktionen: Nullstellen, Pole, Asymptoten im Unendlichen. Zeichnen Sie mit MAPLE den Funktionsgraphen und die Asymptoten.

$$\text{a) } y = \frac{x^2 - 4}{x^2 + 1} \quad \text{b) } y = \frac{x^3 - 6x^2 + 12x - 8}{x^2 - 4} \quad \text{c) } y = \frac{x^3 - 5x^2 + 8x - 4}{x^3 - 6x^2 + 12x - 8} \quad \text{d) } y = \frac{(x - 1)^2}{(x + 1)^2}$$

4.19 Bestimmen Sie mit MAPLE die Null- und Polstellen der Funktion

$$h(x) = \frac{x^3 - 6x^2 - 12x + 49}{(x - 2)(x - 7)},$$

indem Sie Zähler und Nenner in Linearfaktoren zerlegen. Bestimmen Sie die Asymptoten und zeichnen Sie die Funktion zusammen mit ihren Asymptoten in ein Schaubild. Zeichnen Sie die Funktion in der Nähe der Null- und Polstellen.

4.20 Gegeben ist die rationale Funktion $\frac{x^4 + x^3 - 4x^2 - 4x}{x^4 + x^3 - x^2 - x}$. Formen Sie diese Funktion mit MAPLE in die folgenden Ausdrücke um:

$$\begin{array}{lll} \text{a) } \frac{(x+2)(x+1)(x-2)}{x^3 + x^2 - x - 1} & \text{b) } \frac{x^4 + x^3 - 4x^2 - 4x}{x(x-1)(x+1)^2} & \text{c) } \frac{(x+2)(x-2)}{(x-1)(x+1)} \\ \text{d) } \frac{x^2}{(x-1)(x+1)} - 4 \frac{1}{(x-1)(x+1)} \end{array}$$

4.21 Wird ein Kondensator mit Kapazität C über einen Ohmschen Widerstand R entladen, so nimmt seine Ladung Q exponentiell mit der Zeit ab:

$$Q = Q_0 e^{-\frac{1}{RC}t}.$$

Zu welchem Zeitpunkt sinkt die Ladung unter 10 % ihres Anfangswertes Q_0 ?

4.22 Stromkreis mit Induktivität L und Ohmschen Widerstand R . Beim Einschalten einer Gleichspannungsquelle erreicht der Strom infolge der Selbstinduktion erst nach einiger Zeit den nach dem Ohmschen Gesetz erwarteten Endwert i_0 . Es gilt

$$i(t) = i_0 \left(1 - e^{-\frac{R}{L}t}\right).$$

Berechnen Sie für $i_0 = 4 \text{ A}$, $R = 5 \Omega$, $L = 2.5 \text{ H}$ den Zeitpunkt, bei dem die Stromstärke 95 % des Endzustandes erreicht hat. Skizzieren Sie die Strom-Zeit-Funktion.

4.23 Bestimmen Sie die Parameter a und b der Funktion $y = a e^{-b x} + 2$ so, dass die Punkte $A = (0, 10)$ und $B = (5, 3)$ auf der Kurve liegen.

4.24 Man löse die folgenden Exponentialgleichungen

a) $e^{x^2-2x} = 2$ b) $e^x + 2e^{-x} = 3$

4.25 Welche Lösung besitzt die logarithmische Gleichung

$$\ln \sqrt{x} + 1.5 \cdot \ln(x) = \ln(2x)?$$

4.26 Bei einer gedämpften Schwingung $x(t) = A e^{-\gamma t} \cdot \sin(\omega t + \varphi)$ kann durch Messung der Amplituden zweier aufeinander folgenden Schwingungen die Dämpfung γ bestimmt werden. Ist $T = \frac{1}{100\text{ s}}$ die Periodendauer der gedämpften Schwingung, $x(t_0) = 200$ und $x(t_0 + T) = 100$ die Amplituden zweier aufeinander folgender Schwingungen, so berechne man γ .

4.27 Rechnen Sie vom Grad- ins Bogenmaß bzw. vom Bogen- ins Gradmaß um :

Grad	40,36°	278,19°
Bogen	1.4171	-5.6213

4.28 Man leite aus dem Additionstheorem der Kosinusfunktion die Formel

$$\boxed{\sin^2 x + \cos^2 x = 1} \quad \text{ab.}$$

4.29 Zeichnen Sie den Funktionsverlauf von

$$f(x) = 2 \cdot \cos(2x - \pi).$$

4.30 Man bestimme für die folgenden Funktionen Amplitude A , Periode p , Phasenverschiebung:

a) $y = 2 \cdot \sin(3x - \frac{\pi}{6})$ b) $y = 5 \cdot \cos(2x + 4.2)$
 c) $y = 10 \cdot \sin(\pi x - 3\pi)$ d) $y = 2.4 \cdot \cos(4x - \frac{\pi}{2})$

4.31 Skizzieren Sie den Funktionsverlauf der harmonischen Schwingung:

$$f(t) = 2 \cdot \sin(2t - 4)$$

4.32 Berechnen Sie

a) $\arcsin(1)$ b) $\arcsin(\frac{1}{2}\sqrt{2})$ c) $\arcsin(-\frac{1}{2}\sqrt{3})$ d) $\arcsin(0.481)$
 e) $\arccos(\frac{1}{2})$ f) $\arccos(\frac{1}{2}\sqrt{3})$ g) $\arccos(-1)$ h) $\arccos(0.8531)$
 i) $\arctan(1)$ k) $\arctan(-\sqrt{3})$ l) $\text{arccot}(-\frac{1}{\sqrt{3}})$ m) $\text{arccot}(\frac{1}{3}\sqrt{3})$

4.33 Welchen Wert hat die Größe x ?

a) $\arcsin x = \pi/4$ b) $\arctan x = 0.7749$ c) $\arccos x = 1.021$
 d) $\text{arccot} x = 2.9208$ e) $(\arccos x)^2 = 0.25$

4.34 Man beweise $\sin(\arccos(x)) = \sqrt{1-x^2}$. (Anleitung: Man setze $y = \arccos(x)$.)

4.35 Zeigen Sie, dass für positive a innerhalb des Definitionsbereichs gilt

a) $\arcsin(a) = \arccos \sqrt{1-a^2}$ b) $\arccos(a) = \arcsin \sqrt{1-a^2}$
 c) $\text{arccot}(a) = \arctan(\frac{1}{a})$ d) $\arcsin(a) = \arctan \frac{a}{\sqrt{1-a^2}}$

4.36 Vereinfachen Sie a) $\sin(\arcsin(x))$ b) $\cos(\arccos(x))$ c) $\sin(\arccos(x))$
 d) $\cos(\arcsin(x))$ e) $\sin(\arctan(x))$ f) $\tan(\arccos(x))$

4.37 Geben Sie den Definitionsbereich der folgenden Funktionen an

a) $y = x + \arccos(x)$ b) $y = \sqrt{x} + \arcsin(x)$ c) $y = \frac{\pi}{2} + \arcsin(x-1)$

Zeichnen Sie die Funktionen und bestimmen Sie den Wertebereich.

Kapitel 5
Komplexe Zahlen

5

5	Komplexe Zahlen	191
5.1	Darstellung komplexer Zahlen	194
5.1.1	Algebraische Normalform	194
5.1.2	Trigonometrische Normalform	195
5.1.3	Exponentielle Normalform	196
5.1.4	Umformungen der Normalformen	197
5.2	Komplexe Rechenoperationen	200
5.2.1	Addition	200
5.2.2	Subtraktion	200
5.2.3	Multiplikation	201
5.2.4	Division	203
5.2.5	Potenz	205
5.2.6	Wurzeln	206
5.2.7	Fundamentalsatz der Algebra	207
5.3	Anwendungen	209
5.3.1	Beschreibung harmonischer Schwingungen im Komplexen	209
5.3.2	Superposition gleichfrequenter Schwingungen	210
5.3.3	Beschreibung von RCL-Gliedern bei Wechselströmen	214
5.3.4	Beispiele für RCL-Wechselstromschaltungen	216
5.4	Aufgaben zu komplexen Zahlen	219
 Zusätzliche Abschnitte auf der CD-Rom		
5.5	Komplexe Zahlen mit MAPLE	cd
5.5.1	Darstellung komplexer Zahlen mit MAPLE	cd
5.5.2	Komplexes Rechnen mit MAPLE	cd
5.6	Übertragungsfunktion für RCL-Filterschaltungen	cd
5.6.1	Übertragungsfunktion für lineare Ketten	cd
5.6.2	Dimensionierung von Hoch- und Tiefpässen	cd

5 Komplexe Zahlen

Die komplexen Zahlen stellen bei der Beschreibung von elektrischen Wechselstromschaltungen ein unverzichtbares Hilfsmittel dar. Fast jedes Lehrbuch über die Beschreibung von elektrischen Schaltkreisen hat als einleitendes Kapitel eine Einführung in die komplexen Zahlen. Einer der Gründe liegt darin, dass einfache Regeln von Gleichstrom-Netzwerken sich auf Wechselstrom-Schaltungen übertragen, wenn man komplexe Widerstände einführt.

Hinweis: Auf der CD-Rom befindet sich ein zusätzlicher Abschnitt über die Anwendung der komplexen Zahlen bei der Beschreibung von [RCL-Filterschaltungen](#).

Zunächst behandeln wir die Grundlagen der komplexen Zahlen innerhalb der Mathematik und beginnen mit einer mathematischen Problemstellung: Wie wir im Abschnitt 4.2 über Polynome bereits festgestellt haben, besitzt jedes Polynom vom Grade n in $\mathbb{R}$ höchstens n verschiedene Nullstellen. Aber schon beim quadratischen Polynom $p(x) = x^2 + 1$ zeigt sich, dass dieses Polynom in $\mathbb{R}$ keine Nullstellen besitzt. Löst man die Gleichung $x^2 + 1 = 0$ formal nach x auf, so erhält man

$$x_{1/2} = \pm\sqrt{-1} \notin \mathbb{R}.$$

Es hat sich als außerordentlich erfolgreich erwiesen, den Zahlenbereich der reellen Zahlen zu erweitern, indem man $\sqrt{-1}$ als eine neue Einheit einführt:

$$i := \sqrt{-1} \quad (\text{imaginäre Einheit}).$$

Die Bezeichnung imaginäre Einheit rührt daher, dass sich die Wurzel jeder negativen reellen Zahl als reelles Vielfache dieser Einheit darstellen lässt:

$$\sqrt{-5} = \sqrt{-1 \cdot 5} = \sqrt{-1} \cdot \sqrt{5} = \sqrt{5}i.$$

Alle reellen Vielfachen von i nennt man die *imaginären Zahlen*. Die Kombination von reellen und imaginären Zahlen liefern die komplexen Zahlen:

Definition: Ausdrücke der Form

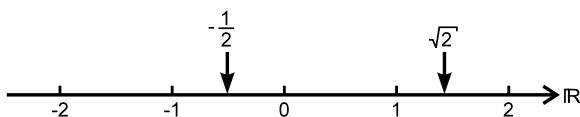
$$c := a + ib \quad \text{mit } a, b \in \mathbb{R}$$

nennt man **komplexe Zahlen** und $\mathbb{C} := \{c = a + ib; a, b \in \mathbb{R}\}$ die Menge der komplexen Zahlen.

Für $b = 0$ ist die Zahl $c = a + 0i = a \in \mathbb{R}$. Die reellen Zahlen sind also in den komplexen enthalten. Die mathematische Bedeutung der komplexen Zahlen liegt darin, dass jedes Polynom vom Grade n genau n Nullstellen besitzt ($\rightarrow$ 5.2.7 Fundamentalsatz der Algebra).

5.1 Darstellung komplexer Zahlen

Jede **reelle** Zahl entspricht einem Punkt auf der Zahlengeraden:



Durch die Definition der komplexen Zahlen als "Paare" $c = a + ib$ hat eine **komplexe** Zahl zwei "Komponenten": eine rein reelle Komponente a und eine imaginäre Komponente ib . Zur Darstellung von komplexen Zahlen geht man also in die Zahlenebene über.

5.1.1 Algebraische Normalform

Komplexe Zahlen

$$c := a + ib \quad \text{mit } a, b \in \mathbb{R}$$

lassen sich mit Hilfe von **zwei** Zahlengeraden veranschaulichen (Abb. 5.1): Wählt man ein Koordinatensystem mit Abszisse a (Vielfaches der Einheit 1) und Ordinate ib (Vielfaches der Einheit i), so ist jede komplexe Zahl ein Punkt dieser Ebene, der sog. *Gaußschen Zahlenebene*.

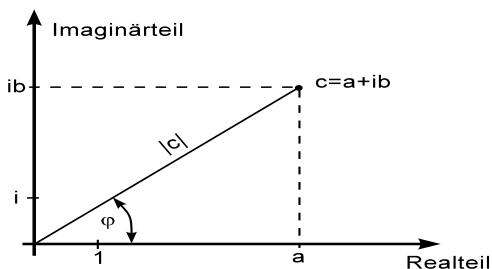


Abb. 5.1. Darstellung der komplexen Zahl $c = a + ib$.

Man nennt

$a = \operatorname{Re}(c)$ den Realteil von c
 $b = \operatorname{Im}(c)$ den Imaginärteil von c .

△ **Achtung:** Sowohl der Real- als auch der Imaginärteil einer komplexen Zahl sind reelle Zahlen. Man beachte daher: Der Imaginärteil einer komplexen Zahl $c = a + ib$ ist **nicht** ib , sondern nur die reelle Größe $\operatorname{Im}(c) = b$!

Man bezeichnet die Darstellung der komplexen Zahl

$$c = a + ib$$

(**Algebraische Normalform**)

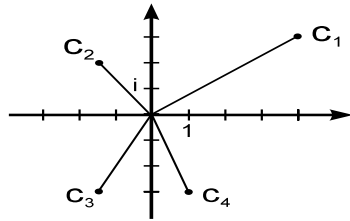
durch Realteil und Imaginärteil als algebraische Normalform. Als den **Betrag** einer komplexen Zahl definieren wir den Abstand zum Nullpunkt

$$|c| := \sqrt{a^2 + b^2} = \sqrt{\operatorname{Re}^2(c) + \operatorname{Im}^2(c)}$$

(**Betrag** von c).

Beispiele 5.1:

- ① $c_1 = 4 + 3i \quad \hookrightarrow \quad |c_1| = 5.$
- ② $c_2 = -\sqrt{2} + 2i \quad \hookrightarrow \quad |c_2| = \sqrt{6}.$
- ③ $c_3 = -\frac{3}{2} - 3i \quad \hookrightarrow \quad |c_3| = \sqrt{\frac{45}{4}}.$
- ④ $c_4 = 1 - 3i \quad \hookrightarrow \quad |c_4| = \sqrt{10}.$



Bemerkungen:

- (1) Zwei komplexe Zahlen $c_1 = a_1 + ib_1$ und $c_2 = a_2 + ib_2$ sind genau dann gleich, wenn $a_1 = a_2$ und $b_1 = b_2$. Realteil und Imaginärteil sind also zwei eindeutig bestimmte Kenngrößen einer komplexen Zahl.
- (2) Eine komplexe Zahl ist also nichts anderes als ein Punkt in der komplexen Zahlenebene.
- (3) Es ist üblich, den vom Ursprung O zum Punkte c weisenden Zeiger (Ortsvektor) ebenfalls mit c zu bezeichnen.

► 5.1.2 Trigonometrische Normalform

Führt man den Winkel φ zwischen dem komplexen Zeiger c und der positiven IR-Achse ein, so gilt nach Abb. 5.1

$$\cos \varphi = \frac{a}{|c|} \quad \text{und} \quad \sin \varphi = \frac{b}{|c|}.$$

Ersetzt man in der algebraischen Normalform $a = |c| \cos \varphi$ und $b = |c| \sin \varphi$,

gilt für die komplexe Zahl

$$c = a + i b = |c| \cos \varphi + i |c| \sin \varphi$$

$$c = |c| (\cos \varphi + i \sin \varphi) \quad (\text{Trigonometrische Normalform}).$$

Man nennt diese Darstellung die *trigonometrische Normalform*, mit

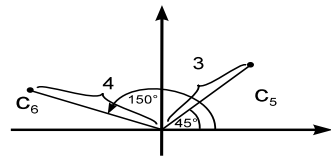
- $|c|$ dem *Betrag* der komplexen Zahl c und
- φ dem *Winkelargument* (Winkel, Argument, Phase) von c .

Für $c = 0$ ist φ nicht erklärt! Die Phase einer komplexen Zahl ist nicht eindeutig, denn bei jeder vollen Umdrehung wird die Phase um 2π bzw. um 360° verändert.

Beispiele 5.2:

① $c_5 = 3 (\cos 45^\circ + i \sin 45^\circ).$

② $c_6 = 4 (\cos 150^\circ + i \sin 150^\circ).$



➤ 5.1.3 Exponentielle Normalform

Ersetzen wir in der trigonometrischen Normalform $(\cos \varphi + i \sin \varphi)$ durch die von Euler (1707-1783) eingeführten Abkürzung

$$e^{i\varphi} := \cos \varphi + i \sin \varphi \quad (\text{Eulersche Formel}),$$

dann lässt sich jede komplexe Zahl schreiben als

$$c = |c| e^{i\varphi} \quad (\text{Exponentialform}).$$

Zunächst sehen wir die Eulersche Formel nur als Abkürzung an. Per Konvention wird das Argument φ bei der Exponentialform **immer** im Bogenmaß angegeben.

Beispiele 5.3:

- ① Exponentielle Normalform von c_5 : $\varphi = 45^\circ \hat{=} \frac{\pi}{4} \hookrightarrow c_5 = 3 e^{i\frac{\pi}{4}}.$
- ② Exponentielle Normalform von c_6 : $\varphi = 150^\circ \hat{=} \frac{5}{6}\pi \hookrightarrow c_6 = 4 e^{i\frac{5}{6}\pi}.$
- ③ Exponentielle Normalform von speziellen komplexen Zahlen:
 $e^{i\frac{\pi}{2}} = i; \quad e^{i\pi} = -1; \quad e^{i\frac{3}{2}\pi} = -i; \quad e^{2\pi i} = 1.$

□

5.1.4 Umformungen der Normalformen

Im Folgenden geben wir die Rechenschritte zur Umformung von den einzelnen Normalformen an. Bei den komplexen Rechenoperationen wählen wir dann immer die geeignete Normalform aus.

Exponentialdarstellung $\Rightarrow$ Trigonometrische Normalform:

Ist eine komplexe Zahl c in der Exponentialform $c = |c| e^{i\varphi}$ gegeben, so folgt mit der Eulerschen Formel direkt die trigonometrische Normalform

$$c = |c| (\cos \varphi + i \sin \varphi).$$

Ist die komplexe Zahl $c = |c| (\cos \varphi + i \sin \varphi)$ in der trigonometrischen Normalform gegeben, so folgt mit der Eulerschen Formel $c = |c| e^{i\varphi}$. Gegebenenfalls muss φ vom Grad- ins Bogenmaß umgerechnet werden.

Beispiele 5.4:

- ① $c_7 = 5 e^{i\frac{3}{4}\pi} \hookrightarrow \varphi = \frac{3}{4}\pi \hat{=} 135^\circ. \Rightarrow c_7 = 5 (\cos 135^\circ + i \sin 135^\circ).$
 ② $c_8 = \sqrt{2} (\cos 60^\circ + i \sin 60^\circ) \hookrightarrow \varphi = 60^\circ \hat{=} \frac{\pi}{3}. \Rightarrow c_8 = \sqrt{2} e^{i\frac{\pi}{3}}. \quad \square$

Trigonometrische Normalform $\Rightarrow$ Algebraische Normalform:

Ist die komplexe Zahl $c = |c| (\cos \varphi + i \sin \varphi)$ in der trigonometrischen Normalform gegeben, folgt durch Ausmultiplizieren und Auswerten der trigonometrischen Funktionen die algebraische Normalform:

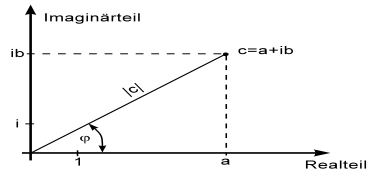
$$c = |c| \cos \varphi + i |c| \sin \varphi$$

mit dem Realteil $|c| \cos \varphi$ und dem Imaginärteil $|c| \sin \varphi$.

Ist die komplexe Zahl in der algebraischen Normalform $c = a + i b$ gegeben, folgt die trigonometrische Normalform, indem der Betrag $|c|$ und der Winkel φ bestimmt werden:

$$|c| = \sqrt{a^2 + b^2}$$

$$\tan \varphi = \frac{b}{a} \Rightarrow \varphi.$$



⚠ Achtung: Bei der Berechnung des Winkels $\tan \varphi = \frac{b}{a}$ durch die Umkehrfunktion $\arctan$ ist zu beachten, dass der Winkel nur im Bereich $[-\frac{\pi}{2}, \frac{\pi}{2}]$ angegeben wird (siehe Kap. 4.7). Der Winkel φ muss dann anhand einer Skizze im Bereich $[0, 2\pi]$ spezifiziert werden.

Beispiele 5.5:

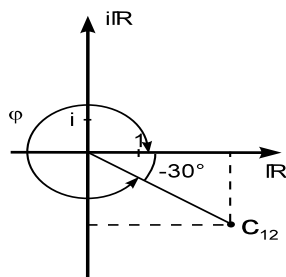
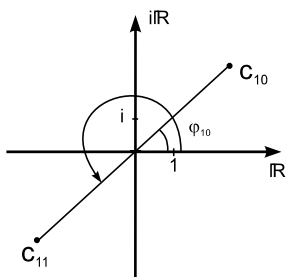
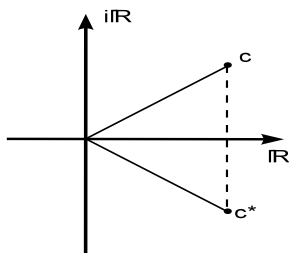
$$\textcircled{1} \quad c_9 = 5(\cos 135^\circ + i \sin 135^\circ) = 5\left(-\frac{1}{2}\sqrt{2}\right) + i 5\frac{1}{2}\sqrt{2} = -\frac{5}{2}\sqrt{2} + i \frac{5}{2}\sqrt{2}.$$

$$\begin{aligned} \textcircled{2} \quad c_{10} &= 4\sqrt{2} + i 4\sqrt{2}. \\ \hookrightarrow |c_{10}| &= \sqrt{16 \cdot 2 + 16 \cdot 2} = \sqrt{64} = 8, \\ &\& \tan \varphi = \frac{4\sqrt{2}}{4\sqrt{2}} = 1 \quad \hookrightarrow \varphi = 45^\circ \triangleq \frac{\pi}{4}. \\ &\Rightarrow c_{10} = 8(\cos 45^\circ + i \sin 45^\circ) = 8e^{i\frac{\pi}{4}}. \end{aligned}$$

$$\begin{aligned} \textcircled{3} \quad c_{11} &= -4\sqrt{2} - i 4\sqrt{2}. \\ \hookrightarrow |c_{11}| &= \sqrt{16 \cdot 2 + 16 \cdot 2} = \sqrt{64} = 8, \\ &\& \tan \varphi = \frac{-4\sqrt{2}}{-4\sqrt{2}} = 1 \quad \hookrightarrow \varphi = 45^\circ + 180^\circ = 225^\circ \triangleq \frac{5}{4}\pi. \\ &\Rightarrow c_{11} = 8(\cos 225^\circ + i \sin 225^\circ) = 8e^{i\frac{5}{4}\pi}. \end{aligned}$$

$$\begin{aligned} \textcircled{4} \quad c_{12} &= \sqrt{3} - i. \\ \hookrightarrow |c_{12}| &= \sqrt{3 + 1} = 2, \\ &\& \tan \varphi = \frac{-1}{\sqrt{3}} = -\frac{1}{3}\sqrt{3} \quad \hookrightarrow \varphi = -30^\circ = 330^\circ \triangleq \frac{11}{6}\pi. \\ &\Rightarrow c_{12} = 2(\cos 330^\circ + i \sin 330^\circ) = 2e^{i\frac{11}{6}\pi}. \end{aligned}$$

□

**⤵ Die komplex konjugierte Zahl**Abb. 5.2. c und c^*

Um die Division von zwei komplexen Zahlen zu bestimmen, benötigen wir noch einen neuen Begriff. Wir führen hierfür zu der komplexen Zahl c die **komplex konjugierte Zahl** c^* (bzw. $\bar{c}$) ein, die aus c durch Spiegelung an der reellen Achse hervorgeht:

Definition: $c^* := a - ib$ heißt die zu $c = a + ib$ **komplex konjugierte Zahl**.

Aufgrund der Definition der komplex konjugierten Zahl folgt

$$\begin{aligned} c = a + i b &\Rightarrow c^* = a - i b. \\ c = |c| (\cos \varphi + i \sin \varphi) &\Rightarrow c^* = |c| (\cos \varphi - i \sin \varphi). \\ c = |c| e^{i\varphi} &\Rightarrow c^* = |c| e^{-i\varphi}. \end{aligned}$$

Man erhält also die zu c komplex konjugierte Zahl sehr einfach, indem man formal i durch $-i$ ersetzt. Es gilt damit natürlich $(c^*)^* = c$.

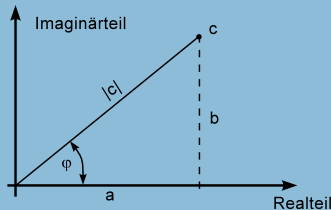
Zusammenfassung:

Die **imaginäre Einheit** $i := \sqrt{-1}$ ist definiert durch die Eigenschaft $i^2 = -1$.

Für komplexe Zahlen gibt es 3 Normalformen:

- (1) $c = a + i b$ **algebraische Normalform**
mit $a = \operatorname{Re}(c)$ (Realteil) und $b = \operatorname{Im}(c)$ (Imaginärteil).
- (2) $c = |c| \cdot (\cos \varphi + i \sin \varphi)$ **trigonometrische Normalform**
mit $|c| = \sqrt{a^2 + b^2}$ (Betrag) und $\tan \varphi = \frac{b}{a}$ (Winkel).
- (3) $c = |c| e^{i\varphi}$ **Exponentialform**
 φ wird hierbei im Bogenmaß angegeben.

Komplexe Zahlen lassen sich in der **Gaußschen Zahlenebene** graphisch darstellen.



Die zu c komplex konjugierte Zahl c^* lautet

$$c^* = a - i b = |c| (\cos \varphi - i \sin \varphi) = |c| e^{-i\varphi}.$$

5.2 Komplexe Rechenoperationen

Was unter Summe, Differenz, Produkt und Quotient zweier komplexer Zahlen zu verstehen ist, wird nicht durch die Konstruktion der komplexen Zahlen festgelegt. Man muss diese Verknüpfungen neu definieren; aber natürlich so, dass für den Spezialfall Imaginärteil gleich Null die bereits festgelegten Verknüpfungen in $\mathbb{R}$ herauskommen.

Seien im Folgenden $c_1 = a_1 + i b_1$ und $c_2 = a_2 + i b_2$ zwei beliebige komplexe Zahlen. Dann definiert man:

➤ 5.2.1 Addition

$$c_1 + c_2 := (a_1 + a_2) + i(b_1 + b_2)$$

Die Addition zweier komplexer Zahlen bedeutet die Addition der Realteile und die Addition der Imaginärteile. Die Addition wird in der algebraischen Normalform durchgeführt.

Beispiele 5.6:

① $c_1 = 9 - 2i$, $c_2 = 4 + i$.

$$c_1 + c_2 = (9 + 4) + i(-2 + 1) = 13 - i.$$

② $c_1 = 3(\cos 30^\circ + i \sin 30^\circ)$, $c_2 = 4 + i$. Um c_1 und c_2 zu addieren, muss die Zahl c_1 erst in die algebraische Normalform umgeformt werden:

$$c_1 = 3 \cos 30^\circ + i 3 \sin 30^\circ = 2,598 + i 1,5.$$

$$\Rightarrow c_1 + c_2 = (2,598 + i 1,5) + (4 + i) = 6,598 + i 2,5. \quad \square$$

➤ 5.2.2 Subtraktion

$$c_1 - c_2 := (a_1 - a_2) + i(b_1 - b_2)$$

Die Subtraktion zweier komplexer Zahlen bedeutet die Subtraktion der Realteile und die Subtraktion der Imaginärteile. Die Subtraktion wird in der algebraischen Normalform durchgeführt.

Beispiele 5.7:

① $c_1 = 9 - 2i$, $c_2 = 4 + i$.

$$c_1 - c_2 = (9 - 2i) - (4 + i) = 9 - 4 + i(-2 - 1) = 5 - i 3.$$

② $c_1 = 2e^{\frac{\pi}{4}i}$, $c_2 = 4 - 2i$. Um c_1 und c_2 voneinander zu subtrahieren, wird c_1 erst in die algebraische Normalform umgeformt:

$$\varphi = \frac{\pi}{4} \triangleq 45^\circ \hookrightarrow c_1 = 2e^{\frac{\pi}{4}i} = 2(\cos 45^\circ + i \sin 45^\circ)$$

$$= 2 \frac{1}{2} \sqrt{2} + i 2 \frac{1}{2} \sqrt{2} = 1,414 + i 1,414.$$

$$\Rightarrow c_1 - c_2 = (1,414 + i 1,414) - (4 - 2i) = -2,586 + 3,414i. \quad \square$$

Geometrische Interpretation. Da die Addition und Subtraktion zweier komplexer Zahlen analog den entsprechenden Regeln der Vektorrechnung erfolgen (nämlich komponentenweise), entspricht die graphische Darstellung der Rechenoperationen dem Kräfteparallelogramm, also der Vektoraddition bzw. -subtraktion.

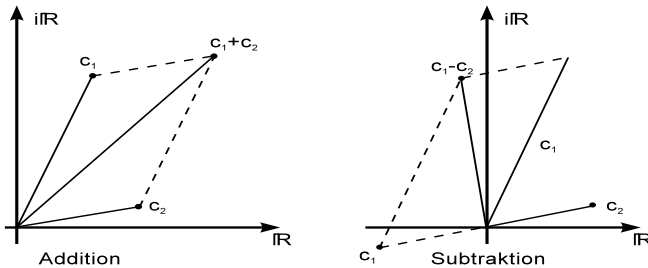


Abb. 5.3. Addition und Subtraktion von komplexen Zahlen

Bemerkung: Obwohl eine komplexe Zahl nur einen Punkt in der komplexen Zahlenebene darstellt, wird wegen obiger Interpretation der "Vektoraddition" eine komplexe Zahl oftmals mit dem *Zeiger* (*Ortsvektor*) identifiziert.

➤ 5.2.3 Multiplikation

$$c_1 \cdot c_2 := (a_1 a_2 - b_1 b_2) + i (a_1 b_2 + b_1 a_2)$$

Diese Formel für die Multiplikation ergibt sich, wenn $(a_1 + i b_1) \cdot (a_2 + i b_2)$ nach dem Distributivgesetz für reelle Zahlen gliedweise ausmultipliziert und die Definition von $i^2 = -1$ ausgenutzt wird:

$$\begin{aligned} c_1 \cdot c_2 &= (a_1 + i b_1) \cdot (a_2 + i b_2) \\ &= (a_1 a_2 + a_1 i b_2 + i b_1 a_2 + i b_1 i b_2) \\ &= a_1 a_2 + i^2 b_1 b_2 + i a_1 b_2 + i b_1 a_2 \\ &= (a_1 a_2 - b_1 b_2) + i (a_1 b_2 + b_1 a_2). \end{aligned}$$

Beispiele 5.8:

① $c_1 = 9 - 2i$, $c_2 = 4 + i$.

$$c_1 \cdot c_2 = (9 - 2i)(4 + i) = (36 + 2) + i(9 - 8) = 38 + i.$$

② Für das Produkt von $c = a + ib$ mit der komplex konjugierten Zahl $c^* = a - ib$ gilt

$$c \cdot c^* = (a + ib)(a - ib) = a^2 + b^2 = |c|^2.$$

Damit erhält man folgende wichtige Formel für $|c|$:

$$|c| = \sqrt{a^2 + b^2} = \sqrt{c \cdot c^*}$$

□

Geometrische Interpretation: Zur geometrischen Interpretation führen wir die Multiplikation nochmals aus, jetzt allerdings gehen wir von der trigonometrischen Normalform von

$$c_1 = |c_1|(\cos \varphi_1 + i \sin \varphi_1) \quad \text{und} \quad c_2 = |c_2|(\cos \varphi_2 + i \sin \varphi_2)$$

aus. Gliedweises ausmultiplizieren liefert

$$c_1 \cdot c_2 = |c_1|(\cos \varphi_1 + i \sin \varphi_1) \cdot |c_2|(\cos \varphi_2 + i \sin \varphi_2)$$

$$= |c_1| |c_2| \{[\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2] + i [\sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2]\}.$$

Wenden wir nun die Additionstheoreme für $\cos(\varphi_1 + \varphi_2)$ und $\sin(\varphi_1 + \varphi_2)$ aus Kapitel 4.6.4 an:

$$\begin{aligned} \cos(\varphi_1 + \varphi_2) &= \cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2 \\ \sin(\varphi_1 + \varphi_2) &= \sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2, \end{aligned}$$

so erhalten wir als Produkt

$$c_1 \cdot c_2 = |c_1| \cdot |c_2| \cdot (\cos(\varphi_1 + \varphi_2) + i \sin(\varphi_1 + \varphi_2)).$$

Die Multiplikation zweier komplexer Zahlen bedeutet die **Multiplikation der Beträge** und die **Addition der Winkel**. Dadurch kann der Punkt $c_1 \cdot c_2$ leicht in der Gaußschen Zahlenebene konstruiert werden.

Für die Darstellung in der Exponentialform folgt

$$c_1 \cdot c_2 = |c_1| e^{i\varphi_1} \cdot |c_2| e^{i\varphi_2} = |c_1| |c_2| e^{i(\varphi_1 + \varphi_2)}.$$

Dies entspricht genau der Eigenschaft der reellen Exponentialfunktion:

$$e^{x_1} \cdot e^{x_2} = e^{x_1 + x_2}.$$

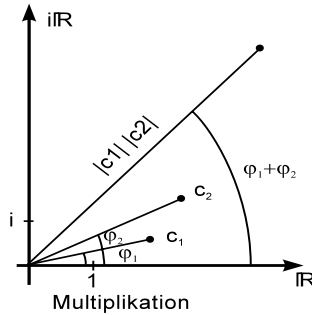


Abb. 5.4. Multiplikation zweier komplexer Zahlen

► 5.2.4 Division

$$\frac{c_1}{c_2} := \frac{a_1 a_2 + b_1 b_2}{(a_2)^2 + (b_2)^2} + i \frac{b_1 a_2 - a_1 b_2}{(a_2)^2 + (b_2)^2} \quad \text{für } c_2 \neq 0$$

Diese Formel für die Division ergibt sich, wenn man formal $\frac{c_1}{c_2}$ mit c_2^* erweitert und Zähler bzw. Nenner ausmultipliziert:

$$\begin{aligned} \frac{c_1}{c_2} &= \frac{c_1}{c_2} \cdot \frac{c_2^*}{c_2^*} = \frac{a_1 + i b_1}{a_2 + i b_2} \cdot \frac{a_2 - i b_2}{a_2 - i b_2} = \frac{(a_1 + i b_1)(a_2 - i b_2)}{(a_2 + i b_2)(a_2 - i b_2)} \\ &= \frac{(a_1 a_2 + b_1 b_2) + i (b_1 a_2 - a_1 b_2)}{(a_2)^2 + (b_2)^2}. \end{aligned}$$

⚠ Auch in $\mathbb{C}$ ist die Division durch $0 = 0 + i 0$ **nicht** erlaubt!

Geometrische Interpretation: Führt man die Division in der trigonometrischen Normalform durch, so erhält man unter Verwendung der trigonometrischen Formeln für $\cos(\varphi_1 - \varphi_2)$ und $\sin(\varphi_1 - \varphi_2)$ analog dem Vorgehen unter Abschnitt 5.2.3

$$\frac{c_1}{c_2} = \frac{|c_1|(\cos \varphi_1 + i \sin \varphi_1)}{|c_2|(\cos \varphi_2 + i \sin \varphi_2)} = \frac{|c_1|}{|c_2|} (\cos(\varphi_1 - \varphi_2) + i \sin(\varphi_1 - \varphi_2))$$

sowie

$$\frac{c_1}{c_2} = \frac{|c_1| e^{i\varphi_1}}{|c_2| e^{i\varphi_2}} = \frac{|c_1|}{|c_2|} e^{i(\varphi_1 - \varphi_2)}.$$

Bei der Quotientenbildung zweier komplexer Zahlen werden die **Beträge dividiert** und die **Winkel subtrahiert**. Damit ist $\frac{c_1}{c_2}$ ebenfalls in der Gaußschen Zahlenebene geometrisch zu konstruieren.

Beispiele 5.9:

① $c_1 = 9 - 2i, c_2 = 4 + i.$

Um $\frac{c_1}{c_2}$ zu berechnen, erweitern wir den Quotienten mit c_2^* und multiplizieren Zähler und Nenner aus:

$$\frac{c_1}{c_2} = \frac{9 - 2i}{4 + i} \cdot \frac{4 - i}{4 - i} = \frac{(9 \cdot 4 - 2 \cdot 1) + i(-2 \cdot 4 - 9 \cdot 1)}{17} = 2 - i.$$

② $c_1 = 8e^{i\frac{4}{3}\pi}, c_2 = 4(\cos 60^\circ + i \sin 60^\circ).$

Um $\frac{c_1}{c_2}$ zu berechnen, stellen wir c_2 in der Exponentialform dar. Da $60^\circ \triangleq \frac{\pi}{3}$ gilt

$$c_2 = 4(\cos 60^\circ + i \sin 60^\circ) = 4e^{i\frac{\pi}{3}}.$$

Damit folgt

$$\frac{c_1}{c_2} = \frac{8e^{i\frac{4}{3}\pi}}{4e^{i\frac{\pi}{3}}} = 2e^{i(\frac{4}{3}\pi - \frac{\pi}{3})} = 2e^{i\pi} = -2. \quad \square$$

Beispiel 5.10. Gegeben seien $c_1 = 1 + i\sqrt{3}$ und $c_2 = -\sqrt{3} + 3i$. Man berechne (i) $c_1 \cdot c_2$ und (ii) $\frac{c_1}{c_2}$. (iii) Man bestimme die exponentielle Normalform der Zahlen und führe nochmals die (iv) Multiplikation bzw. (v) die Division durch.

$$(i) \quad c_1 \cdot c_2 = (1 + i\sqrt{3})(-\sqrt{3} + 3i) = (-\sqrt{3} - 3\sqrt{3}) + i(3 - 3) = -4\sqrt{3}.$$

$$(ii) \quad \frac{c_1}{c_2} = \frac{1 + i\sqrt{3}}{-\sqrt{3} + 3i} \cdot \frac{-\sqrt{3} - 3i}{-\sqrt{3} - 3i} = \frac{-\sqrt{3} + 3\sqrt{3} - 3i - 3i}{3 + 9} = \frac{\sqrt{3}}{6} - \frac{1}{2}i.$$

(iii) Darstellung von c_1 und c_2 in exponentieller Normalform

$$|c_1| = \sqrt{1 + 3} = 2; \tan \varphi = \frac{\sqrt{3}}{1} = \sqrt{3} \Rightarrow \varphi = 60^\circ \triangleq \frac{\pi}{3} \Rightarrow c_1 = 2e^{i\frac{\pi}{3}}.$$

$$|c_2| = 2\sqrt{3}; \tan \varphi = -\frac{3}{\sqrt{3}} \Rightarrow \varphi = \pi - \frac{\pi}{3} = \frac{2}{3}\pi \Rightarrow c_2 = 2\sqrt{3}e^{i\frac{2}{3}\pi}.$$

$$(iv) \quad c_1 \cdot c_2 = 2e^{i\frac{\pi}{3}} \cdot 2\sqrt{3}e^{i\frac{2}{3}\pi} = 4\sqrt{3}e^{i\pi} = -4\sqrt{3}.$$

$$(v) \quad \frac{c_1}{c_2} = \frac{2}{2\sqrt{3}}e^{i(\frac{\pi}{3} - \frac{2}{3}\pi)} = \frac{1}{\sqrt{3}}e^{-i\frac{\pi}{3}}. \quad \square$$

Zusammenfassung:

Addition, Subtraktion und Multiplikation werden formal wie bei reellen Zahlen ausgeführt, wobei $i^2 = -1$ zu ersetzen ist.

Die Division $\frac{c_1}{c_2}$ wird durch Erweiterung mit c_2^* berechnet. Die Ergebnisse werden in die Form $a + ib$ ($a, b \in \mathbb{R}$) gebracht.

Multiplikation und Division lassen sich in der trigonometrischen bzw. exponentiellen Normalform sehr einfach ausführen: Bei der Multiplikation werden die Beträge multipliziert und die Winkel addiert, während bei der Division die Beträge dividiert und die Winkel subtrahiert werden.

► 5.2.5 Potenz

Die Potenz c^n ($n \in \mathbb{N}$) einer komplexen Zahl gestaltet sich in der trigonometrischen bzw. exponentiellen Normalform als besonders einfach. Gehen wir von der komplexen Zahl c in der exponentiellen Normalform aus: $c = |c| e^{i\varphi}$. Dann gilt

$$c^2 = c \cdot c = |c| e^{i\varphi} \cdot |c| e^{i\varphi} = |c|^2 e^{i2\varphi}$$

$$c^3 = c^2 \cdot c = |c|^2 e^{i2\varphi} \cdot |c| e^{i\varphi} = |c|^3 e^{i3\varphi}$$

usw.

Durch vollständige Induktion weist man direkt nach, dass gilt

$$c = |c| (\cos \varphi + i \sin \varphi) \quad \Rightarrow \quad c^n = |c|^n (\cos(n\varphi) + i \sin(n\varphi))$$

bzw.

$$c = |c| e^{i\varphi} \quad \Rightarrow \quad c^n = |c|^n e^{in\varphi}.$$

Diese sog. *Moivresche Formel* besagen, dass man c^n dadurch erhält, indem der Betrag potenziert und der Winkel mit n multipliziert wird.

Beispiele 5.11:

① Gesucht ist $(2\sqrt{2} + i2\sqrt{2})^5$.

Um die komplexe Zahl $c = 2\sqrt{2} + i2\sqrt{2}$ mit 5 zu potenzieren, müssen wir sie zuerst in der exponentiellen Normalform darstellen:

$$\left. \begin{aligned} |c| &= \sqrt{4 \cdot 2 + 4 \cdot 2} = \sqrt{16} = 4 \\ \tan \varphi &= \frac{2\sqrt{2}}{2\sqrt{2}} = 1 \hookrightarrow \varphi = \frac{\pi}{4} \end{aligned} \right\} \Rightarrow c = 4 e^{i\frac{\pi}{4}}$$

$$\Rightarrow c^5 = 4^5 e^{i\frac{\pi}{4} \cdot 5} = 1024 e^{i\frac{5}{4}\pi}.$$

② Gesucht ist $(\sqrt{3} - i)^6$.

Nach Beispiel 5.5 ④ ist $c = \sqrt{3} - i = 2 e^{i\frac{11}{6}\pi}$.

$$\Rightarrow c^6 = \left(2 e^{i\frac{11}{6}\pi}\right)^6 = 2^6 e^{i\frac{11}{6}\pi \cdot 6} = 64 e^{i11\pi} = -64.$$

□

5.2.6 Wurzeln

Für $c = |c| (\cos \varphi + i \sin \varphi) = |c| e^{i\varphi}$ ist die n -te **Wurzel** ($n \in \mathbb{N}$) gegeben durch

$$\begin{aligned} c^{\frac{1}{n}} &= \left\{ \sqrt[n]{|c|} \left(\cos \left(\frac{\varphi + k \cdot 360^\circ}{n} \right) + i \sin \left(\frac{\varphi + k \cdot 360^\circ}{n} \right) \right); \right. \\ &\quad \left. k = 0, 1, \dots, n-1 \right\} \\ &= \left\{ \sqrt[n]{|c|} e^{i \frac{\varphi + k \cdot 2\pi}{n}}; k = 0, 1, 2, \dots, n-1 \right\}, \quad (*) \end{aligned}$$

wenn $\sqrt[n]{|c|}$ die **reelle** n -te Wurzel von $|c| \geq 0$.

Begründung: Um zu zeigen, dass die komplexen Zahlen

$$W_k := \sqrt[n]{|c|} e^{i \frac{\varphi + k \cdot 2\pi}{n}}; k = 0, 1, 2, \dots, n-1$$

n -te Wurzel von c sind, genügt es zu zeigen, dass $(W_k)^n = c$. Denn die n -te Wurzel einer komplexen Zahl hat die universelle Eigenschaft, dass sie zur n -ten Potenz genommen genau c ergeben muss! Dies ist aber aufgrund der Rechenregeln für das Potenzieren offensichtlich:

$$(W_k)^n = \left(\sqrt[n]{|c|} \right)^n e^{i \frac{\varphi + k \cdot 2\pi}{n} \cdot n} = |c| e^{i(\varphi + k \cdot 2\pi)} = |c| e^{i\varphi},$$

wenn man beachtet, dass $e^{i(\varphi + k \cdot 2\pi)} = e^{i\varphi}$ für $k \in \mathbb{N}$ ist. □

Die n -ten Wurzeln W_k sind für $k = 0, \dots, n-1$ voneinander verschieden, wiederholen sich aber für $k \geq n$. Man beachte also, dass die n -te Potenz einer komplexen Zahl **eindeutig**, die n -ten Wurzeln aber **mehrdeutig** sind.

Beispiele 5.12:

$$\begin{aligned} \textcircled{1} \quad (4\sqrt{2} + i 4\sqrt{2})^{\frac{1}{3}} &= (8 e^{i \frac{\pi}{4}})^{\frac{1}{3}} = \{ \sqrt[3]{8} e^{i \frac{\frac{\pi}{4} + k \cdot 2\pi}{3}}; k = 0, 1, 2 \} \\ &= \{ 2 e^{i \frac{\pi}{12}}, 2 e^{i \frac{9}{12} \pi}, 2 e^{i \frac{17}{12} \pi} \}. \end{aligned}$$

$$\begin{aligned} \textcircled{2} \quad (-1)^{\frac{1}{5}} &= (1 e^{i\pi})^{\frac{1}{5}} = \{ \sqrt[5]{1} e^{i \frac{\pi + k \cdot 2\pi}{5}}; k = 0, 1, 2, 3, 4 \} \\ &= \{ e^{i \frac{\pi}{5}}, e^{i \frac{3}{5} \pi}, e^{i\pi}, e^{i \frac{7}{5} \pi}, e^{i \frac{9}{5} \pi} \}. \end{aligned} \quad \square$$

Sonderfall: Die n -ten Wurzeln aus 1: Jede komplexe Lösung von $Z^n = 1$ heißt n -te *Einheitswurzel*. Mit Formel (*) folgt für $c = 1$:

$$1^{\frac{1}{n}} = (1 e^{i0})^{\frac{1}{n}} = \{ 1, e^{i \frac{2\pi}{n}}, e^{i \frac{4\pi}{n}}, \dots, e^{i \frac{2\pi(n-1)}{n}} \}.$$

Der Betrag dieser Zahlen ist jeweils 1, d.h. die n -ten Einheitswurzeln liegen auf dem Einheitskreis. Die Differenz der Winkel ist jeweils $\frac{2\pi}{n}$, so dass sie nacheinander durch Drehung um $\frac{2\pi}{n}$ aus der 1 hervorgehen.

Beispiel 5.13. Gesucht sind alle 9.-ten Einheitswurzeln:

$$\begin{aligned}(1)^{\frac{1}{9}} &= (1 e^{i0})^{\frac{1}{9}} = \{\sqrt[9]{1} e^{i \frac{0+k \cdot 2\pi}{9}}; k = 0, \dots, 8\} \\ &= \{1 e^0, e^{i \frac{2\pi}{9}}, e^{i \frac{4\pi}{9}}, \dots, e^{i \frac{16\pi}{9}}\}.\end{aligned}$$

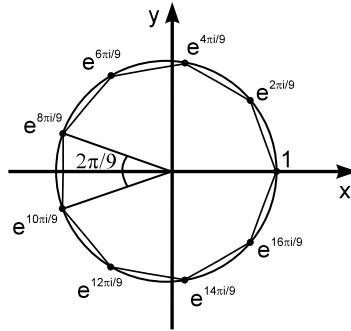


Abb. 5.5. 9-te Einheitswurzel von $c = 1$

Satz: Für $n > 1$ gilt:

$$\sum_{k=0}^{n-1} e^{i \frac{\varphi + k \cdot 2\pi}{n}} = 0.$$

Begründung: Dieser Satz ist aufgrund seiner geometrischen Eigenschaft offensichtlich, da $e^{i \frac{\varphi + k \cdot 2\pi}{n}}$ die n -te Einheitswurzel der komplexen Zahl $e^{i\varphi}$ darstellt. Summiert man alle n Einheitswurzeln auf (Vektoraddition), so ergibt die Summe Null; formal erhält man diese Aussage über die geometrische Reihe, denn

$$\begin{aligned}\sum_{k=0}^{n-1} e^{i \frac{\varphi + k \cdot 2\pi}{n}} &= \sum_{k=0}^{n-1} e^{i \frac{\varphi}{n}} e^{i k \frac{2\pi}{n}} = e^{i \frac{\varphi}{n}} \sum_{k=0}^{n-1} \left(e^{i \frac{2\pi}{n}}\right)^k = e^{i \frac{\varphi}{n}} \frac{1 - \left(e^{i \frac{2\pi}{n}}\right)^n}{1 - e^{i \frac{2\pi}{n}}} = 0, \\ \text{da } 1 - \left(e^{i \frac{2\pi}{n}}\right)^n &= 1 - e^{i 2\pi} = 1 - 1 = 0.\end{aligned}$$

□



Visualisierung mit MAPLE: Auf der CD-Rom befinden sich Worksheets, um **komplexe Zahlen** und die **komplexen Rechenoperationen** graphisch darzustellen bzw. in Form von Animationen zu visualisieren.

► 5.2.7 Fundamentalsatz der Algebra

Wir interpretieren die Mehrdeutigkeit der n -ten Wurzel folgendermaßen: Jedes Polynom n -ten Grades der Form $p(z) = z^n - a$ ($n \in \mathbb{N}$, $a \in \mathbb{C}$) hat genau n Nullstellen, nämlich die n -ten Wurzeln von a . Diese Eigenschaft lässt sich auf beliebige komplexe Polynome vom Grade n verallgemeinern. Dies ist der Inhalt des *Fundamentalsatzes der Algebra*, der auf F. Gauß (1797) zurückgeht:

Satz: Jedes komplexe Polynom n -ten Grades

$$p(z) = a_n z^n + a_{n-1} z^{n-1} + \dots + a_1 z + a_0$$

$$(a_k \in \mathbb{C}, a_n \neq 0, z \in \mathbb{C})$$

besitzt **genau** n Nullstellen.

Zusatz: Sind die Koeffizienten von $p(z)$ reell (d.h. $a_k \in \mathbb{R}$), so sind die Nullstellen reell oder sie treten paarweise komplex konjugiert auf.

Der Fundamentalsatz stellt zwar sicher, dass jedes Polynom n -ten Grades n Nullstellen besitzt, er sagt aber nichts darüber aus, wie diese Nullstellen zu finden sind. Es gibt auch im Komplexen außer in einfachen Spezialfällen keine allgemeine Formel, wie die Nullstellen berechnet werden können. Somit bleibt wie im Reellen: Entweder die Nullstellen zu erraten und durch Polynomdivision den Grad zu reduzieren oder sie numerisch zu bestimmen.

Beispiel 5.14. Gesucht sind die Nullstellen von $p(z) = z^3 - 2z - 4$.

Der Fundamentalsatz besagt, dass es genau 3 Nullstellen gibt. Um eine Nullstelle zu erhalten, probieren wir $z = 0, \pm 1, \pm 2$: $\hookrightarrow z = 2$ ist eine Nullstelle. Polynomdivision:

$$\begin{array}{r} (z^3 \quad -2z \quad -4) : (z - 2) = z^2 + 2z + 2. \\ \underline{z^3 \quad -2z^2} \\ 2z^2 \quad -2z \\ \underline{2z^2 \quad -4z} \\ 2z \quad -4 \\ \underline{2z \quad -4} \\ 0 \end{array}$$

$$\Rightarrow (z^3 - 2z - 4) = (z - 2)(z^2 + 2z + 2).$$

Die quadratische Formel liefert $z_{2/3} = -1 \pm \sqrt{1 - 2} = -1 \pm i$.

Die Nullstellen des Polynoms sind also: $2, -1 + i, -1 - i$. □

Bemerkung: Der Zusatz zum Fundamentalsatz lässt sich direkt nachrechnen: Ist $p(z) = a_n z^n + a_{n-1} z^{n-1} + \dots + a_1 z + a_0$ ein reelles Polynom und z_0 eine Nullstelle von p , dann ist z_0^* ebenfalls eine Nullstelle von p :

$$\begin{aligned} p(z_0^*) &= a_n (z_0^*)^n + a_{n-1} (z_0^*)^{n-1} + \dots + a_1 z_0^* + a_0 \\ &= a_n (z_0^n)^* + a_{n-1} (z_0^{n-1})^* + \dots + a_1 (z_0)^* + a_0 \\ &= (a_n z_0^n + a_{n-1} z_0^{n-1} + \dots + a_1 z_0 + a_0)^* = (p(z_0))^* = 0^* = 0. \quad \square \end{aligned}$$

5.3 Anwendungen

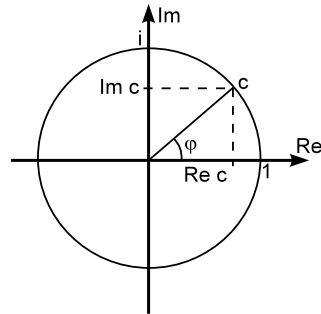
► 5.3.1 Beschreibung harmonischer Schwingungen im Komplexen

Das aus der Mechanik bekannte Federpendel hat die Eigenschaft, dass bei einer ungedämpften Schwingung die Auslenkung aus der Ruhelage $s(t)$ den zeitlichen Verlauf $s(t) = A \cos(\omega t + \varphi)$ besitzt. Das System schwingt mit der Frequenz $\omega = \sqrt{\frac{D}{m}}$, wenn D die Federkonstante und m die Masse ist. Diese Funktion besitzt eine zeitlich konstante Maximalamplitude A und die Nullphase φ . Die Schwingungsdauer beträgt $T = \frac{2\pi}{\omega}$. Eine periodische Bewegung mit einer Frequenz ω und zeitlich konstanter Maximalamplitude A nennt man **harmonische Schwingung**. Das zum Federpendel elektrische Analogon ist der Spannungsverlauf $U(t) = U_0 \cos(\omega t + \varphi)$ in einem LC-Wechselstromkreis mit der Frequenz $\omega = \sqrt{\frac{1}{LC}}$ bzw. der Schwingungsdauer $T = 2\pi\sqrt{LC}$.

Zur Beschreibung von harmonischen Schwingungen im Komplexen betrachten wir zunächst die komplexe Zahl

$$c = \cos \varphi + i \sin \varphi = e^{i\varphi}.$$

Da $|c| = 1$, ist c eine Zahl auf dem Einheitskreis. Projiziert man den Punkt c auf die reelle Achse, erhält man den Realteil von c : $\operatorname{Re}(c) = \cos \varphi$; projiziert man den Punkt c auf die imaginäre Achse, so erhält man den Imaginärteil von c : $\operatorname{Im}(c) = \sin \varphi$.



Variiert der Winkel φ als Funktion der Zeit $\varphi = \omega \cdot t$ ($\omega = \frac{2\pi}{T}$ konstante Kreisfrequenz), durchläuft $e^{i\varphi} = e^{i\omega t}$ für $0 \leq t \leq T$ den Einheitskreis in der komplexen Ebene.

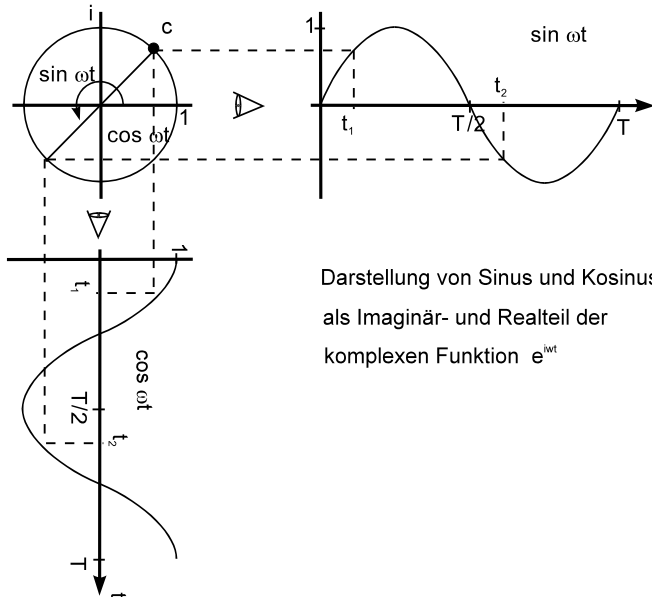
$\cos \omega t$ und $\sin \omega t$ sind die Projektionen des komplexen Zeigers $e^{i\omega t}$ auf die reelle bzw. auf die imaginäre Achse.



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich die Prozedur **projektion**, welche die Projektionen von $e^{i\omega t}$ auf die x - bzw. y -Achse zeichnet. Der im Einheitskreis laufenden Zeiger $e^{i\omega t}$ wird zusammen mit seinem Real- und Imaginärteil animiert dargestellt, indem die Variable t von 0 bis $\frac{2\pi}{T}$ variiert.

Mit einer Nullphase φ_0 folgt die komplexe Darstellung in der Form

$$\begin{aligned}\cos(\omega t + \varphi_0) &= \operatorname{Re}(e^{i(\omega t + \varphi_0)}) \\ \sin(\omega t + \varphi_0) &= \operatorname{Im}(e^{i(\omega t + \varphi_0)}).\end{aligned}$$



Darstellung von Sinus und Kosinus
als Imaginär- und Realteil der
komplexen Funktion $e^{i\omega t}$

Allgemein lässt sich damit eine harmonische Schwingung $y(t)$ im Komplexen schreiben als

$$\begin{aligned}\hat{y}(t) &= A(\cos(\omega t + \varphi_0) + i \sin(\omega t + \varphi_0)) \\ &= A e^{i(\omega t + \varphi_0)} = A e^{i\varphi_0} e^{i\omega t}.\end{aligned}$$

Also ist die komplexe Beschreibung einer harmonischen Schwingung

$$\hat{y}(t) = A e^{i\varphi_0} e^{i\omega t}$$

gegeben durch die *komplexe Amplitude* $A e^{i\varphi_0}$ und dem reinen Zeitanteil $e^{i\omega t}$.

► 5.3.2 Superposition gleichfrequenter Schwingungen

Im Folgenden werden wir die Überlagerung (*Superposition*) zweier **gleich-frequenter** harmonischer Schwingungen im Komplexen berechnen. Gegeben

seien z.B. zwei Wechselspannungen

$$\begin{aligned}u_1(t) &= u_1 \sin(\omega t + \varphi_1) \\ u_2(t) &= u_2 \sin(\omega t + \varphi_2).\end{aligned}$$

Gesucht ist die Überlagerung

$$u(t) = u_1(t) + u_2(t) = A \sin(\omega t + \varphi)$$

mit Amplitude A und Phase φ .

Zur Berechnung der Überlagerung interpretieren wir $u_1(t)$ als Imaginärteil der komplexen Schwingung $\hat{u}_1(t) = u_1 e^{i(\omega t + \varphi_1)}$ und $u_2(t)$ als Imaginärteil von $\hat{u}_2(t) = u_2 e^{i(\omega t + \varphi_2)}$ und führen die Überlagerung im Komplexen durch. Anschließend nehmen wir von dem Ergebnis die imaginäre Komponente; sie entspricht dann $u_1(t) + u_2(t)$.

Methode:

$$u_1(t) + u_2(t) = \operatorname{Im} \hat{u}_1(t) + \operatorname{Im} \hat{u}_2(t) = \operatorname{Im} (\hat{u}_1(t) + \hat{u}_2(t)) = \operatorname{Im} \hat{u}(t) = u(t).$$

$$\begin{array}{lcl} \text{Übergang} & & \\ \hline \text{ins Komplexe} & \hat{u}_1(t) & = u_1 e^{i(\omega t + \varphi_1)} = u_1 e^{i\varphi_1} e^{i\omega t}, \\ & \hat{u}_2(t) & = u_2 e^{i(\omega t + \varphi_2)} = u_2 e^{i\varphi_2} e^{i\omega t}. \end{array}$$

$$\begin{array}{lcl} \text{Komplexe} & \hat{u}(t) & = \hat{u}_1(t) + \hat{u}_2(t) \\ \hline \text{Addition} & & = u_1 e^{i\varphi_1} e^{i\omega t} + u_2 e^{i\varphi_2} e^{i\omega t} \\ & & = (u_1 e^{i\varphi_1} + u_2 e^{i\varphi_2}) e^{i\omega t} \\ & & = A e^{i\varphi} e^{i\omega t} = A e^{i(\omega t + \varphi)}. \end{array}$$

Die komplexe Amplitude der Superposition $A e^{i\varphi}$ ergibt sich als Summe der beiden Einzelamplituden $u_1 e^{i\varphi_1}$ und $u_2 e^{i\varphi_2}$. Die Überlagerung zweier gleichfrequenter Schwingungen entspricht der vektoriellen Addition der komplexen Amplituden. Die komplexe Addition kann damit sowohl zeichnerisch als auch rechnerisch durchgeführt werden.

$$\begin{array}{lcl} \text{Übergang} & & \\ \hline \text{ins Reelle} & u(t) = u_1(t) + u_2(t) & = \operatorname{Im} (A e^{i(\omega t + \varphi)}) = A \sin(\omega t + \varphi). \end{array}$$

Bemerkungen:

- (1) Auf dieselbe Weise erhält man die Überlagerung zweier Kosinusschwingungen $u_1(t) = u_1 \cos(\omega t + \varphi_1)$, $u_2(t) = u_2 \cos(\omega t + \varphi_2)$, indem diese Schwingungen als *Realteil* der entsprechenden komplexen Schwingungen interpretiert werden. $u(t) = \operatorname{Re} (A e^{i(\omega t + \varphi)})$ liefert dann den Kosinusannteil der Superposition.

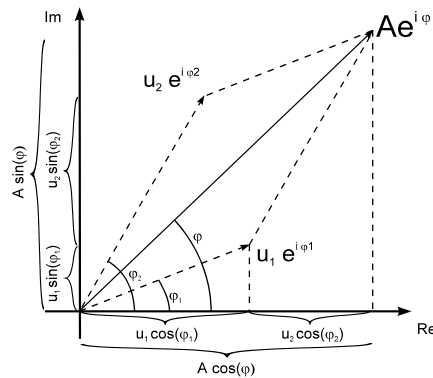


Abb. 5.6. Graphische Addition der komplexen Amplituden

- (2) Ist eine Schwingung in Kosinusdarstellung $u_1(t) = a_1 \cos(\omega t + \varphi_1)$ und die andere in der Sinusdarstellung $u_2(t) = a_2 \sin(\omega t + \varphi_2)$ gegeben, so muss eine gemeinsame Darstellungsform gewählt werden. Entweder man schreibt

$$u_1(t) = a_1 \cos(\omega t + \varphi_1) = a_1 \sin\left(\omega t + \varphi_1 + \frac{\pi}{2}\right)$$

und führt die Überlagerung in der Sinusform durch oder man schreibt

$$u_2(t) = a_2 \sin(\omega t + \varphi_2) = a_2 \cos\left(\omega t + \varphi_2 - \frac{\pi}{2}\right)$$

für die andere Schwingung und führt die Überlagerung in der Kosinusform durch.

- (3) **⚠ Achtung:** Die Überlagerung zweier harmonischer Schwingungen **unterschiedlicher** Frequenzen führt i.A. nicht mehr zu einer periodischen Funktion. Nur im Fall, dass das Verhältnis der Frequenzen eine gebrochenrationale Zahl ist, erhält man wieder eine periodische Funktion, aber auch dann keine harmonische mehr. Siehe auch das zugehörige [MAPLE-Worksheet](#).

Zusammenfassend gilt also

Satz: Besitzen zwei harmonische Schwingungen die **gleiche** Frequenz ω

$$u_1(t) = u_1 \sin(\omega t + \varphi_1), \quad u_2(t) = u_2 \sin(\omega t + \varphi_2),$$

dann ist die Superposition wieder eine harmonische Schwingung mit

$$u(t) = u_1(t) + u_2(t) = A \sin(\omega t + \varphi).$$

Beispiel 5.15 (Mit MAPLE-Worksheet). Gesucht ist die Überlagerung der beiden Wechselspannungen

$$u_1(t) = 4 \sin(2t) \quad \text{und} \quad u_2(t) = 3 \cos\left(2t - \frac{\pi}{6}\right).$$

Bevor man diese beiden harmonischen Funktionen überlagert, stellt man z.B. $u_2(t)$ als Sinusfunktion dar:

$$u_2(t) = 3 \cos\left(2t - \frac{\pi}{6}\right) = 3 \sin\left(2t - \frac{\pi}{6} + \frac{\pi}{2}\right) = 3 \sin\left(2t + \frac{\pi}{3}\right).$$

$$\begin{array}{l} \text{Übergang} \\ \text{ins Komplexe} \end{array} \quad \begin{array}{l} \hat{u}_1(t) = 4 e^{i2t} \\ \hat{u}_2(t) = 3 e^{i(2t + \frac{\pi}{3})} = 3 e^{i\frac{\pi}{3}} e^{i2t}. \end{array}$$

$$\begin{array}{l} \text{Komplexe} \\ \text{Addition} \end{array} \quad \begin{array}{l} \hat{u}(t) = \hat{u}_1(t) + \hat{u}_2(t) \\ = 4 e^{i2t} + 3 e^{i\frac{\pi}{3}} e^{i2t} = (4 + 3 e^{i\frac{\pi}{3}}) e^{i2t}. \end{array}$$

Addition der komplexen Amplituden in der algebraischen Normalform
 $c = 4 + 3 e^{i\frac{\pi}{3}} = 4 + 3 \left(\cos \frac{\pi}{3} + i \sin \frac{\pi}{3}\right) = 4 + 3 \left(\frac{1}{2} + i \frac{1}{2} \sqrt{3}\right) = 5,5 + i 2,6.$

Darstellung von c in Exponentialform

$$\begin{aligned} |c| &= \sqrt{5,5^2 + 2,6^2} = 6,08, \\ \tan \varphi &= \frac{\operatorname{Im} c}{\operatorname{Re} c} = \frac{2,6}{5,5} \quad \hookrightarrow \varphi = 25,28^\circ \hat{=} 0,44. \\ \Rightarrow c &= A e^{i\varphi} \quad \text{mit } A = 6,08 \quad \varphi = 0,44 \end{aligned}$$

$$\Rightarrow \hat{u}(t) = 6,08 e^{i0,44} e^{i2t} = 6,08 e^{i(2t + 0,44)}.$$

$$\begin{array}{l} \text{Übergang} \\ \text{ins Reelle} \end{array} \quad u(t) = \operatorname{Im} \hat{u}(t) = 6,08 \sin(2t + 0,44).$$

In Abb. 5.7 ist die Überlagerung der beiden Schwingungen graphisch dargestellt:

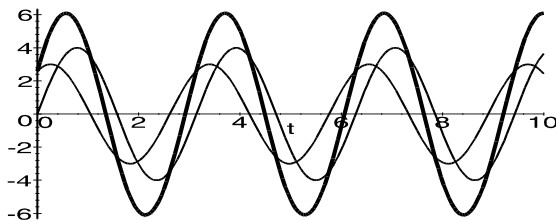


Abb. 5.7. Überlagerung zweier gleichfrequenter Schwingungen

Dieses Verfahren lässt sich leicht auf den Fall der Überlagerung von mehr als zwei harmonischen Schwingungen mit gleichen Frequenzen übertragen. $\square$

► 5.3.3 Beschreibung von RCL-Gliedern bei Wechselströmen

Wir betrachten elektrische Netzwerke, die sich aus Ohmschen Widerständen, Kapazitäten und Induktivitäten zusammensetzen. In Wechselstromkreisen besitzen die Spannungen $U(t)$ und die Ströme $I(t)$ zeitlich einen sinus- oder kosinusförmigen Verlauf:

$$U(t) = U_0 \cos(\omega t + \varphi_1) \quad , \quad I(t) = I_0 \cos(\omega t + \varphi_2) .$$

Wir gehen zu der komplexen Formulierung über und fassen sie als Realteile der komplexen Funktionen

$$\begin{aligned} \hat{U}(t) &= U_0 e^{i(\omega t + \varphi_1)} = U_0 e^{i\varphi_1} e^{i\omega t} = \hat{U}_0 e^{i\omega t} \\ \hat{I}(t) &= I_0 e^{i(\omega t + \varphi_2)} = I_0 e^{i\varphi_2} e^{i\omega t} = \hat{I}_0 e^{i\omega t} \end{aligned}$$

auf. Im Folgenden zeigen wir, dass sich das Ohmsche Gesetz auf die induktiven und kapazitiven Schaltelemente überträgt, wenn man diese komplexe Formulierung wählt.

(1) Ohmscher Widerstand R . Für einen Ohmschen Widerstand ist der Zusammenhang zwischen Spannung und Strom gegeben durch $U(t) = R I(t)$. Dieses Gesetz gilt auch für einen komplexen Wechselstrom $\hat{I}(t) = \hat{I}_0 e^{i\omega t}$.

$$\hookrightarrow \hat{U}(t) = R \hat{I}_0 e^{i\omega t} = R \hat{I}(t) .$$

Ein *Ohmscher Widerstand* wird durch den reellen Widerstand R beschrieben. Strom und Spannung sind in Phase.

(2) Kapazität C . Bei einem Kondensator mit Kapazität C besteht folgender Zusammenhang zwischen Ladung Q und angelegter Spannung U :

$$\boxed{Q = C \cdot U} \quad \hookrightarrow \quad I(t) = \frac{d}{dt} Q(t) = C \cdot \dot{U}(t) .$$

Speziell für $\hat{U}(t) = \hat{U}_0 e^{i\omega t}$ folgt

$$\hat{I}(t) = C \cdot \left(\hat{U}_0 e^{i\omega t} \right)' = C \cdot \hat{U}_0 e^{i\omega t} i\omega = C \cdot i\omega \hat{U}(t) .$$

Also ist der *komplexe Widerstand*

$$\hat{R}_C := \frac{\hat{U}(t)}{\hat{I}(t)} = \frac{1}{i\omega C} = -i \frac{1}{\omega C} .$$

Einer Kapazität wird der komplexe Widerstand $\hat{R}_C = \frac{1}{i\omega C}$ zugeordnet. Spannung und Strom sind um -90° verschoben.

(3) Induktivität L . Bei einer Spule mit Induktivität L ist der Zusammenhang zwischen Strom und induzierter Spannung durch das **Induktionsgesetz**

$$U(t) = L \frac{dI(t)}{dt}$$

gegeben. Speziell für $\hat{I}(t) = \hat{I}_0 e^{i\omega t}$ folgt

$$\hat{U}(t) = L \left(\hat{I}_0 e^{i\omega t} \right)' = L \hat{I}_0 e^{i\omega t} i\omega = i\omega L \hat{I}(t).$$

Einer Spule mit Induktivität L wird der *komplexe Widerstand*

$$\hat{R}_L := \frac{\hat{U}(t)}{\hat{I}(t)} = i\omega L$$

zugeordnet. $i\omega L$ liegt auf der positiven imaginären Achse. Die Phase zwischen Spannung und Strom beträgt $+90^\circ$; die Spannung eilt dem Strom um 90° voraus.

Bemerkung: Bei diesen Überlegungen wurde die Formel $(e^{i\omega t})' = i\omega e^{i\omega t}$ benutzt. Diese Gesetzmäßigkeit werden wir in Kap. 9.5.5 nachprüfen.

Zusammenfassung: Für RCL-Netzwerke gelten bei Wechselspannungen, $\hat{U}(t) = \hat{U}_0 e^{i\omega t}$, bzw. Wechselströmen, $\hat{I}(t) = \hat{I}_0 e^{i\omega t}$, Ohmsche Gesetze der Form $\hat{U}(t) = \hat{R} \hat{I}(t)$, wenn den einzelnen Schaltelementen **komplexe Widerstände (Impedanzen)** $\hat{R}$ zugeordnet werden:

Ohmscher Widerstand R	$\hat{R}_\Omega =$	R
Kapazität C	$\hat{R}_C =$	$\frac{1}{i\omega C}$
Induktivität L	$\hat{R}_L =$	$i\omega L$

Folgerung: Mit den Kirchhoffschen Regeln ergibt sich für die Ersatzschaltung zweier komplexer Widerstände $\hat{R}_1$ und $\hat{R}_2$ durch einen **komplexen Gesamtwiderstand (= Ersatzwiderstand)** $\hat{R}$:

(a) Reihenschaltung	$\hat{R} =$	$\hat{R}_1 + \hat{R}_2.$
(b) Parallelschaltung	$\frac{1}{\hat{R}} =$	$\frac{1}{\hat{R}_1} + \frac{1}{\hat{R}_2}$ bzw. $\hat{R} = \frac{\hat{R}_1 \hat{R}_2}{\hat{R}_1 + \hat{R}_2}.$

$\text{Re } \hat{R}$ heißt der *Wirkwiderstand*, $\text{Im } \hat{R}$ der *Blindwiderstand* und $|\hat{R}|$ der *reelle Scheinwiderstand*.

Im Wechselstromkreis dürfen also die bekannten Regeln für die Ersatzschaltung von Widerständen wie im Gleichstromkreis verwendet werden, wenn bei Kapazität und Induktivität zu komplexen Widerständen übergegangen wird!

➤ 5.3.4 Beispiele für RCL-Wechselstromschaltungen

Beispiel 5.16 (RCL-Reihenschaltung, mit MAPLE-Worksheet):

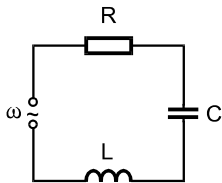


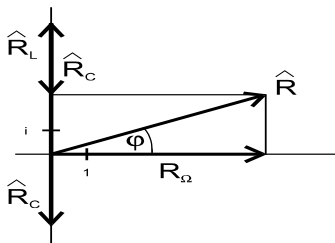
Abb. 5.8. RCL-Kreis

Nebenstehendes Bild zeigt eine Reihenschaltung aus je einem Ohmschen Widerstand R_Ω , einer Kapazität C und einer Induktivität L . Es addieren sich die komplexen Einzelwiderstände zum komplexen Gesamtwiderstand

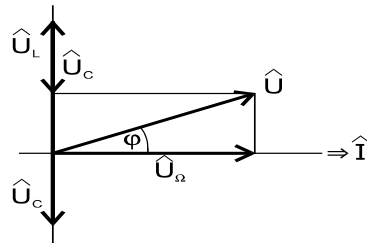
$$\hat{R} = R_\Omega + \hat{R}_C + \hat{R}_L = R_\Omega + \frac{1}{i\omega C} + i\omega L$$

$$\hat{R} = R_\Omega + i\left(\omega L - \frac{1}{\omega C}\right).$$

Die Addition ist graphisch durch das Zeigerdiagramm gegeben.



Zeigerdiagramm



Spannungsdiagramm

Der *Blindwiderstand* ist

$$\text{Im } \hat{R} = \omega L - \frac{1}{\omega C},$$

der *Wirkwiderstand* ist

$$\text{Re } \hat{R} = R_\Omega$$

und der **reelle Scheinwiderstand**

$$R = |\hat{R}| = \sqrt{R_\Omega^2 + \left(\omega L - \frac{1}{\omega C}\right)^2}.$$

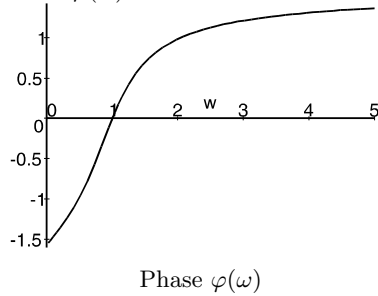
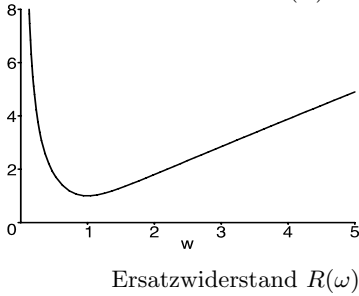
Die **Phase** zwischen Spannung und Strom erhält man aus

$$\tan \varphi = \frac{\text{Im } \hat{R}}{\text{Re } \hat{R}} = \frac{\omega L - \frac{1}{\omega C}}{R_\Omega}.$$

Diskussion: Multipliziert man die Widerstände jeweils mit $\hat{I}$, erhält man das zugehörige Spannungsdiagramm:

- (1) U_Ω fällt am Ohmschen Widerstand ab und ist mit dem Strom I in Phase.
- (2) U_L fällt an der Induktivität ab. U_L eilt dem Strom um 90° voraus.
- (3) U_C fällt an der Kapazität ab. U_C hinkt dem Strom um 90° nach.

Für $R = 1$, $L = 1$ und $C = 1$ erhält man die folgende graphische Darstellung für den Ersatzwiderstand $R(\omega)$ bzw. die Phase $\varphi(\omega)$:



□

Beispiel 5.17 (LC-Parallelkreis, mit [MAPLE-Worksheet](#)):

Für die in Abbildung 5.9 gezeichnete Schaltung berechnet man den komplexen Ersatzwiderstand, indem zuerst L und R_2 ersetzt werden durch den Reihenersatzwiderstand $R_r = R_2 + i\omega L$. R_r liegt parallel zu C , so dass sich die Leitwerte addieren

$$Z_p = i\omega C + \frac{1}{i\omega L + R_2}.$$

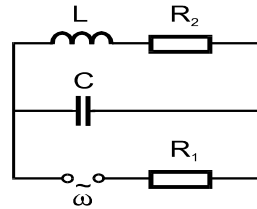


Abb. 5.9. LC-Parallelkreis

Der komplexe Gesamtwiderstand setzt sich nun zusammen aus der Summe von R_1 und $R_p = \frac{1}{Z_p}$:

$$R_{ges} = R_1 + \frac{1}{Z_p} = R_1 + \frac{1}{i\omega C + \frac{1}{i\omega L + R_2}}$$

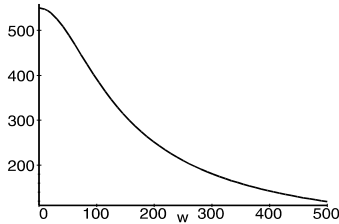
$$R_{ges} = \frac{R_1 \omega^2 C L - R_1 \omega C R_2 i - R_1 - \omega L i - R_2}{\omega^2 C L - \omega C R_2 i - 1}.$$

Man erkennt in dieser Darstellung, dass der Gesamtwiderstand eine komplexe rationale Funktion in ω ist und 2 der höchste auftretende Exponent. Dies spiegelt die Tatsache wider, dass der Schaltkreis zwei Energiespeicher, nämlich C und L besitzt. Für die Werte $C = 20 \cdot 10^{-6}$, $L = 20 \cdot 10^{-3}$, $R_1 = 50$, $R_2 = 500$

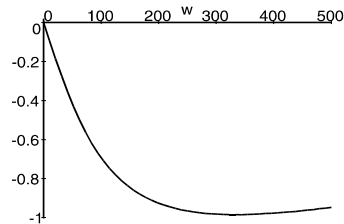
ergibt sich der Gesamtwiderstand als Funktion in ω

$$R_{ges} = \frac{0.8000 \cdot 10^{-11} \omega^4 + 0.004960 \omega^2 + 550}{0.1600 \cdot 10^{-12} \omega^4 + 0.00009920 \omega^2 + 1} + \frac{(-0.8000 \cdot 10^{-8} \omega^3 - 4.980 \omega) i}{0.1600 \cdot 10^{-12} \omega^4 + 0.00009920 \omega^2 + 1}.$$

Die Kurvenverläufe von Gesamtwiderstand und Phase in Abhängigkeit von ω sind gegeben durch



Gesamtwiderstand $R(\omega)$



Phase $\varphi(\omega)$

□

MAPLE-Worksheets zu Kapitel 5



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 5 mit MAPLE zur Verfügung. Sie können direkt durch Anklicken aus der pdf-Version des Buches gestartet werden.

- Darstellung komplexer Zahlen mit MAPLE
- Komplexes Rechnen mit MAPLE
- Visualisierung der komplexen Rechenoperationen
- Überlagerung von Schwingungen
- RCL-Wechselstromkreise mit MAPLE
- Übertragungsverhalten von Filterschaltungen
- MAPLE-Lösungen zu den Aufgaben

5.4 Aufgaben zu komplexen Zahlen

- 5.1 Geben Sie die Exponentialform der folgenden komplexen Zahlen an
 a) $3\sqrt{3} + 3i$ b) $-2 - 2i$ c) $1 - \sqrt{3}i$ d) 5 e) $-5i$ f) -1
- 5.2 Wie lautet die trigonometrische und algebraische Normalform von
 a) $3\sqrt{2}e^{i\frac{\pi}{4}}$ b) $2e^{i\frac{2\pi}{3}}$ c) $e^{i\pi}$ d) $4e^{i\frac{4\pi}{3}}$
- 5.3 Welches sind die zugehörigen komplex konjugierten Zahlen
 a) $3 + \sqrt{2}i$ b) $4(\cos 125^\circ + i \sin 125^\circ)$ c) $5e^{i\frac{3}{2}\pi}$ d) $\sqrt{3}e^{i0.734}$
- 5.4 Man bestimme die trigonometrische Normalform von
 a) $-1 + \sqrt{3}i$ b) $-1 + i$ c) $\sqrt{2} + \sqrt{2}i$ d) $-3 - 4i$
- 5.5 Berechnen Sie
 a) $2(5 - 3i) - 3(-2 + i) + 5(i - 3)$ b) $(3 - 2i)^3$ c) $\frac{5}{3 - 4i} + \frac{10}{4 + 3i}$
 d) $\left(\frac{1 - i}{1 + i}\right)^{10}$ e) $\left|\frac{2 - 4i}{5 - 7i}\right|^2$ f) $\frac{(1 + i)(2 + 3i)(4 - 2i)}{(1 + 2i)^2(1 - i)}$
- 5.6 Sei $z_1 = 1 - i$, $z_2 = -2 + 4i$, $z_3 = \sqrt{3} - 2i$. Wie lautet die algebraische Normalform von
 a) $z_1^2 + 2z_1 - 3$ b) $|2z_2 - 3z_1|^2$ c) $(z_3 - z_3^*)^5$
 d) $|z_1 z_2^* + z_2 z_1^*|$ e) $\left|\frac{z_1 + z_2 + 1}{z_1 - z_2 + i}\right|$ f) $\frac{1}{2} \left(\frac{z_3}{z_3^*} + \frac{z_3^*}{z_3}\right)$
 g) $((z_2 + z_3)(z_1 - z_3))^*$ h) $|z_1^2 + z_2^*|^2 + |z_3^* - z_2|^2$ i) $\operatorname{Im} \left\{ \frac{z_1 z_2}{z_3} \right\}$
- 5.7 Berechnen Sie
 a) $(-1 + \sqrt{3}i)^{10}$ b) $[2(\cos 45^\circ + i \sin 45^\circ)]^3$ c) $(3\sqrt{3} + 3i)^6$ d) $\left(2e^{i\frac{5}{3}\pi}\right)^7$
- 5.8 Geben Sie im Komplexen alle Lösungen an von
 a) $z^4 + 81 = 0$,
 b) $z^6 + 1 = \sqrt{3}i$
- 5.9 Bestimmen Sie alle komplexen Lösungen von
 a) $z^5 - 2z^4 - z^3 + 6z - 4 = 0$
 b) $4x^4 + 4x^3 - 7x^2 + x - 2 = 0$
- 5.10 Lösen Sie Aufgaben 5.1 - 5.9 mit MAPLE.
- 5.11 Wie lauten der Real- und Imaginärteil der folgenden komplexen Zahlen
 a) $\frac{-2 + 7i}{15i}$ b) $\frac{1 + i}{1 - i}$ c) $\frac{1 - i}{1 + 2i} - \frac{1 + 3i}{1 - 2i}$ d) $\frac{2e^{i\frac{\pi}{4}}}{(1 + i)(2 + i)}$ e) $2e^{i120^\circ}$
 f) $3e^{i\frac{5\pi}{6}}$ g) $-5e^{-i\frac{\pi}{2}}$ h) $7e^{i\pi}$ i) $\frac{2 - i}{2 + i} \cdot e^{-i\frac{\pi}{3}}$
 Wie groß sind jeweils Betrag und Winkel?
- 5.12 Wie heißen die folgenden komplexen Zahlen in Exponentialform? (Verwenden Sie zur Berechnung MAPLE.)
 a) $-1 - i$ b) $-1 + i$ c) $3 + 4i$ d) $-3 - 4i$ e) $2i$ f) -2 g) $1 - 2i$

5.13 Es sei $z = x + iy$ und z^* die zu z konjugiert komplexe Zahl. Bestimmen Sie mit MAPLE

a) $a = \left| \frac{z}{z^*} \right|$ b) $b = \operatorname{Re} \{z^{-2}\}$ c) $c = \operatorname{Im} \{z^{*3}\}$ d) $d = \operatorname{Im} \{(z^3)^*\}$

5.14 Berechnen Sie mit MAPLE

a) $\left(\frac{3+4i}{5}\right)^{10}$ b) $\left(i + \frac{1}{1+i}\right)^6$ c) $\left[(1+i) \cdot e^{-i\frac{\pi}{6}}\right]^9$

5.15 Berechnen Sie mit MAPLE alle reellen und komplexen Lösungen der Gleichungen

a) $z^3 = i$ b) $z^2 = -1 + i\sqrt{3}$ c) $32z^5 - 243 = 0$
 d) $z^3 + \frac{4}{1+i} = 0$ e) $z^4 + \frac{1+2e^{i\frac{\pi}{2}}}{2+e^{-i\frac{\pi}{2}}} = 0$ f) $z^2 - 2iz + 3 = 0$

5.16 Bestimmen Sie mit MAPLE alle Nullstellen der Funktion $z^4 - 3z^3 + 2z^2 + 2z - 4$.

5.17 a) Berechnen Sie den komplexen und reellen Scheinwiderstand für die in Abb. 1a skizzierte Reihenschaltung ($R = 100\Omega$, $C = 20\mu F$, $L = 0.2H$, $\omega = 10^6 \frac{1}{s}$).

b) Bestimmen Sie den komplexen und reellen Scheinwiderstand für die in Abb. 1b skizzierte Parallelschaltung ($R = 100\Omega$, $L = 0.5H$, $\omega = 500 \frac{1}{s}$).

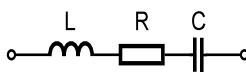


Abb. 1a

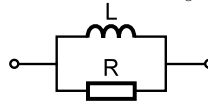


Abb. 1b

5.18 a) Man berechne den komplexen Scheinwiderstand der in Abb. 2a dargestellten Schaltung als Funktion von ω .

b) Man berechne den komplexen Scheinwiderstand der in Abb. 2b dargestellten Schaltung bei einer Kreisfrequenz $\omega = 300 s^{-1}$ für die Parameter $R_1 = 50\Omega$, $L_1 = 1H$, $R_2 = 300\Omega$, $C_1 = 10\mu F$, $R_3 = 20\Omega$, $L_2 = 1.5H$.

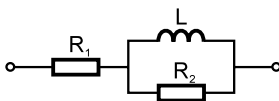


Abb. 2a

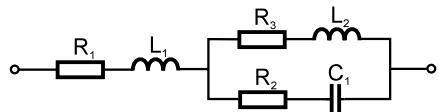


Abb. 2b

5.19 Gegeben sind die beiden Wechselspannungen $u_1(t)$ und $u_2(t)$. Man bestimme die durch Superposition entstehende resultierende Wechselspannung ($\omega = 314 \frac{1}{s}$):

$$u_1(t) = 100 V \cdot \sin(\omega t) \quad u_2(t) = 150 V \cdot \cos\left(\omega t - \frac{\pi}{4}\right)$$

und zeichne alle drei Graphen in ein Schaubild.

5.20 Die mechanischen Schwingungen $y_1(t) = 20 cm \cdot \sin\left(\pi t + \frac{\pi}{10}\right)$ und $y_2(t) = 15 cm \cdot \cos\left(\pi t + \frac{\pi}{6}\right)$ werden ungestört zur Überlagerung gebracht. Wie lautet die resultierende Schwingung? (Man rechne in der Kosinusdarstellung!)

5.21 Man zeige zeichnerisch, dass

$$3 \cos\left(\omega t + \frac{\pi}{6}\right) + 2 \cos\left(\omega t + \frac{\pi}{4}\right) = A \cos(\omega t + \varphi)$$

mit $A \approx 5$, $\varphi \approx 36^\circ$.

The background of the page is a light gray with several large, overlapping circles of varying sizes. A large, stylized number '6' is positioned in the upper left quadrant, partially overlapping the circles. The text 'Kapitel 6' and 'Grenzwert und Stetigkeit' is centered in the upper half of the page.

Kapitel 6
Grenzwert und Stetigkeit

6

6

6	Grenzwert und Stetigkeit	221
6.1	Reelle Zahlenfolgen	223
6.2	Funktionsgrenzwert	229
6.3	Stetigkeit einer Funktion	235
6.4	Intervallhalbierungs-Methode	237
6.5	Aufgaben zu Grenzwert und Stetigkeit	240

Zusätzliche Abschnitte auf der CD-Rom

6.6	MAPLE: Folgen und Grenzwerte	cd
6.6.1	Ermittlung von Grenzwerten mit MAPLE	cd
6.6.2	Graphische Darstellung von Funktionsfolgen mit MAPLE...	cd
6.6.3	Berechnung von Funktionsgrenzwerten mit MAPLE	cd
6.6.4	Bisektionsverfahren mit MAPLE	cd

6 Grenzwert und Stetigkeit

Grundlegend für das gesamte Kapitel sind Grenzwerte von Zahlenfolgen. Auf dieser Grundlage baut die Konstruktion des Funktionsgrenzwertes auf, der wiederum für den Begriff der Stetigkeit benötigt wird. Etwas lax formuliert sind die stetigen Funktionen die Funktionen, die bei einem zusammenhängenden Definitionsbereich keine Sprungstelle aufweisen, d.h. ohne Unterbrechung gezeichnet werden können. Für diese stetigen Funktionen werden wir das Bisektionsverfahren einführen, um numerisch die Nullstellen dieser Funktionen zu bestimmen.

6.1 Reelle Zahlenfolgen

6.1

In den meisten Tests zur Erfassung der Denkfähigkeit von Schülern und Studenten kommt eine Aufgabenstellung der Form vor: Gegeben ist

$$-\frac{1}{2}, \frac{1}{4}, -\frac{1}{6}, \frac{1}{8}, \dots$$

Man möge vier weitere Glieder dieser Folge angeben. Gemeint ist natürlich $-\frac{1}{10}, \frac{1}{12}, -\frac{1}{14}, \frac{1}{16}$. Etwas schwieriger ist die Aufgabe, eine Gesetzmäßigkeit zu finden, um das 100. Glied der Folge zu bestimmen. Hierbei ist dann nach der Formel $(-1)^n \frac{1}{2n}$ gefragt, in die man anschließend $n = 100$ einsetzt. Eventuell wird auch gefragt, welchem Wert die Folgenglieder für große n beliebig nahe kommen. Dann ist der Grenzwert der Folge gesucht. Wir definieren verallgemeinernd:

Definition: (Zahlenfolge). Unter einer **reellen Zahlenfolge** versteht man eine geordnete Menge reeller Zahlen, indiziert mit $1, 2, 3, \dots$.

Notation: $(a_n)_n = a_1, a_2, a_3, \dots, a_n, \dots$.

Die Zahlen $a_1, a_2, \dots$ heißen *Glieder* der Folge, a_n das *n-te Glied* bzw. das *allgemeine Glied* der Folge (*Bildungsgesetz*).

Beispiele 6.1:

① $(a_n)_n = 1, 2, 3, 4, \dots, n, \dots;$

$$a_n = n.$$

② $(a_n)_n = 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots, \frac{1}{n}, \dots;$

$$a_n = \frac{1}{n}.$$

③ $(a_n)_n = -1, +1, -1, +1, -1, +1, \dots;$

$$a_n = (-1)^n.$$

④ $(a_n)_n = -\frac{1}{2}, \frac{1}{4}, -\frac{1}{6}, \frac{1}{8}, -\frac{1}{10}, \dots;$

$$a_n = (-1)^n \frac{1}{2n}.$$

⑤ $(a_n)_n = 0.1, 0.11, 0.111, 0.1111, \dots;$

$$a_1 = 0.1 \text{ und}$$

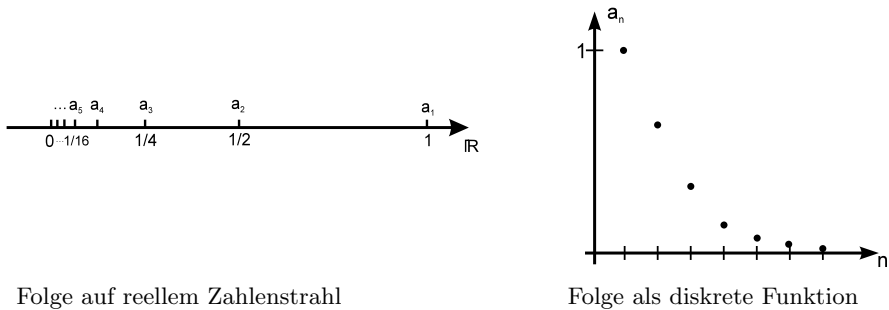
$$a_n = a_{n-1} + 10^{-n} \text{ für } n \geq 2.$$

Eine Zahlenfolge kann als **diskrete Funktion** $F : \mathbb{N} \rightarrow \mathbb{R}$ mit $n \mapsto a(n) = a_n$ aufgefasst werden. Die Funktion F ordnet jeder natürlichen Zahl n genau eine reelle Zahl $F(n) = a_n$ zu. Hierbei tritt die Variable n als Index auf; die Funktionswerte sind nummeriert.

Darstellung von Zahlenfolgen: Die Glieder einer Folge sind darstellbar auf der reellen *Zahlengeraden*. Z.B. für $a_n = \frac{1}{2^{n-1}}$ erhalten wir die Folge

$$(a_n)_n = 1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \frac{1}{32}, \dots$$

die unten auf dem Zahlenstrahl dargestellt ist. Gemäß der Interpretation als diskrete Funktion können Folgen auch über den Funktionsgraphen (siehe rechte Abb.) dargestellt werden. Da die Funktion nur für $n \in \mathbb{N}$ definiert ist, dürfen die Punkte nicht verbunden werden!



⑤ **Grenzwert einer Folge**

Um das Verhalten der Folge $a_n = 1 - \frac{1}{n}$ ($n \in \mathbb{N}$) für große n zu diskutieren, erstellen wir eine Wertetabelle

n	1	2	3	4	...	10	...	100	...	1000	...	10000
a_n	0	$\frac{1}{2}$	$\frac{2}{3}$	$\frac{3}{4}$	...	0.9	...	0.99	...	0.999	...	0.9999

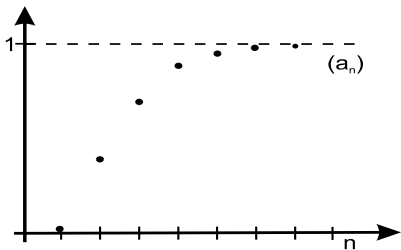


Abb. 6.1. Grenzwert der Folge $1 - \frac{1}{n}$

Die Eigenschaften dieser Folge sind, dass alle Glieder a_n kleiner als 1 sind und dass mit wachsendem n die Glieder a_n sich an die Zahl 1 annähern. Damit wird der Abstand zwischen den Folgengliedern a_n und dem Wert 1 mit wachsendem n kleiner:

$$|a_n - 1| \rightarrow 0 \quad \text{für } n \rightarrow \infty.$$

Die Zahl 1 wird als *Grenzwert* der Folge $a_n = 1 - \frac{1}{n}$ bezeichnet.



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich ein **MAPLE-Worksheet**, bei dem man selbst Folgen spezifiziert und diese Folgen dann - sofern sie einen Grenzwert besitzt - in Form einer Animation dargestellt werden.

Definition: (Grenzwert)

- (1) Eine reelle Zahl a heißt **Grenzwert** oder **Limes** der Zahlenfolge $(a_n)_{n \in \mathbb{N}}$, wenn es zu jedem $\varepsilon > 0$ eine natürliche Zahl n_0 gibt, so dass für alle $n > n_0$ stets gilt

$$|a_n - a| < \varepsilon.$$

- (2) Eine Folge heißt **konvergent**, wenn sie einen Grenzwert besitzt. Wir verwenden dann die Notation

$$a_n \xrightarrow{n \rightarrow \infty} a \quad \text{oder} \quad \lim_{n \rightarrow \infty} a_n = a.$$

- (3) Eine Folge heißt **divergent**, wenn sie keinen Grenzwert besitzt.

Definition (1) besagt, dass a der Grenzwert einer Folge ist, wenn der Abstand von Folgengliedern zum Grenzwert, $|a_n - a|$, beliebig klein (ε) gewählt werden kann und alle Folgenglieder a_n ab n_0 einen noch kleineren Abstand zum Grenzwert a besitzen. Anschaulich formuliert bedeutet dies:

Folgerung: Eine Folge a_n konvergiert, wenn es einen Grenzwert a gibt, so dass der Abstand

$$d = |a_n - a| \rightarrow 0 \quad \text{für} \quad n \rightarrow \infty.$$

Bemerkungen:

- (1) Konvergiert eine Zahlenfolge gegen einen Grenzwert a , dann hängt der Index n_0 von der Wahl des Abstandes ε ab.
- (2) Der Grenzwert einer Zahlenfolge ist eindeutig.
- (3) Divergiert eine Folge, so muss nicht notwendigerweise $a_n \rightarrow \pm\infty$ gelten.
- (4) Wir werden Sätze kennen lernen, mit denen man den Grenzwert einer Folge direkt berechnen kann, ohne auf die obige Definition zurückgreifen zu müssen.
- (5) Der Grenzwert einer Folge wird in der Regel nie von den Folgengliedern erreicht.

Beispiele 6.2:

① Die Folge

$$(a_n)_n = \left(\frac{1}{n}\right)_n = 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots$$

konvergiert gegen den Grenzwert 0 : $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} \frac{1}{n} = 0.$

Man bezeichnet Folgen, die gegen den Grenzwert 0 konvergieren, als **Nullfolgen**.

② Die Folge

$$(a_n)_n = \left(1 + \frac{1}{2^n}\right)_n = 1.5, 1.25, 1.125, 1.0625, 1.03125, \dots$$

konvergiert gegen 1, da für den Abstand der Folgenglieder a_n zu 1 gilt

$$d = |a_n - a| = \left|1 + \frac{1}{2^n} - 1\right| = \frac{1}{2^n} \xrightarrow{n \rightarrow \infty} 0.$$

③ Die Folge

$$(a_n)_n = (n)_n = 1, 2, 3, 4, \dots$$

ist unbeschränkt wachsend und daher divergent.

④ $\triangle$ Die Folge

$$(a_n)_n = ((-1)^n)_n = -1, 1, -1, 1, -1, \dots$$

hat zwei sog. *Häufungspunkte*, nämlich 1 und -1 . Sie konvergiert aber nicht gegen **einen** Grenzwert. Daher ist sie divergent.

⑤ $x \in \mathbb{R}$ fest.

$$(a_n)_n = (x^n)_n = x, x^2, x^3, x^4, x^5, \dots, x^n, \dots$$

Für x -Werte mit $-1 < x < 1$ konvergiert die Folge gegen Null, für $x = 1$ gegen 1, für andere x -Werte divergiert die Folge. $\square$

Um nachzuprüfen, dass Zahlenfolgen konvergent sind, muss nach der Definition von Konvergenz der Grenzwert bereits bekannt sein, da der Abstand $d = |a_n - a|$ bestimmt werden muss. Das folgende *Monotoniekriterium* macht eine Aussage über die Konvergenz einer Folge, ohne dass der Grenzwert bekannt ist. Es besagt, dass eine monoton wachsende Folge, die nach oben hin beschränkt ist, stets einen Grenzwert besitzt. Eine Folge, die monoton fällt und nach unten beschränkt ist, besitzt ebenfalls einen Grenzwert.

Monotonie-Kriterium:

(1) Sei $(a_n)_n$ eine Folge mit den Eigenschaften

$$(i) \ a_n \leq a_{n+1} \text{ (Monotonie),} \quad (ii) \ a_n \leq A \text{ (Beschränktheit),}$$

dann konvergiert die Folge gegen einen Grenzwert $a \leq A$.

(2) Sei $(a_n)_n$ eine Folge mit den Eigenschaften

$$(i) \ a_n \geq a_{n+1} \text{ (Monotonie),} \quad (ii) \ a_n \geq A \text{ (Beschränktheit),}$$

dann konvergiert die Folge gegen einen Grenzwert $a \geq A$.

Beispiel 6.3 (Exponentialfolge, mit [MAPLE-Worksheet](#)): Die Folge

$$a_n = \left(1 + \frac{1}{n}\right)^n$$

ist konvergent, da sie eine monoton wachsende Folge darstellt, die nach oben durch 3 beschränkt ist (ohne Beweis):

n	1	10^1	10^2	10^3	10^4	10^5
a_n	2	2.59374	2.70481	2.71692	2.71814	2.71826

Der Grenzwert e heißt *Eulersche Zahl*

$$e = 2.71828\ 18284\ 59045\ 23536\ 0287 \dots$$

Wir stellen den Grenzwert zusammen mit einer ε -Umgebung als Funktions-schaubild für die Folge $(1 + \frac{1}{n})^n$ graphisch dar.

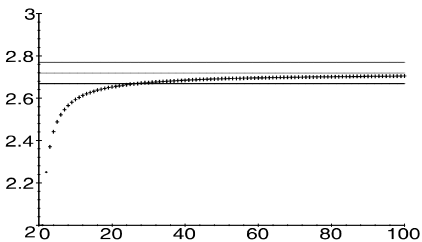


Abb. 6.2. Zum Grenzwert der Folge $(1 + \frac{1}{n})^n$

Die Eulersche Zahl e tritt in vielen naturwissenschaftlichen Zusammenhängen auf. Aus mathematischer Sicht ist sie eine der bedeutsamsten reellen Zahlen. Die Exponentialfunktion basiert auf e als Basis. Wir werden im Kapitel Taylor-Reihen in [Beispiel 9.33](#) eine alternative Methode kennen lernen, um die Zahl

e durch eine schneller konvergente Folge zu bestimmen:

$$e = \sum_{n=0}^{\infty} \frac{1}{n!} = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!} + \dots + \frac{1}{n!} + \dots \quad \square$$

Beispiel 6.4 (Babylonisches Wurzelziehen, mit [MAPLE-Worksheet](#)):

Die rekursiv definierte Folge

$$a_0 = a \quad , \quad a_{n+1} = \frac{1}{2} \left(a_n + \frac{a}{a_n} \right) \quad (*)$$

ist für jedes $a > 0$ eine monoton fallende Folge, die nach unten durch $\sqrt{a}$ beschränkt ist (ohne Beweis).

Der Grenzwert der Folge bestimmt sich aus der Definitionsgleichung von a_n (*), indem auf beiden Seiten der Gleichung der Limes $n \rightarrow \infty$ gebildet wird. Sei der Grenzwert der Folge $b := \lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} a_{n+1}$, so folgt für b mit den Limesrechenregeln

$$\begin{aligned} \lim_{n \rightarrow \infty} a_{n+1} &= \lim_{n \rightarrow \infty} \frac{1}{2} \left(a_n + \frac{a}{a_n} \right) \\ \Rightarrow b &= \frac{1}{2} \left(b + \frac{a}{b} \right). \end{aligned}$$

Löst man diese Gleichung nach b auf, folgt

$$b^2 = a \quad \Rightarrow \quad \boxed{b = \sqrt{a}.}$$

Somit stellt obige Folge ein Näherungsverfahren zur Berechnung von Quadratwurzeln dar, das schon den Babyloniern bekannt war. Tatsächlich ist dies ein Spezialfall des Newton-Verfahrens, das wir in 7.9 einführen werden.

Die folgende Wertetabelle verdeutlicht die schnelle Konvergenz der Folge für $a = 2$:

n	1	2	3	4	5
a_n	1.5	1.416666666	1.414215686	1.414213562	1.414213562

Nach 4 Iterationen ist $\sqrt{2} = 1.414213562$ bis auf 9 Stellen genau berechnet! $\square$

Das Monotonie-Kriterium sichert zwar die Konvergenz einer Folge, aber es liefert nicht den Grenzwert. Die Limesrechenregeln bei Folgen bieten eine Möglichkeit, den Grenzwert einer Folge für viele aber nicht alle Fälle zu berechnen:

Limesrechenregeln bei Folgen:

Seien $(a_n)_n$ und $(b_n)_n$ konvergente Folgen mit $\lim_{n \rightarrow \infty} a_n = a$ und $\lim_{n \rightarrow \infty} b_n = b$. Sei $c \in \mathbb{R}$. Dann gilt

$$(L1) \quad \lim_{n \rightarrow \infty} c a_n = c \lim_{n \rightarrow \infty} a_n = c \cdot a$$

$$(L2) \quad \lim_{n \rightarrow \infty} (a_n \pm b_n) = \lim_{n \rightarrow \infty} a_n \pm \lim_{n \rightarrow \infty} b_n = a \pm b$$

$$(L3) \quad \lim_{n \rightarrow \infty} (a_n \cdot b_n) = \lim_{n \rightarrow \infty} a_n \cdot \lim_{n \rightarrow \infty} b_n = a \cdot b$$

$$(L4) \quad \lim_{n \rightarrow \infty} \left(\frac{a_n}{b_n} \right) = \lim_{n \rightarrow \infty} a_n / \lim_{n \rightarrow \infty} b_n = \frac{a}{b}, \quad \text{falls } b_n, b \neq 0.$$

Beispiele 6.5 (Ermittlung von Grenzwerten):

$$\textcircled{1} \quad a_n = \frac{4n^3 - 6}{6n^3 + 2n^2} = \frac{4n^3 - 6}{6n^3 + 2n^2} \cdot \frac{\frac{1}{n^3}}{\frac{1}{n^3}} = \frac{4 - \frac{6}{n^3}}{6 + \frac{2}{n}} \xrightarrow{n \rightarrow \infty} \frac{4}{6} = \frac{2}{3}.$$

$$\textcircled{2} \quad a_n = \frac{n-1}{2n^2+1} = \frac{n-1}{2n^2+1} \cdot \frac{\frac{1}{n^2}}{\frac{1}{n^2}} = \frac{\frac{1}{n} - \frac{1}{n^2}}{2 + \frac{1}{n^2}} \xrightarrow{n \rightarrow \infty} \frac{0}{2} = 0.$$

$$\textcircled{3} \quad a_n = \frac{3^{n+1} + 2^n}{3^n + 1} = \frac{3^{n+1} + 2^n}{3^n + 1} \cdot \frac{\frac{1}{3^n}}{\frac{1}{3^n}} = \frac{3 + \left(\frac{2}{3}\right)^n}{1 + \left(\frac{1}{3}\right)^n} \xrightarrow{n \rightarrow \infty} \frac{3}{1} = 3,$$

da $\left(\frac{2}{3}\right)^n \rightarrow 0$ und $\left(\frac{1}{3}\right)^n \rightarrow 0$ für $n \rightarrow \infty$ nach Beispiel 6.2 ⑤.

Man beachte, dass die Umformungen notwendig sind, da die Limesrechenregeln nur für **konvergente** Folgen gelten. Für alle Beispiele wird der Quotient mit dem Kehrwert des größten Terms erweitert. $\square$

6.2 Funktionsgrenzwert**6.2**

In Abschnitt 6.1 werden Grenzwerte von Zahlenfolgen $(x_n)_{n \in \mathbb{N}}$ untersucht. Dieser Begriff wird nun direkt auf *Funktionsgrenzwerte* ausgedehnt, indem Folgen der Form $(f(x_n))_{n \in \mathbb{N}}$ betrachtet werden. Zur Einführung untersuchen wir das Verhalten der Funktion $f(x) = x^2$ an der Stelle $x_0 = 2$. Dazu wählen wir die Folge

$$(x_n)_n = 1.9, 1.99, 1.999, 1.9999, \dots \xrightarrow{n \rightarrow \infty} 2$$

und berechnen zu jedem Folgenglied den Funktionswert

$$(f(x_n))_n = 3.61, 3.9601, 3.996, 3.9996, \dots \xrightarrow{n \rightarrow \infty} 4.$$

Die **Folge der Funktionswerte** konvergiert gegen den Wert 4.

Um den Funktionsgrenzwert zu gegebener Funktion f an einer Stelle x_0 zu erhalten, wählt man sich eine Zahlenfolge $x_n \xrightarrow{n \rightarrow \infty} x_0$ aus dem Definitionsbereich von f und wendet die Funktion f auf x_n an. Dann untersucht man die Konvergenzeigenschaften der Folge $(f(x_n))_n$ (= Grenzwertuntersuchung der Funktion an der Stelle x_0). In unserem Beispiel gilt auch für jede andere Folge $(x_n)_n$, die gegen den Wert 2 konvergiert, dass $f(x_n) \xrightarrow{n \rightarrow \infty} 4$. Man schreibt daher:

$$\lim_{n \rightarrow \infty} f(x_n) = \lim_{\substack{x \rightarrow 2 \\ (x < 2)}} f(x) = \lim_{\substack{x \rightarrow 2 \\ (x < 2)}} x^2 = 4.$$

Da die Folgenglieder $x < 2$, nennt man diesen Grenzwert den *linksseitigen Grenzwert* von $f(x) = x^2$ an der Stelle $x_0 = 2$.

Beispiel 6.6 (Mit MAPLE-Worksheet). Man kann diesen Sachverhalt anschaulich darstellen, indem sowohl die Folge $(x_n)_n$ als auch die Funktionsfolge $(f(x_n))_n$ in ein Schaubild gezeichnet werden. Zur übersichtlicheren Darstellung wählen wir nun die Folge $x_n = 2 - \frac{1}{n^2} \xrightarrow{n \rightarrow \infty} 2$:

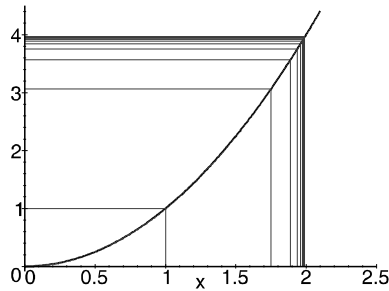


Abb. 6.3. Linksseitiger Funktionsgrenzwert bei $x_0 = 2$

Man erkennt, dass die x_n -Werte sich der Zahl 2 von links annähern; die Funktionswerte $f(x_n)$ dem y -Wert 4. $\square$

Analog erhält man den *rechtsseitigen* Grenzwert der Funktion bei $x_0 = 2$, indem man als Zahlenfolge z.B.

$$(x_n)_n = 2.1, 2.01, 2.001, 2.0001, \dots \rightarrow 2$$

wählt. Dazu ist die zugehörige Funktionsfolge

$$(f(x_n))_n = 4.41, 4.041, 4.004, 4.0004, \dots \rightarrow 4.$$

Auch hier gilt allgemeiner, dass der Funktionsgrenzwert unabhängig von der gewählten Zahlenfolge x_n ist. Man schreibt für den rechtsseitigen Grenzwert

$$\lim_{n \rightarrow \infty} f(x_n) = \lim_{\substack{x \rightarrow 2 \\ (x > 2)}} f(x) = \lim_{\substack{x \rightarrow 2 \\ (x > 2)}} x^2 = 4.$$

Für die Funktion $f(x) = x^2$ existieren also sowohl der linksseitige als auch der rechtsseitige Grenzwert der Funktion und beide sind gleich 4.

Definition: (Funktionsgrenzwert). Eine Funktion f sei in einer Umgebung von x_0 definiert. Gilt für **jede** im Definitionsbereich der Funktion liegende Folge $(x_n)_n$, die gegen x_0 konvergiert, stets

$$\lim_{n \rightarrow \infty} f(x_n) = g \in \mathbb{R},$$

so heißt g der **Grenzwert** von $f(x)$ für $x_n \xrightarrow{n \rightarrow \infty} x_0$.

Schreibweise: $\lim_{n \rightarrow \infty} f(x_n) = \lim_{x \rightarrow x_0} f(x) = g$, wenn $x_n \xrightarrow{n \rightarrow \infty} x_0$.

Bemerkungen:

- (1) Es wird **nicht** gefordert, dass x_0 aus dem Definitionsbereich der Funktion ist.
- (2) Der Grenzübergang $x \rightarrow x_0$ bedeutet, dass x der Stelle x_0 beliebig nahe kommt, **ohne** den Wert x_0 anzunehmen!
- (3) Es kann der Fall eintreten, dass, obwohl $x_0 \notin \mathbb{D}$, der Funktionsgrenzwert existiert, d.h. der linksseitige mit dem rechtsseitigen Grenzwert übereinstimmt.
- (4) Der *linksseitige Grenzwert* wird auch oftmals bezeichnet mit

$$g_l := \lim_{\substack{x \rightarrow x_0 \\ (x < x_0)}} f(x) = \lim_{h \rightarrow 0} f(x_0 - h)$$

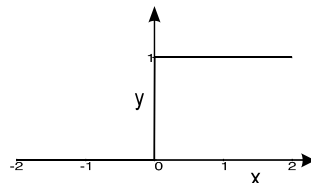
und der *rechtsseitige Grenzwert* mit

$$g_r := \lim_{\substack{x \rightarrow x_0 \\ (x > x_0)}} f(x) = \lim_{h \rightarrow 0} f(x_0 + h).$$

Beispiele 6.7:

- ① Die **Heaviside-Funktion**:

$$f: \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f(x) = \begin{cases} 0 & \text{für } x < 0 \\ 1 & \text{für } x \geq 0 \end{cases}$$



Die Heaviside-Funktion ist die Funktion, die für negative x -Werte Null und für positive x -Werte den Funktionswert 1 besitzt. Sie wird in den Anwen-

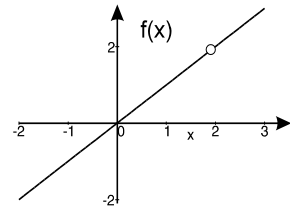
dungen auch oftmals mit Sprungfunktion $S(x)$ bzw. als Einschaltfunktion bezeichnet. Die Heaviside-Funktion besitzt bei $x_0 = 0$ **keinen** Grenzwert, da der rechtsseitige Grenzwert nicht mit dem linksseitigen übereinstimmt:

$$g_l = \lim_{h \rightarrow 0} f(x_0 - h) = \lim_{h \rightarrow 0} f(-h) = \lim_{h \rightarrow 0} 0 = 0,$$

$$g_r = \lim_{h \rightarrow 0} f(x_0 + h) = \lim_{h \rightarrow 0} f(h) = \lim_{h \rightarrow 0} 1 = 1.$$

② Für die Funktion

$$f : \mathbb{R} \setminus \{2\} \rightarrow \mathbb{R} \text{ mit } f(x) = \frac{x^2 - 2x}{x - 2}$$



existiert der Funktionsgrenzwert an der Stelle $x_0 = 2$, obwohl $x_0 \notin \mathbb{D}$:

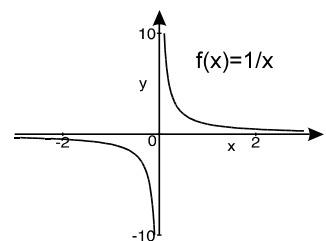
$$g_l = \lim_{\substack{x \rightarrow x_0 \\ (x < x_0)}} f(x) = \lim_{\substack{x \rightarrow 2 \\ (x < 2)}} \frac{x^2 - 2x}{x - 2} = \lim_{\substack{x \rightarrow 2 \\ (x < 2)}} x = 2,$$

$$g_r = \lim_{\substack{x \rightarrow x_0 \\ (x > x_0)}} f(x) = \lim_{\substack{x \rightarrow 2 \\ (x > 2)}} \frac{x^2 - 2x}{x - 2} = \lim_{\substack{x \rightarrow 2 \\ (x > 2)}} x = 2.$$

Der Faktor $(x - 2)$ ist im Zähler und Nenner enthalten und kann damit gekürzt werden.

③ Die Funktion

$$f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R} \text{ mit } f(x) = \frac{1}{x}$$



besitzt in $x_0 = 0$ **keinen** Grenzwert, denn

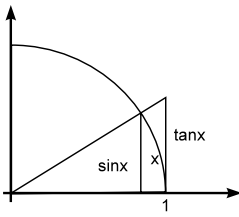
$$g_l = \lim_{h \rightarrow 0} f(x_0 - h) = \lim_{h \rightarrow 0} = -\frac{1}{h} \rightarrow -\infty,$$

$$g_r = \lim_{h \rightarrow 0} f(x_0 + h) = \lim_{h \rightarrow 0} = \frac{1}{h} \rightarrow +\infty.$$

□

Nicht für alle Funktionsgrenzwerte ist die Frage der Konvergenz so einfach zu beantworten wie in den obigen Beispielen. Man braucht dann in der Regel zusätzliche geometrische Überlegungen.

Beispiel 6.8. $\lim_{x \rightarrow 0} \frac{\sin(x)}{x} = ?$



Geometrisch entspricht der Grenzwert der Funktion $f(x) = \frac{\sin x}{x}$ an der Stelle $x_0 = 0$ der Tatsache, dass im Einheitskreis für $0 < x < \frac{\pi}{2}$ gilt: $\tan x > x > \sin x$

$$\Rightarrow \frac{1}{\cos x} > \frac{x}{\sin x} > 1 \Rightarrow \cos x < \frac{\sin x}{x} < 1$$

$$\Rightarrow 1 = \lim_{x \rightarrow 0} \cos x \leq \lim_{x \rightarrow 0} \frac{\sin x}{x} \leq 1.$$

Man erhält also insgesamt

$$\lim_{x \rightarrow 0} \frac{\sin x}{x} = 1$$

□

Beispiel 6.9. Ähnliche geometrische Überlegungen führen auf die Formel

$$\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = 1.$$

Denn für positive x -Werte ist

$$1 + x < e^x < 1 + x + x^2.$$

Damit ist $x < e^x - 1 < x(x + 1)$ bzw. $1 < \frac{e^x - 1}{x} < x + 1$. Der Grenzübergang $x \rightarrow 0$ liefert dann die behauptete Formel. □

Verhalten der Funktion für Folgen $x \rightarrow \pm\infty$: Gilt für eine beliebige Folge $(x_n)_{n \in \mathbb{N}}$ aus dem Definitionsbereich von f mit $x_n \xrightarrow{n \rightarrow \infty} \infty$, dass $f(x_n) \xrightarrow{n \rightarrow \infty} g$ konvergiert, so heißt g der Grenzwert der Funktion für $x \rightarrow \infty$:

$$\lim_{x \rightarrow \infty} f(x) = g.$$

Um die Funktionsgrenzwerte zu berechnen, sowohl für $x_n \rightarrow x_0$ als auch für $x_n \rightarrow \pm\infty$, gelten dieselben Rechenregeln wie für reelle Zahlenfolgen:

Rechenregeln für Funktionsgrenzwerte: Unter der Voraussetzung, dass die Grenzwerte $\lim_{x \rightarrow x_0} f(x)$ und $\lim_{x \rightarrow x_0} g(x)$ existieren, gelten folgende Regeln:

$$(F1) \quad \lim_{x \rightarrow x_0} c f(x) = c \lim_{x \rightarrow x_0} f(x).$$

$$(F2) \quad \lim_{x \rightarrow x_0} (f(x) \pm g(x)) = \lim_{x \rightarrow x_0} f(x) \pm \lim_{x \rightarrow x_0} g(x).$$

$$(F3) \quad \lim_{x \rightarrow x_0} (f(x) \cdot g(x)) = \lim_{x \rightarrow x_0} f(x) \cdot \lim_{x \rightarrow x_0} g(x).$$

$$(F4) \quad \lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{\lim_{x \rightarrow x_0} f(x)}{\lim_{x \rightarrow x_0} g(x)}, \quad \text{falls } \lim_{x \rightarrow x_0} g(x) \neq 0.$$

Bemerkungen:

- (1) Die Regeln gelten auch für Grenzwerte von Funktionen für $x \rightarrow \pm\infty$, falls die Grenzwerte $\lim_{x \rightarrow \pm\infty} f(x)$ und $\lim_{x \rightarrow \pm\infty} g(x)$ existieren.
- (2) Für Grenzwerte vom Typ $\frac{0}{0}$ und $\frac{\infty}{\infty}$ gelten die *Regeln von l'Hospital*, auf die in Abschnitt 7.7.3 näher eingegangen wird!

Beispiele 6.10:

$$\textcircled{1} \quad \lim_{x \rightarrow 0} \frac{x^2 - 2x + 5}{\cos x} = \frac{\lim_{x \rightarrow 0} (x^2 - 2x + 5)}{\lim_{x \rightarrow 0} \cos x} = \frac{5}{1} = 5.$$

$$\textcircled{2} \quad \lim_{x \rightarrow \infty} \frac{2x^2 + 4}{x^2 - 1} = \lim_{x \rightarrow \infty} \frac{2x^2 + 4}{x^2 - 1} \cdot \frac{\frac{1}{x^2}}{\frac{1}{x^2}} = \lim_{x \rightarrow \infty} \frac{2 + \frac{4}{x^2}}{1 - \frac{1}{x^2}} = \frac{2}{1} = 2.$$

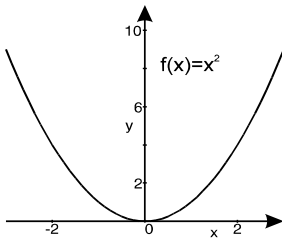
$$\textcircled{3} \quad \lim_{x \rightarrow 1} \frac{x - 1}{x^2 - 1} = \lim_{x \rightarrow 1} \frac{x - 1}{(x - 1)(x + 1)} = \lim_{x \rightarrow 1} \frac{1}{x + 1} = \frac{1}{2}.$$

$$\textcircled{4} \quad \lim_{x \rightarrow \infty} \frac{4 + 2x}{x^2 + 1} = \lim_{x \rightarrow \infty} \frac{4 + 2x}{x^2 + 1} \cdot \frac{\frac{1}{x^2}}{\frac{1}{x^2}} = \lim_{x \rightarrow \infty} \frac{\frac{4}{x^2} + \frac{2}{x}}{1 + \frac{1}{x^2}} = \frac{0}{1} = 0.$$

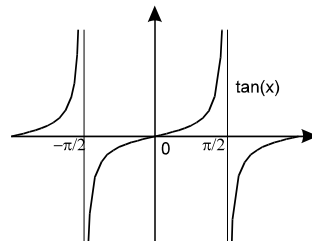
$$\begin{aligned} \textcircled{5} \quad \lim_{x \rightarrow 0} \frac{\sqrt{x+1} - 1}{x} &= \lim_{x \rightarrow 0} \frac{(\sqrt{x+1} - 1)(\sqrt{x+1} + 1)}{x(\sqrt{x+1} + 1)} \\ &= \lim_{x \rightarrow 0} \frac{(x+1) - 1}{x(\sqrt{x+1} + 1)} = \lim_{x \rightarrow 0} \frac{1}{\sqrt{x+1} + 1} = \frac{1}{2}. \quad \square \end{aligned}$$

6.3 Stetigkeit einer Funktion

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt *stetig*, wenn der Graph keine Sprünge aufweist. Mit dieser Erklärung hat man sich lange Zeit begnügt, und für die meisten Anwendungen reicht diese anschauliche Interpretation aus. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist demnach stetig. Um auch Grenzfälle wie z.B. die Funktion $f : \mathbb{R} \setminus \{\frac{\pi}{2} + k\pi, k \in \mathbb{Z}\} \rightarrow \mathbb{R}$ mit $f(x) = \tan x$ klassifizieren zu können, benötigt man die folgende präzise Definition:



Quadratfunktion



Tangens

Definition: (Stetigkeit). Ist $x_0 \in \mathbb{D}$ und ist die Funktion f in einer Umgebung von x_0 definiert. Die Funktion f heißt **stetig** in x_0 , wenn der Funktionsgrenzwert in x_0 existiert und mit dem Funktionswert $f(x_0)$ übereinstimmt.

kurz: f ist in $x_0 \in \mathbb{D}$ **stetig**, wenn

$$\lim_{h \rightarrow 0} f(x_0 + h) = \lim_{h \rightarrow 0} f(x_0 - h) = f(x_0).$$

Bemerkungen:

- (1) Die Stetigkeit im Punkte x_0 setzt voraus, dass $x_0 \in \mathbb{D}$. Stellen, an denen f nicht definiert ist, sind *Definitionslücken*. Dort wird die Stetigkeit nicht untersucht.
- (2) Ist f in jedem Punkt $x \in \mathbb{D}$ stetig, so nennt man f eine *stetige Funktion*.
- (3) Man kann die Stetigkeit einer Funktion bei x_0 auch umformulieren:

$$\lim_{x \rightarrow x_0} f(x) = f\left(\lim_{x \rightarrow x_0} x\right) = f(x_0).$$

Bei Stetigkeit dürfen Grenzwertbildung und Funktionsauswertung vertauscht werden.

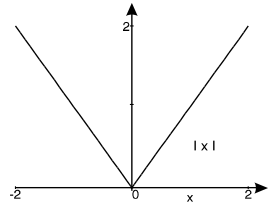
Beispiele 6.11:

- ① **Polynome** $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = a_0 + a_1 x + \dots + a_n x^n$ sind in jedem Punkt $x \in \mathbb{R}$ **stetig**.

- ② Die **Betragsfunktion**

$$f : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } f(x) = |x|$$

ist auch bei $x_0 = 0$ **stetig**, da



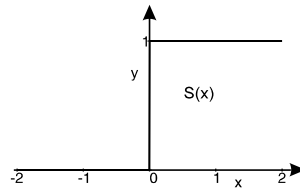
$$\lim_{h \rightarrow 0} f(x_0 + h) = \lim_{h \rightarrow 0} f(h) = \lim_{h \rightarrow 0} h = 0$$

$$\lim_{h \rightarrow 0} f(x_0 - h) = \lim_{h \rightarrow 0} f(-h) = \lim_{h \rightarrow 0} |-h| = 0$$

$$f(0) = 0.$$

- ③ Die **Sprungfunktion (Heavisidefunktion)**, die zur Beschreibung von Einschaltvorgängen dient, $S : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$S(x) = \begin{cases} 1 & \text{für } x \geq 0 \\ 0 & \text{für } x < 0 \end{cases}$$



ist bei $x_0 = 0$ **nicht stetig**, da sie einen Sprung aufweist:

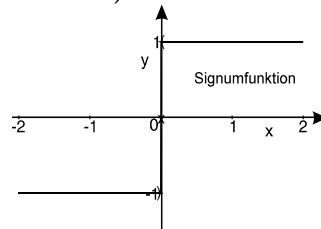
$$\lim_{h \rightarrow 0} f(x_0 + h) = \lim_{h \rightarrow 0} S(h) = 1,$$

$$\lim_{h \rightarrow 0} f(x_0 - h) = \lim_{h \rightarrow 0} S(-h) = 0.$$

- ④ Die **Vorzeichenfunktion (Signumfunktion)**

$$\text{sign} : \mathbb{R} \rightarrow \mathbb{R} \text{ mit}$$

$$\text{sign}(x) = \begin{cases} 1 & \text{für } x > 0 \\ 0 & \text{für } x = 0 \\ -1 & \text{für } x < 0 \end{cases}$$

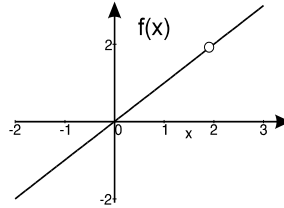


ist an der Stelle $x_0 = 0$ **nicht stetig**.

⑤ Die Funktion

$$f : \mathbb{R} \setminus \{2\} \rightarrow \mathbb{R}$$

$$\text{mit } f(x) = \frac{x^2 - 2x}{x - 2}$$



hat an der Stelle $x_0 = 2$ eine Definitionslücke. Nach Beispiel 6.7 ② existieren in $x_0 = 2$ der rechtsseitige und linksseitige Grenzwert und stimmen überein. Man definiert die **stetige Erweiterung** von f :

$$\tilde{f} : \mathbb{R} \rightarrow \mathbb{R} \quad \text{mit} \quad \tilde{f}(x) = \begin{cases} f(x) & \text{für } x \neq 2 \\ 2 & \text{für } x = 2 \end{cases}.$$

Dann ist die Funktion $\tilde{f}$ in x_0 stetig und damit für alle $x \in \mathbb{R}$ stetig. Oftmals verzichtet man auf die Notation $\tilde{f}$ und verwendet als Bezeichnung für die stetige Erweiterung wieder den Funktionsnamen f . $\square$

6.4 Intervallhalbierungs-Methode

6.4

Grundlage für eine einfache numerische Methode zur Bestimmung von Nullstellen einer Funktion bildet der folgende, anschauliche Satz: Jede stetige Funktion, die auf einem Intervall $[a, b]$ einen Vorzeichenwechsel hat, besitzt in diesem Intervall eine Nullstelle (siehe Abb. 6.4):

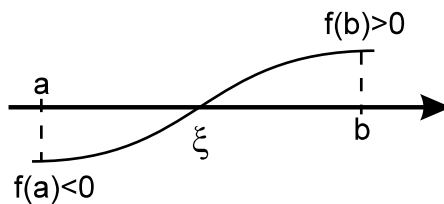


Abb. 6.4. Intervallhalbierungs-Methode

Zwischenwertsatz: Sei $f : [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion mit $(f(a) < 0 \text{ und } f(b) > 0)$ oder $(f(a) > 0 \text{ und } f(b) < 0)$. Dann existiert eine Zwischenstelle $\xi \in (a, b)$ mit der Eigenschaft

$$f(\xi) = 0.$$

Die Idee des Beweises besteht darin, dass man zu gegebenem Intervall $[a, b]$ die Intervallmitte m bestimmt und die Funktionswerte

$$f(a), f(m), f(b)$$

miteinander vergleicht: Man ersetzt den Intervallrand durch m , dessen Funktionswert dasselbe Vorzeichen wie $f(m)$ besitzt. Anschließend wiederholt man die Vorgehensweise auf das halbierte Intervall usw.

Dieses Verfahren liefert direkt einen Algorithmus (Iteration), um eine Nullstelle im Intervall $[a, b]$ zu berechnen. Die Voraussetzung ist nur, dass die Funktionswerte an den Intervallgrenzen ein unterschiedliches Vorzeichen besitzen. Da man ständig die Intervalllänge halbiert, nennt man das Verfahren auch Intervallhalbierungs-Methode.



Hinweis: Auf der CD-Rom befindet sich ein eigenes Kapitel über das [numerische Lösen von Gleichungen](#). In dem Abschnitt über die [Intervallhalbierungs-Methode](#) findet man auch einen ausführlichen Beweis des Zwischenwertsatzes, der daraus resultierenden Konvergenz des Verfahrens, die Beschreibung des numerischen Algorithmus zusammen mit der MAPLE-Prozedur [bise](#). □

Bemerkungen:

- (1) $\triangle$ Die Funktion f kann mehrere Nullstellen im Intervall $[a, b]$ besitzen. Die Bisektionsmethode liefert aber nur eine.
- (2) Man nennt Algorithmen, welche die Nullstelle in einem immer kleiner werdenden Intervall einschließen, auch *Einschließungsalgorithmen*. Ausgehend von einem Startintervall, wird dieses Intervall systematisch durch den gleichen Algorithmus verkleinert. Man nennt diesen Prozess eine *Iteration*.
- (3) Die Anzahl der benötigten Iterationen für eine vorgegebene Genauigkeit kann vor der Rechnung abgeschätzt werden, denn es gilt für den Abstand der Nullstelle von der linken Intervallgrenze

$$|a_n - \xi| \leq \frac{b-a}{2^n} \quad \text{für } n = 1, 2, 3, \dots$$

Ist f eine stetige Funktion auf $[0, 1]$ mit Vorzeichenwechsel und soll die Nullstelle $\xi \in (0, 1)$ bis auf 4 Dezimalstellen genau bestimmt werden, so darf der Abstand von $|a_n - \xi|$ nicht größer sein als $9 \cdot 10^{-5}$. n muss also so gewählt werden, dass $\frac{b-a}{2^n} \leq 9 \cdot 10^{-5} \Rightarrow n = 14$.

(Bei einer Genauigkeit von 6 Dezimalstellen ist n immerhin schon 21.)

Beispiel 6.12 (Mit MAPLE-Prozedur): Gegeben ist die Funktion

$$f(x) = x^3 - \sqrt{x^2 + 1}.$$

Gesucht ist die Nullstelle im Intervall $[1, 2]$. Da die Funktionswerte an den Intervallgrenzen unterschiedliches Vorzeichen besitzen ($f(1) = -0.4142$ und $f(2) = 5.7639$), erhält man mit dem Bisektionsverfahren die Nullstelle:

n	a	b	$f(\frac{a+b}{2})$
	1.0	2.0	1.5722
1	1.0	1.5	0.3523
2	1.0	1.25	-0.0813
3	1.125	1.25	0.1220
4	1.125	1.1875	0.0171
5	1.125	1.1562	-0.0329
6	1.1406	1.1562	-0.0081
7	1.1484	1.1562	0.0044
8	1.1484	1.1523	0.0018
9	1.1503	1.1523	0.0012
10	1.1503	1.1513	$-2 \cdot 10^{-4}$
11	1.1508	1.1513	$5 \cdot 10^{-4}$
12	1.1508	1.1511	$1 \cdot 10^{-4}$
13	1.1508	1.1510	$-7 \cdot 10^{-5}$

$$\Rightarrow \xi \approx 1.1509$$



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich eine erweiterte MAPLE-Prozedur, **bise_ext**, die den Konvergenzprozess in Form einer Animation visualisiert. □

MAPLE-Worksheets zu Kapitel 6



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 6 mit MAPLE zur Verfügung.

- Zahlenfolgen
- Babylonisches Wurzelziehen
- Funktionsfolgen
- Bisektionsmethode
- MAPLE-Lösungen zu den Aufgaben

6.5 Aufgaben zu Grenzwert und Stetigkeit

6.1 Für welche $n \in \mathbb{N}$ gelten die Ungleichungen

$$\text{a) } \left| \frac{1}{n^2} \right| < 10^{-6} \quad \text{b) } \left| \frac{1}{n^2} + 1 \right| < 1 + 10^{-8} \quad \text{c) } \left| \frac{1}{n+1} \right| < 10^{-10}$$

6.2 Bestimmen Sie - falls möglich - den Grenzwert der Zahlenfolgen für $n \rightarrow \infty$

$$\begin{array}{lll} \text{a) } a_n = \frac{2n+1}{4n^3+4n+1} & \text{b) } a_n = \frac{n^2+4}{n} & \text{c) } a_n = \frac{n^2+4n-1}{n^2-3n} \\ \text{d) } a_n = \frac{1}{n} \frac{5n^3+4n+1}{n^2} & \text{e) } a_n = \frac{1}{n} \left(\sum_{k=1}^n k \right) - \frac{1}{2}n & \text{f) } a_n = \frac{4(n+1)^4}{3n^4+3n^2+5} \\ \text{g) } a_n = \frac{4n+1}{5n-1} & \text{h) } a_n = \frac{3n^2+4n}{\sqrt[3]{n^6+n^4+1}} & \text{i) } \sin \left(\frac{\pi n^3+n^2}{2(n^3+4)} \right) \end{array}$$

6.3 Zeigen Sie, dass die Folge $a_n = \frac{2n+1}{4n}$ konvergiert.6.4 a) Gegeben sei die rekursiv definierte Folge a_n mit $a_0 := 0$ und $a_{n+1} := \frac{1}{3}(a_n + 1)$ für $n \geq 0$. Bestimmen Sie explizit das allgemeine Glied a_n und berechnen Sie den Grenzwert der Folge.b) Zeigen Sie: Ist $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| < 1$, dann ist a_n eine Nullfolge.

6.5 Man beweise die Limesrechenregeln

$$\begin{aligned} (I1) \quad \lim_{n \rightarrow \infty} (a_n + b_n) &= \lim_{n \rightarrow \infty} a_n + \lim_{n \rightarrow \infty} b_n \\ (I2) \quad \lim_{n \rightarrow \infty} (a_n \cdot b_n) &= \left(\lim_{n \rightarrow \infty} a_n \right) \cdot \left(\lim_{n \rightarrow \infty} b_n \right) \end{aligned}$$

6.6 Man berechne mit MAPLE die Grenzwerte

$$\text{a) } \lim_{n \rightarrow \infty} \frac{1}{n} \left(1 + \frac{1}{2} + \dots + \frac{1}{n} \right) \quad \text{b) } \lim_{n \rightarrow \infty} \frac{n}{\sqrt[n]{n!}}$$

6.7 Berechnen Sie die Grenzwerte der Funktionen:

$$\text{a) } \lim_{x \rightarrow 1} (x^3 + 5x^2 - 3x + 4) \quad \text{b) } \lim_{x \rightarrow 0} \frac{x^2 - 2x}{x^2 + 3x} \quad \text{c) } \lim_{x \rightarrow \infty} \frac{x}{1 + x^2}$$

6.8 Berechnen Sie die Funktionsgrenzwerte

$$\begin{array}{lll} \text{a) } \lim_{x \rightarrow 1} \frac{x^2 - 1}{x^2 + 1} & \text{b) } \lim_{x \rightarrow -3} \frac{x^2 - x - 12}{x + 3} & \text{c) } \lim_{x \rightarrow 0} \frac{\sin(2x)}{\sin(x)} \\ \text{d) } \lim_{x \rightarrow 2} \frac{(x-2)(3x+1)}{4x-8} & \text{e) } \lim_{x \rightarrow 0} \frac{\sqrt{1+x}-1}{x} & \text{f) } \lim_{x \rightarrow \infty} \frac{x^2}{x^2 - 4x + 1} \end{array}$$

6.9 Welchen Grenzwert besitzt die Funktion $f(x) = \frac{1-x}{1-\sqrt{x}}$ für $x \rightarrow 1$?6.10 Man zeige, dass die Funktion $f(x) = \begin{cases} x & x \leq 0 \\ x-2 & x > 0 \end{cases}$ an der Stelle $x_0 = 0$ unstetig ist.6.11 Man zeige, dass die Funktion $f(x) = \begin{cases} \frac{x^2-1}{x-1} & \text{für } x \neq 1 \\ 2 & \text{für } x = 1 \end{cases}$ an der Stelle $x_0 = 1$ stetig ist.6.12 Lassen sich die Definitionslücken der Funktion $f(x) = \frac{x^2 - x}{x^3 - x^2 + x - 1}$ stetig heben?

Kapitel 7
Differenzialrechnung

7

7	Differenzialrechnung	241
7.1	Einführung	243
7.2	Rechenregeln bei der Differenziation	249
7.2.1	Faktorregel	249
7.2.2	Summenregel	249
7.2.3	Produktregel	250
7.2.4	Quotientenregel	251
7.2.5	Kettenregel	252
7.2.6	Begründung der Formeln 7.2.1 - 7.2.5	254
7.2.7	Ableitung der Umkehrfunktion	255
7.2.8	Logarithmische Differenziation	258
7.2.9	Implizite Differenziation	260
7.3	Anwendungsbeispiele aus Physik und Technik	262
7.3.1	Kinematik	262
7.3.2	Induktionsgesetz	263
7.3.3	Elektrostatik	264
7.4	Differenzial einer Funktion	265
7.4.1	Linearisierung von Funktionen	266
7.4.2	Fehlerrechnung	268
7.5	Anwendungen in der Mathematik	270
7.5.1	Geometrische Bedeutung von f' und f''	270
7.5.2	Relative Extremalwerte	271
7.5.3	Wendepunkte und Sattelpunkte	272
7.5.4	Kurvendiskussion	274
7.6	Extremwertaufgaben (Optimierungsprobleme)	277
7.7	Sätze der Differenzialrechnung	281
7.7.1	Satz über die Exponentialfunktion	282
7.7.2	Mittelwertsatz	283
7.7.3	Die Regeln von l'Hospital	284
7.8	Spektrum eines strahlenden schwarzen Körpers	287
7.9	Newton-Verfahren	289
7.10	Aufgaben zur Differenzialrechnung	293

Zusätzliche Abschnitte auf der CD-Rom

7.11	Differenziation mit MAPLE	cd
------	---------------------------	----

7 Differenzialrechnung

Eine der wichtigsten Aufgaben in der angewandten Mathematik ist die Berechnung der Ableitung einer Funktion. Viele physikalische Gesetzmäßigkeiten lassen sich nur über die Differenziation einer physikalischen Größe beschreiben. Ist beispielsweise bei einem nichtlinearen Bewegungsvorgang das *Weg-Zeit-Gesetz* $s(t)$ gegeben, dann ist die Geschwindigkeit $v(t)$ die Ableitung des Weg-Zeit-Gesetzes nach der Zeit t . Die konkrete Bestimmung der Geschwindigkeit setzt rechentechnisch voraus, dass man die Funktion $s(t)$ ableiten kann.

Immer dann, wenn sich physikalische Größen mit der Zeit nichtlinear ändern, benötigt man zur Beschreibung dieser Änderung die Ableitung. Aber auch die Fehlerrechnung, die Bestimmung der Extremwerte einer Funktion bzw. die Optimierungsaufgaben führen auf das Problem der Ableitung einer Funktion. Der Ableitungsbegriff ist also motiviert durch die physikalische Beschreibung von Bewegungsabläufen (Geschwindigkeit, Beschleunigung) und durch die mathematische Beschreibung von Kurven (Tangente, Kurvendiskussion).

Hinweis: Auf der CD-Rom befindet sich ein zusätzliches Kapitel über das [numerische Differenzieren](#) und Methoden beim numerischen Differenzieren.

7.1 Einführung

7.1

Eine der wichtigsten Methoden in der angewandten Mathematik ist die der Differenziation einer Funktion. Viele physikalische Gesetzmäßigkeiten lassen sich nur über die Differenziation einer physikalischen Größe beschreiben. Wir betrachten zur Einführung zwei einfache Weg-Zeit-Gesetze aus der Kinematik.

Anwendungsbeispiel 7.1 (Gleichförmige Bewegung).

In Abb. 7.1 ist das Weg-Zeit-Gesetz einer gleichförmigen Bewegung dargestellt. Ändert sich der Weg in der Zeiteinheit Δt um Δs , so bewegt sich der Körper mit der **Geschwindigkeit**

$$v := \frac{\Delta s}{\Delta t}.$$

Die Geschwindigkeit hängt weder vom Zeitpunkt t noch vom Messintervall Δt ab.

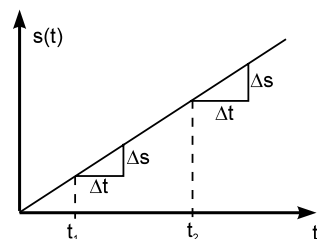


Abb. 7.1. Gleichförmige Bewegung

Anwendungsbeispiel 7.2 (Gleichförmig beschleunigten Bewegung).

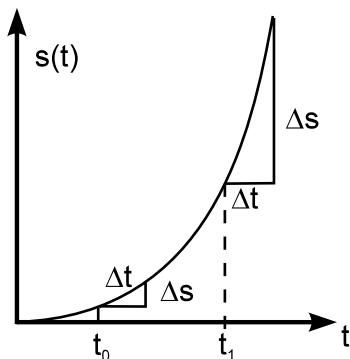


Abb. 7.2. Beschleunigte Bewegung

In Abb. 7.2 ist das Weg-Zeit-Gesetz einer gleichförmig beschleunigten Bewegung dargestellt.

$$s(t) = \frac{1}{2} g t^2.$$

Durch eine hypothetische Messreihe werden wir den Wert der Geschwindigkeit zum Zeitpunkt $t_0 = 1$ bestimmen. Es zeigt sich, dass die "gemessene" Geschwindigkeit v nun vom Messintervall Δt abhängt: Je kleiner das Messintervall Δt gewählt wird, umso genauer bestimmt sich der Wert der Geschwindigkeit zu diesem Zeitpunkt.

Wir berechnen das Verhältnis $\frac{\Delta s}{\Delta t}$ in Abhängigkeit von Δt :

$$v(t_0) = \frac{\Delta s}{\Delta t} = \frac{\frac{1}{2} g (t_0 + \Delta t)^2 - \frac{1}{2} g t_0^2}{\Delta t} = \frac{1}{2} g \frac{(t_0 + \Delta t)^2 - t_0^2}{\Delta t}.$$

Messintervall	Geschwindigkeit
$\Delta t = 1$	$v(1) = \frac{1}{2} g \frac{4-1}{1} = \frac{3}{2} g$
$\Delta t = \frac{1}{2}$	$v(1) = \frac{1}{2} g \frac{(1+\frac{1}{2})^2 - 1}{\frac{1}{2}} = \frac{5}{4} g$
$\Delta t = \frac{1}{4}$	$v(1) = \frac{1}{2} g \frac{(1+\frac{1}{4})^2 - 1}{\frac{1}{4}} = \frac{9}{8} g$
$\Delta t = \frac{1}{8}$	$v(1) = \frac{1}{2} g \frac{(1+\frac{1}{8})^2 - 1}{\frac{1}{8}} = \frac{17}{16} g$
$\Delta t = \frac{1}{16}$	$v(1) = \frac{1}{2} g \frac{(1+\frac{1}{16})^2 - 1}{\frac{1}{16}} = \frac{33}{32} g$

Beobachtung: Für kleiner werdende Messintervalle Δt geht die Geschwindigkeit gegen den Wert g . Die Mathematik stellt ein Hilfsmittel bereit, die sog. *Differenzialrechnung*, Δt gegen Null streben zu lassen.

$$v(t_0) := \lim_{\Delta t \rightarrow 0} \frac{s(t_0 + \Delta t) - s(t_0)}{\Delta t} = \left. \frac{ds}{dt} \right|_{t_0} \quad (\text{Momentangeschwindigkeit})$$

Für beliebiges t_0 gilt

$$\begin{aligned} v(t_0) &= \lim_{\Delta t \rightarrow 0} \frac{s(t_0 + \Delta t) - s(t_0)}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{1}{2} g \frac{(t_0 + \Delta t)^2 - t_0^2}{\Delta t} \\ &= \lim_{\Delta t \rightarrow 0} \frac{1}{2} g \frac{2 \Delta t t_0 + (\Delta t)^2}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{1}{2} g (2 t_0 + \Delta t) = g t_0 \end{aligned}$$

□

Wir übertragen diese Vorgehensweise auf beliebige Funktionen:

Definition: (Ableitung einer Funktion). Eine Funktion $f : \mathbb{D} \rightarrow \mathbb{R}$ heißt im Punkte $x_0 \in \mathbb{D}$ **differenzierbar**, falls der Grenzwert

$$f'(x_0) := \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x} := \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0}$$

existiert. Man bezeichnet ihn als erste **Ableitung** der Funktion f im Punkte x_0 .

Eine Funktion heißt **differenzierbar**, falls sie in jedem Punkt $x \in \mathbb{D}$ differenzierbar ist.

Bemerkungen:

- (1) Man notiert die Ableitung der Funktion $y = f(x)$ auch durch die Symbole

$$f'(x), y', \frac{dy}{dx}, \frac{df(x)}{dx}.$$

- (2) Die Bestimmung der Ableitung bezeichnet man als *Differenziation* der Funktion $f(x)$.
 (3) Den Quotienten

$$\frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}$$

nennt man den *Differenzenquotient* und

$$\lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}$$

den *Differenzialquotient*.

- (4) Für $\Delta x \rightarrow 0$ geht sowohl der Zähler als auch der Nenner des Differenzenquotienten gegen Null. Man hat also den Fall $\frac{0}{0}$! Um die Ableitung für elementare Funktionen bestimmen zu können, muss man bei konkret gegebener Funktion den Zähler so lange (geschickt) umformen, bis man Δx ausklammern und mit Δx im Nenner kürzen kann. Anschließend lässt man im verbleibenden Term $\Delta x \rightarrow 0$ gehen (siehe Beispiele 7.3).

Geometrische Interpretation der Ableitung: Gegeben sei die Funktion $f(x)$ und $x_0 \in \mathbb{D}$. Dann ist die **Steigung der Sekante** durch die Punkte $P(x_0, f(x_0))$ und $Q(x_0 + \Delta x, f(x_0 + \Delta x))$ gegeben durch

$$\frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x} \quad (\text{Differenzenquotient}).$$

Durch den Grenzübergang $\Delta x \rightarrow 0$ bleibt P fest, aber Q wandert auf der Kurve und strebt gegen P . Die Steigung der Sekante geht so in die Steigung der Tangente im Punkte $P(x_0, f(x_0))$ über. **Die Ableitung einer Funktion im Punkte x_0 ist demnach die Steigung der Tangente an die Kurve im Punkte $(x_0, f(x_0))$.**

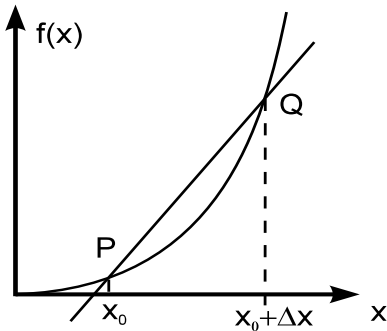


Abb. 7.3. Differenzenquotient als Sekantensteigung



Visualisierung mit MAPLE: In der [MAPLE-Animation](#) erkennt man, dass der Punkt Q auf der Kurve von f entlang zum Punkt P wandert. Dabei nähert sich die Sekante der Tangente an und die Sekantensteigung geht in die Tangentensteigung über.

Numerische Differenziation: Numerisch kann der Grenzübergang $h \rightarrow 0$ bei der Ableitung nicht durchgeführt werden, da dies sofort zu einem *Overflow* führen würde. Man würde ja durch 0 dividieren. Um die Ableitung einer Funktion f an der Stelle x_0 näherungsweise zu berechnen, ersetzt man die Ableitung einer Funktion f im Punkte x_0 durch die Sekantensteigung

$$f'(x_0) \approx \frac{f(x_0 + h) - f(x_0)}{h}$$

mit $h > 0$. Dies ist die sog. **rechtsseitige Differenzenformel**. Eine genauere Näherung erhält man, wenn man den Mittelwert der rechtsseitigen und linksseitigen Formel nimmt. Dies ergibt die **zentrale Differenzenformel**:

$$f'(x) \approx \frac{1}{2} \frac{f(x_0 + h) - f(x_0 - h)}{h}$$

Hinweis: Auf der CD-Rom befindet sich ein eigenes Kapitel über das [numerische Differenzieren](#) und den Eigenschaften der diskreten Formeln.

Beispiele 7.3:

- ① $f(x) = c \Rightarrow f'(x) = 0$ (Konstante Funktion):

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{c - c}{h} = 0.$$

- ② $f(x) = x \Rightarrow f'(x) = 1$ (Identische Funktion):

$$f'(x) = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{x+h-x}{h} = 1.$$

- ③ $f(x) = x^2 \Rightarrow f'(x) = 2x$ (Quadratische Funktion):

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{(x+h)^2 - x^2}{h} \\ &= \lim_{h \rightarrow 0} \frac{2hx + h^2}{h} = \lim_{h \rightarrow 0} (2x + h) = 2x. \end{aligned}$$

- ④ $f(x) = \frac{1}{x} \Rightarrow f'(x) = -\frac{1}{x^2}$ ($x \neq 0$):

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{1}{h} (f(x+h) - f(x)) = \lim_{h \rightarrow 0} \frac{1}{h} \left(\frac{1}{x+h} - \frac{1}{x} \right) \\ &= \lim_{h \rightarrow 0} \frac{1}{h} \frac{-h}{x(x+h)} = -\frac{1}{x^2}. \end{aligned}$$

- ⑤ $f(x) = e^x \Rightarrow f'(x) = e^x$ (Exponentialfunktion):

$$f'(x) = \lim_{h \rightarrow 0} \frac{1}{h} (e^{x+h} - e^x) = \lim_{h \rightarrow 0} \frac{1}{h} e^x (e^h - 1) = e^x \lim_{h \rightarrow 0} \frac{1}{h} (e^h - 1).$$

Nach Beispiel 6.9 ist $\lim_{h \rightarrow 0} \frac{1}{h} (e^h - 1) = 1$, so dass $\boxed{(e^x)' = e^x}$.

- ⑤ $f(x) = \sin(x) \Rightarrow f'(x) = \cos(x)$ (Sinusfunktion):

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{1}{h} (\sin(x+h) - \sin(x)) = \lim_{h \rightarrow 0} \frac{1}{h} (2 \cos\left(\frac{2x+h}{2}\right) \sin \frac{h}{2}) \\ &= \left\{ \lim_{h \rightarrow 0} \cos\left(\frac{2x+h}{2}\right) \right\} \cdot \left\{ \lim_{h \rightarrow 0} \frac{1}{h} \sin \frac{h}{2} \right\} = \cos(x), \end{aligned}$$

da nach Beispiel 6.8 $\lim_{h \rightarrow 0} \frac{\sin h}{h} = 1$.

Tabelle 7.1: Ableitungen elementarer Funktionen

	$f(x)$	$f'(x)$
Potenzfunktion	x^n	$n x^{n-1}$
Trigonometrische Funktionen	$\sin(x)$	$\cos(x)$
	$\cos(x)$	$-\sin(x)$
	$\tan(x)$	$\frac{1}{\cos^2(x)}$
	$\cot(x)$	$-\frac{1}{\sin^2(x)}$
Arkusfunktionen	$\arcsin(x)$	$\frac{1}{\sqrt{1-x^2}}$
	$\arccos(x)$	$-\frac{1}{\sqrt{1-x^2}}$
	$\arctan(x)$	$\frac{1}{1+x^2}$
	$\operatorname{arccot}(x)$	$-\frac{1}{1+x^2}$
Exponentialfunktion	e^x	e^x
	a^x	$a^x \ln a$
Logarithmusfunktion	$\ln(x)$	$\frac{1}{x}$
	$\log_a(x)$	$\frac{1}{\ln(a) \cdot x}$
Hyperbelfunktionen	$\sinh(x)$	$\cosh(x)$
	$\cosh(x)$	$\sinh(x)$
	$\tanh(x)$	$\frac{1}{\cosh^2(x)}$
	$\coth(x)$	$-\frac{1}{\sinh^2(x)}$
Areafunktionen	$\operatorname{ar\,sinh}(x)$	$\frac{1}{\sqrt{x^2+1}}$
	$\operatorname{ar\,cosh}(x)$	$\frac{1}{\sqrt{x^2-1}}$
	$\operatorname{ar\,tanh}(x)$	$\frac{1}{1-x^2}$ für $ x < 1$
	$\operatorname{ar\,coth}(x)$	$\frac{1}{1-x^2}$ für $ x > 1$

Höhere Ableitungen von f : Existiert zu einer Funktion $f: \mathbb{ID} \rightarrow \mathbb{IR}$ die Ableitung $f': \mathbb{ID} \rightarrow \mathbb{IR}$ und ist $f'(x)$ wieder differenzierbar, so bezeichnet man deren Ableitung als *zweite Ableitung* $f''(x)$ von f . $f''(x)$ ist die Ableitung der Funktion $f'(x)$, d.h. es gilt $f'' = (f')'$. Durch wiederholtes Differenzieren gelangt man schließlich zu den Ableitungen höherer Ordnung. Man schreibt:

1. Ableitung $y' = f'(x) = \frac{d}{dx} f(x).$
2. Ableitung $y'' = f''(x) = \frac{d^2}{dx^2} f(x).$
- $\vdots$
- n -te Ableitung $y^{(n)} = f^{(n)}(x) = \frac{d^n}{dx^n} f(x).$

7.2 Rechenregeln bei der Differenziation

7.2

Um komplizierte Ausdrücke und funktionale Zusammenhänge differenzieren zu können, benötigt man Ableitungs- bzw. Differenziationsregeln. Im Folgenden werden die wichtigsten Ableitungsregeln vorgestellt und jeweils an Beispielen verdeutlicht. Wir gehen davon aus, dass alle in den Formeln auftretenden Funktionen differenzierbar sind.

7.2.1 Faktorregel

Faktorregel: Ein konstanter Faktor darf beim Differenzieren vorgezogen werden:

$$y = c f(x) \Rightarrow y' = c f'(x)$$

Beispiele 7.4:

- ① $x(t) = 220 \sin(t) \Rightarrow x'(t) = 220 (\sin(t))' = 220 \cos(t).$
- ② $f(x) = 9x^5 \Rightarrow f'(x) = 9(x^5)' = 45x^4.$
- ③ $s(t) = -5e^t \Rightarrow s'(t) = -5(e^t)' = -5e^t. \quad \square$

7.2.2 Summenregel

Summenregel: Eine Summe von Funktionen wird gliedweise differenziert:

$$y = f_1(x) + f_2(x) \Rightarrow y' = f_1'(x) + f_2'(x)$$

$$y' = n x^{n-1} \cdot x + x^n \cdot 1 = n x^n + x^n = (n+1) x^n.$$

Bemerkungen:

- (1) Die Produktregel kann auch auf Funktionen angewendet werden, die sich aus mehr als 2 Faktoren zusammensetzen:

$$\begin{aligned} y = u \cdot v \cdot w &\Rightarrow y' = (u \cdot v)' \cdot w + u \cdot v \cdot w' \\ &= (u' \cdot v + u \cdot v') \cdot w + u \cdot v \cdot w' \\ &= u' \cdot v \cdot w + u \cdot v' \cdot w + u \cdot v \cdot w'. \end{aligned}$$

- (2) In Verallgemeinerung der Produktregel findet man für die n -te Ableitung eines Produktes die *Leibnizsche Regel*:

$$y = u \cdot v \Rightarrow y^{(n)} = \sum_{i=0}^n \binom{n}{i} u^i v^{n-i}$$

➤ **7.2.4 Quotientenregel**

Quotientenregel: Die Ableitung einer Funktion y , die sich als Quotient zweier Funktionen $\frac{u(x)}{v(x)}$ darstellen lässt, berechnet man durch:

$$y = \frac{u(x)}{v(x)} \Rightarrow y' = \frac{u'(x) \cdot v(x) - u(x) \cdot v'(x)}{v^2(x)}$$

Beispiele 7.8:

$$\begin{aligned} \textcircled{1} \quad y &= \frac{4x^2 + 3x + 5}{2x^6 + 2x^7} : \\ \Rightarrow y' &= \frac{(4x^2 + 3x + 5)' \cdot (2x^6 + 2x^7) - (4x^2 + 3x + 5) \cdot (2x^6 + 2x^7)'}{(2x^6 + 2x^7)^2} \\ &= \frac{(8x + 3) \cdot (2x^6 + 2x^7) - (4x^2 + 3x + 5) \cdot (12x^5 + 14x^6)}{(2x^6)^2 (1 + x)^2} \\ &= \frac{(8x + 3)x \cdot (1 + x) - (4x^2 + 3x + 5) \cdot (6 + 7x)}{2x^7 (1 + x)^2}. \end{aligned}$$

$$\begin{aligned} \textcircled{2} \quad y &= \tan x = \frac{\sin x}{\cos x} : \\ \Rightarrow y' &= \frac{(\sin x)' \cos x - \sin x (\cos x)'}{\cos^2 x} \\ &= \frac{\cos^2 x + \sin^2 x}{\cos^2 x} = \frac{1}{\cos^2 x} = 1 + \tan^2 x. \end{aligned}$$

$$\textcircled{3} \quad y = \cot x = \frac{\cos x}{\sin x} :$$

$$\begin{aligned}\Rightarrow y' &= \frac{(\cos x)' \sin x - \cos x (\sin x)'}{\sin^2 x} \\ &= \frac{-\sin^2 x - \cos^2 x}{\sin^2 x} = \frac{-1}{\sin^2 x} = -1 - \cot^2 x.\end{aligned}$$

$$\textcircled{4} \quad y = \frac{1}{x^n} = x^{-n}.$$

$$\Rightarrow y' = \frac{-n x^{n-1}}{x^{2n}} = -n \frac{1}{x^{n+1}} = -n x^{-n-1} \quad (n > 0).$$

□

Bemerkung: Nach Beispiel 7.7 und 7.8 $\textcircled{4}$ gilt die Potenzregel in der Form

$$\boxed{y = x^n \quad \Rightarrow \quad y' = n x^{n-1}} \quad \text{für } n \in \mathbb{Z}.$$

7.2.5 Kettenregel

Die im Folgenden diskutierte *Kettenregel* gehört zu den wichtigsten (und anfänglich schwierigsten) Ableitungsregeln. Bisher haben wir jeweils elementare Funktionen wie z.B. $\sin t$, $\cos t$, $e^x \cdot \sin x$ etc. differenziert. In den Anwendungen treten aber Funktionen der Form $f(\omega t + \varphi)$ auf. Bei der Auswertung wird zuerst die "innere" Funktion $\omega t + \varphi$ berechnet und anschließend die "äußere" Funktion f auf das Ergebnis angewendet. Die Kettenregel macht eine Aussage darüber, wie verschachtelte (= verkettete) Funktionen $y = f(g(x))$ differenziert werden.

Kettenregel: Die Ableitung einer verketteten Funktion $y = f(g(x))$ erhält man aus dem Produkt der äußeren und inneren Ableitung:

$$y = f(g(x)) \quad \Rightarrow \quad y' = f'(g(x)) \cdot g'(x).$$

In Worten besagt die Kettenregel, dass verschachtelte Funktionen $y = f(g(x))$ differenziert werden, indem zunächst die Ableitung der äußeren Funktion f' gebildet und an der Stelle $g(x)$ ausgewertet wird (= äußere Ableitung) und anschließend das Ergebnis mit der Ableitung der inneren Funktion $g'(x)$ (= innere Ableitung) multipliziert wird.

Beispiele 7.9:

$$\begin{aligned} \textcircled{1} \quad s(t) &= (t^2 - 3)^4. \\ \Rightarrow s'(t) &= 4(t^2 - 3)^3 \cdot (t^2 - 3)' = 4(t^2 - 3)^3 \cdot 2t = 8t(t^2 - 3)^3. \end{aligned}$$

$$\begin{aligned} \textcircled{2} \quad y &= e^{5x+4}. \\ \Rightarrow y' &= e^{5x+4} \cdot (5x+4)' = e^{5x+4} \cdot 5. \end{aligned}$$

$$\begin{aligned} \textcircled{3} \quad f(x) &= \ln(x^2 + 4x + 2). \\ \Rightarrow f'(x) &= \frac{1}{x^2 + 4x + 2} \cdot (x^2 + 4x + 2)' = \frac{2x + 4}{x^2 + 4x + 2}. \end{aligned}$$

$$\begin{aligned} \textcircled{4} \quad x(t) &= x_0 \sin(\omega t + \varphi). \\ \Rightarrow x'(t) &= x_0 \cos(\omega t + \varphi) \cdot (\omega t + \varphi)' = x_0 \cos(\omega t + \varphi) \cdot \omega. \quad \square \end{aligned}$$

Oftmals wird die Kettenregel in der Literatur mit einer Substitution durchgeführt: Wenn $y = f(g(x))$, so setzt man $u = g(x)$ (innere Funktion) und $y = f(u)$ (äußere Funktion). Die Ableitung ist dann

$$y' = f'(u) \cdot u'(x) \quad \text{bzw.} \quad \boxed{\frac{dy}{dx} = \frac{dy}{du} \cdot \frac{du}{dx}}.$$

Beispiele 7.10:

① $y = \ln(1 + x^2)$. Wir setzen $u = 1 + x^2$ (innere Funktion) und $f(u) = \ln u$ (äußere Funktion). Wegen $f'(u) = \frac{1}{u}$ und $u'(x) = 2x$ folgt dann

$$y' = f'(u) \cdot u'(x) = \frac{1}{u} \cdot 2x = \frac{1}{1 + x^2} \cdot 2x.$$

② $y = \sqrt[3]{(x^2 + 4x + 10)^2} = (x^2 + 4x + 10)^{\frac{2}{3}}$. Mit $u = x^2 + 4x + 10$ (innere Funktion) und $f(u) = u^{\frac{2}{3}}$ (äußere Funktion) folgt mit $f'(u) = \frac{2}{3} u^{-\frac{1}{3}}$ und $u'(x) = 2x + 4$:

$$y' = f'(u) \cdot u'(x) = \frac{2}{3} u^{-\frac{1}{3}} (2x + 4) = \frac{2}{3} (x^2 + 4x + 10)^{-\frac{1}{3}} (2x + 4).$$

③ $y = e^{x \cdot \sin x}$. Mit $u = x \cdot \sin x$ und $f(u) = e^u$ folgt wegen $f'(u) = e^u$ und $u'(x) = 1 \cdot \sin x + x \cos x$:

$$y' = f'(u) \cdot u'(x) = e^u \cdot (\sin x + x \cos x) = e^{x \sin x} (\sin x + x \cos x).$$

Die Kettenregel bedeutet, immer von Außen nach Innen zu differenzieren. Sie kann auch bei mehrmals verschachtelten Funktionen angewendet werden.

Beispiele 7.11:

① $y = \ln(\sin(2x - 3))$. Mehrmaliges Anwenden der Kettenregel liefert

$$\begin{aligned} y' &= \frac{1}{\sin(2x - 3)} \cdot (\sin(2x - 3))' = \frac{1}{\sin(2x - 3)} \cdot \cos(2x - 3) \cdot 2 \\ &= 2 \cot(2x - 3). \end{aligned}$$

② $y = \exp(\sin(4x + 2x^2))$. Mehrmaliges Anwenden der Kettenregel liefert

$$\begin{aligned} y' &= \exp(\sin(4x + 2x^2)) \cdot (\sin(4x + 2x^2))' \\ &= \exp(\sin(4x + 2x^2)) \cdot \cos(4x + 2x^2) \cdot (4x + 2x^2)' \\ &= \exp(\sin(4x + 2x^2)) \cdot \cos(4x + 2x^2) \cdot (4 + 4x). \quad \square \end{aligned}$$

➤ **7.2.6 Begründung der Formeln 7.2.1 - 7.2.5**

(1) Aus

$$\frac{1}{h} (c f(x + h) - c f(x)) = c \frac{1}{h} (f(x + h) - f(x))$$

folgt durch Grenzübergang und Limesrechenregel $L1$ die Faktorregel.

(2) Aus

$$\begin{aligned} &\frac{1}{h} ([f_1(x + h) + f_2(x + h)] - [f_1(x) + f_2(x)]) \\ &= \frac{1}{h} (f_1(x + h) - f_1(x)) + \frac{1}{h} (f_2(x + h) - f_2(x)) \end{aligned}$$

folgt durch Grenzübergang und Limesrechenregel $L2$ die Summenregel.

(3) Aus

$$\begin{aligned} &\frac{1}{h} (u(x + h) v(x + h) - u(x) v(x)) \\ &= \frac{1}{h} (u(x + h) v(x + h) - u(x) v(x + h) \\ &\quad + u(x) v(x + h) - u(x) v(x)) \\ &= \frac{1}{h} (u(x + h) - u(x)) v(x + h) + u(x) \frac{1}{h} (v(x + h) - v(x)) \end{aligned}$$

folgt durch Grenzübergang und Limesrechenregeln ($L2$), ($L3$) die Produktregel.

(4) Wegen

$$\frac{1}{h} \left(\frac{1}{v(x+h)} - \frac{1}{v(x)} \right) = -\frac{1}{h} \frac{v(x+h) - v(x)}{v(x+h) v(x)} = -\frac{\frac{1}{h} (v(x+h) - v(x))}{v(x+h) v(x)}$$

folgt durch Grenzübergang und Limesrechenregel (L4)

$$\boxed{\left(\frac{1}{v} \right)' = -\frac{v'}{v^2}.$$

Mit der Produktregel erhält man dann für die Ableitung eines Quotienten

$$\left(\frac{u}{v} \right)' = \left(u \cdot \frac{1}{v} \right)' = u' \frac{1}{v} + u \left(\frac{1}{v} \right)' = \frac{u'}{v} - u \frac{v'}{v^2} = \frac{u' v - u v'}{v^2}.$$

(5) Aus

$$\begin{aligned} & \frac{1}{h} (f(g(x+h)) - f(g(x))) \\ &= \frac{f(g(x+h)) - f(g(x))}{g(x+h) - g(x)} \cdot \frac{1}{h} (g(x+h) - g(x)) \end{aligned}$$

folgt durch Grenzübergang und Limesrechenregel (L3) die Behauptung. $\square$

► 7.2.7 Ableitung der Umkehrfunktion

Wir betrachten die Problemstellung: Gegeben ist eine invertierbare Funktion $y = f(x)$ und deren Ableitung $y' = f'(x)$. Gesucht ist die Ableitung der Umkehrfunktion $g(x) = f^{-1}(x)$. Um diese Ableitung zu berechnen, setzen wir

$$y = f(x)$$

und lösen diese Gleichung nach x auf:

$$x = f^{-1}(y) = g(y).$$

Da

$$y = f(g(y))$$

folgt durch Differenziation $\frac{d}{dy}$ und der Kettenregel

$$1 = \frac{d}{dy} y = \frac{d}{dy} f(g(y)) = f'(g(y)) \cdot g'(y) = f'(x) \cdot g'(y).$$

Löst man diese Gleichung nach $g'(y)$ auf, ist die Ableitung der Umkehrfunktion gegeben durch

$$g'(y) = \frac{1}{f'(x)} = \frac{1}{f'(g(y))}.$$

Anschließend vertauscht man die Variablen x und y miteinander.

Ableitung der Umkehrfunktion: Die Funktion $f(x)$ sei differenzierbar mit der Ableitung $f'(x)$ und besitze die Umkehrfunktion g . Die Ableitung der Umkehrfunktion $g(y)$ ist berechenbar durch

$$g'(y) = \frac{1}{f'(x)}. \quad (*)$$

Ersetzt man die Variable x durch $g(y)$ und vertauscht anschließend formal auf beiden Seiten der Gleichung die Variable x mit y , so folgt $g'(x)$.

Durch diese Formel ist man insbesondere in der Lage die Ableitungen der Umkehrfunktion von e^x und den trigonometrischen Funktionen zu bestimmen:

Beispiele 7.12:

① Ableitung von $\ln x$:

Gegeben: $y = f(x) = e^x$ mit $f'(x) = e^x$.

Gesucht: Ableitung der Umkehrfunktion $g(x) = \ln x$.

Ansatz: $y = e^x$

Auflösen nach x : $x = g(y) = \ln y$

$$\hookrightarrow g'(y) = \frac{1}{f'(x)} = \frac{1}{e^x} = \frac{1}{y}$$

Vertauschen der Variablen: $g'(x) = (\ln x)' = \frac{1}{x}$.

② Ableitung von $\arctan x$:

Gegeben: $y = f(x) = \tan x$ mit $f'(x) = \frac{1}{\cos^2 x} = \tan^2 x + 1$.

Gesucht: Ableitung der Umkehrfunktion $g(x) = \arctan x$.

Ansatz: $y = \tan x$

Auflösen nach x : $x = g(y) = \arctan y$

$$\hookrightarrow g'(y) = \frac{1}{f'(x)} = \frac{1}{\tan^2 x + 1} = \frac{1}{y^2 + 1}$$

Vertauschen der Variablen: $g'(x) = (\arctan x)' = \frac{1}{x^2 + 1}$.

□

Bemerkungen:

- (1) Es ist wichtig die Ableitung der Funktion f wieder durch die Funktion f auszudrücken, denn wählt man in Beispiel 7.12 ② für f' z.B. $f'(x) = \frac{1}{\cos^2 x}$, so erhält man das Ergebnis

$$g'(y) = \frac{1}{f'(x)} = \cos^2 x = \cos^2(\arctan y).$$

- (2) Auf analoge Weise wie $\arctan$ werden die Ableitungen von $\arcsin$, $\arccos$, arccot berechnet (siehe Übungsaufgaben).

Beispiele 7.13 (Ableitung der Hyperbel- und Areefunktionen):

- ① Die Ableitungen der Funktionen

$$\sinh: \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \sinh(x) := \frac{1}{2}(e^x - e^{-x}) \quad (\text{Sinushyperbolikus})$$

$$\cosh: \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \cosh(x) := \frac{1}{2}(e^x + e^{-x}) \quad (\text{Kosinushyperbolikus})$$

berechnet man direkt über die Ableitung der Exponentialfunktion. Es gilt

$$\sinh'(x) = \cosh(x) \quad \text{und} \quad \cosh'(x) = \sinh(x).$$

- ② Mit Sinus- und Kosinushyperbolikus gilt der Zusammenhang

$$\cosh^2(x) - \sinh^2(x) = 1.$$

Denn durch direktes Einsetzen von Sinus- und Kosinushyperbolikus in die linke Seite der Gleichung folgt

$$\cosh^2(x) - \sinh^2(x) = \left(\frac{1}{2}(e^x + e^{-x})\right)^2 - \left(\frac{1}{2}(e^x - e^{-x})\right)^2 = 1.$$

- ③ Über $\sinh$ und $\cosh$ berechnet man die Ableitung von **Tangenshyperbolicus**, der definiert ist durch

$$\tanh: \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \tanh(x) := \frac{\sinh(x)}{\cosh(x)}.$$

Denn durch Anwenden der Quotientenregel gilt

$$\tanh'(x) = \frac{\cosh^2(x) - \sinh^2(x)}{\cosh^2(x)}.$$

Verwendet man nun die Identität von ②, folgt

$$\tanh'(x) = \frac{1}{\cosh^2(x)} = 1 - \tanh^2(x).$$

- ④ Analog zu Tangenshyperbolicus definiert man den **Kotangenshyperbolicus**

$$\coth : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R} \text{ mit } \coth(x) := \frac{\cosh(x)}{\sinh(x)}$$

mit der Ableitung

$$\coth'(x) = -\frac{1}{\sinh^2(x)} = 1 - \coth^2(x).$$

- ⑤ Die **Arefunktionen** sind die Umkehrfunktionen von $\sinh$, $\cosh$, $\tanh$ und $\coth$. Die Ableitung z.B. der Funktion $\operatorname{ar} \sinh(x)$ folgt über die Ableitung der Umkehrfunktion:

$$\text{Ansatz: } y = \sinh(x) \quad \hookrightarrow \quad f'(x) = \cosh(x)$$

$$\text{Auflösen nach } x: \quad x = g(y) = \operatorname{ar} \sinh(x)$$

$$\hookrightarrow g'(y) = \frac{1}{f'(x)} = \frac{1}{\cosh(x)} = \frac{1}{\sqrt{1+\sinh^2 x}} = \frac{1}{\sqrt{1+y^2}}.$$

Vertauschen der Variablen:

$$\operatorname{ar} \sinh'(x) = \frac{1}{\sqrt{1+x^2}}.$$

➤ 7.2.8 Logarithmische Differenziation

Die Funktion $f(x) = x^{\cos x}$ ist mit den Differenziationsregeln 7.2.1 - 7.2.7 nicht unmittelbar ableitbar. Wenn man auf diese Gleichung jedoch den Logarithmus anwendet

$$\ln f(x) = \ln x^{\cos x} = \cos x \cdot \ln x$$

und mit der Kettenregel nach x differenziert

$$\frac{1}{f(x)} \cdot f'(x) = -\sin x \cdot \ln x + \cos x \cdot \frac{1}{x}$$

folgt

$$f'(x) = f(x) \cdot \left[-\sin x \cdot \ln x + \cos x \cdot \frac{1}{x} \right].$$

Man bezeichnet diese Vorgehensweise als **logarithmische Differenziation**, da man die logarithmierte Funktion differenziert.

Logarithmische Differenziation: Eine Funktion y mit

$$y = [u(x)]^{v(x)}$$

($u(x) > 0$) wird logarithmisch differenziert, indem man

(1) die Funktionsgleichung logarithmiert: $\ln y = v(x) \cdot \ln(u(x))$;

(2) die logarithmierte Gleichung differenziert:

$$\frac{1}{y} \cdot y' = v'(x) \cdot \ln(u(x)) + v(x) \cdot \frac{1}{u(x)} \cdot u'(x);$$

(3) diese Gleichung nach y' auflöst.

Bemerkung: Man beachte, dass beim Differenzieren der logarithmierten Gleichung die Funktion y von x abhängt. Dadurch muss die Kettenregel beim Differenzieren der linken Seite angewendet werden:

$$\frac{d}{dx} \ln y(x) = \frac{1}{y(x)} \cdot y'(x).$$

Beispiele 7.14:

① $y = x^x$.

Logarithmieren: $\ln y = \ln x^x = x \ln x$.

Differenzieren: $\frac{1}{y} y' = 1 \cdot \ln x + x \cdot \frac{1}{x} = \ln x + 1$.

Auflösen: $y' = y (\ln x + 1) = x^x (\ln x + 1)$.

② In vielen praktischen Beispielen besteht die zu differenzierende Funktion aus komplizierten Ausdrücken von Produkten und Quotienten. Zwar kann durch Anwenden der Produkt- und Quotientenregel die Ableitung berechnet werden, jedoch wird durch logarithmisches Differenzieren die Aufgabe vereinfacht und die Berechnung auf elegante Weise durchführbar. Sei

$$y = \frac{\sin(x-2) e^{2x}}{(x-1)^3 (x^2+3)^5}.$$

Logarithmieren:

$$\begin{aligned}\ln y &= \ln [\sin(x-2) e^{2x}] - \ln [(x-1)^3 (x^2+3)^5] \\ &= \ln(\sin(x-2)) + \ln(e^{2x}) - \ln(x-1)^3 - \ln(x^2+3)^5 \\ &= \ln(\sin(x-2)) + 2x - 3 \ln(x-1) - 5 \ln(x^2+3).\end{aligned}$$

Differenzieren:

$$\frac{1}{y} \cdot y' = \frac{1}{\sin(x-2)} \cdot \cos(x-2) + 2 - 3 \frac{1}{x-1} - 5 \frac{1}{x^2+3} \cdot 2x.$$

Auflösen:

$$y' = y \cdot \left[\cot(x-2) + 2 - \frac{3}{x-1} - \frac{10x}{x^2+3} \right].$$

- ③ Mit der logarithmischen Differenziation kann auch die Ableitung der allgemeinen Potenzfunktion $y = x^\alpha$ ($x > 0, \alpha \in \mathbb{R}$ fest) berechnet werden:

Logarithmieren: $\ln y = \ln x^\alpha = \alpha \ln x.$

Differenzieren: $\frac{1}{y} \cdot y' = \alpha \cdot \frac{1}{x}.$

Auflösen: $y' = y \alpha \frac{1}{x} = x^\alpha \alpha x^{-1} = \alpha x^{\alpha-1}.$

$$\Rightarrow y = x^\alpha \Rightarrow y' = \alpha x^{\alpha-1} \quad (\alpha \in \mathbb{R}).$$

□

7.2.9 Implizite Differenziation

In manchen Anwendungen ist eine Funktion $f(x)$ nur in einer *impliziten Form* $F(x, f(x)) = 0$ gegeben und die Bestimmungsgleichung nur schwer oder gar nicht explizit nach $y = f(x)$ auflösbar. Die Ableitung solcher implizit gegebener Funktionen kann mit Hilfe der Kettenregel berechnet werden:

Beispiel 7.15. Gegeben ist die **Kreisgleichung**

$$F(x, y) = (x-4)^2 + (y+5)^2 - 25 = 0,$$

Kreis um den Mittelpunkt $(x_0, y_0) = (4, -5)$ mit Radius 5. Gesucht ist die Steigung im Punkte $(x, y) = (7, -1)$.

Wir differenzieren jeden einzelnen Term der Gleichung nach x . Man beachte, dass hierbei $y = y(x)$ von der Variablen x abhängt! Wenn die Funktion F

identisch Null ist, dann ist auch die Ableitung von F nach x Null:

$$2(x-4) + 2(y+5) \cdot y' - 0 = 0.$$

Durch Auflösen nach y' folgt

$$y' = -\frac{x-4}{y+5}$$

und nach Einsetzen des Punktes $(x, y) = (7, -1)$ ist die Steigung

$$y' = -\frac{7-4}{-1+5} = -\frac{3}{4}. \quad \square$$

Implizite Differenziation: Ist eine Funktion $y(x)$ *implizit* gegeben durch

$$F(x, y(x)) = 0,$$

so erhält man die Ableitung der Funktion y , indem F gliedweise nach x differenziert wird. Jeder Term, der y enthält, muss unter Verwendung der Kettenregel differenziert werden. Anschließend wird die differenzierte Gleichung nach y' aufgelöst.

Beispiele 7.16:

- ① Gesucht ist die Ableitung der Funktion $y(x)$, die implizit durch die Gleichung $e^y - e^{2x} = x \cdot y$ definiert ist:

$$\text{Aus } e^y - e^{2x} = x \cdot y \quad \text{folgt} \quad e^y - e^{2x} - x \cdot y = 0.$$

$$\text{Differenziation: } e^y \cdot y' - e^{2x} \cdot 2 - (1 \cdot y + x \cdot y') = 0$$

$$\Rightarrow (e^y - x) y' - 2e^{2x} - y = 0.$$

$$\text{Auflösen: } y' = \frac{2e^{2x} + y}{e^y - x}.$$

- ② Gesucht ist die Ableitung der Funktion $y(x)$, die implizit durch die Gleichung $x \cdot \sin 2y = 1 - 3y^2$ definiert ist:

$$x \cdot \sin 2y = 1 - 3y^2 \quad \Rightarrow \quad x \cdot \sin 2y - 1 + 3y^2 = 0.$$

$$\text{Differenziation: } 1 \cdot \sin 2y + x \cos 2y \cdot 2y' - 0 + 3 \cdot 2y \cdot y' = 0.$$

$$\text{Auflösen: } y' = \frac{-\sin 2y}{2x \cos 2y + 6y}. \quad \square$$

7.3 Anwendungsbeispiele aus Physik und Technik

7.3.1 Kinematik

Bewegungsabläufe lassen sich durch das **Weg-Zeit-Gesetz** $s = s(t)$ beschreiben. Die Momentangeschwindigkeit ist die Ableitung des Weg-Zeit-Gesetzes nach der Zeit

$$v(t) := \dot{s}(t) = \frac{d}{dt} s(t) \quad (\text{Geschwindigkeit})$$

und die Beschleunigung gibt die Änderung der Geschwindigkeit an:

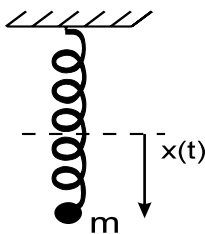
$$a(t) := \dot{v}(t) = \frac{d}{dt} v(t) = \ddot{s}(t). \quad (\text{Beschleunigung})$$

(1) Für den **freien Fall** gilt

$$s(t) = \frac{1}{2} g t^2 + v_0 t + s_0,$$

wenn s_0 die Anfangsposition und v_0 die Anfangsgeschwindigkeit. Es gilt hier für die Geschwindigkeit und Beschleunigung

$$\begin{aligned} v(t) &= \dot{s}(t) = g t + v_0, \\ a(t) &= \dot{v}(t) = g = \text{const.} \end{aligned}$$



(2) Ein durch Luftreibung **gedämpftes Federpendel** schwingt mit

$$x(t) = x_0 e^{-\gamma t} \cos(\omega t),$$

wenn x_0 die Anfangsauslenkung, γ der Reibungskoeffizient und ω die Schwingungsfrequenz. Die Geschwindigkeit und Beschleunigung sind

Abb. 7.4.

Federpendel $v(t) = \dot{x}(t) = -\gamma x_0 e^{-\gamma t} \cos(\omega t) - \omega x_0 e^{-\gamma t} \sin(\omega t),$

$$\begin{aligned} a(t) = \dot{v}(t) &= \gamma^2 x_0 e^{-\gamma t} \cos(\omega t) + \gamma \omega x_0 e^{-\gamma t} \sin(\omega t) \\ &\quad + \gamma \omega x_0 e^{-\gamma t} \sin(\omega t) - \omega^2 x_0 e^{-\gamma t} \cos(\omega t) \end{aligned}$$

$$a(t) = (\gamma^2 - \omega^2) x_0 e^{-\gamma t} \cos(\omega t) + 2\gamma\omega x_0 e^{-\gamma t} \sin(\omega t).$$

Für den Spezialfall ohne Reibung ($\gamma = 0$) ist

$$x(t) = x_0 \cos(\omega t)$$

und

$$\ddot{x}(t) = a(t) = -\omega^2 x_0 \cos(\omega t) = -\omega^2 x(t).$$

Dann ist die Rückstellkraft der Feder

$$F = m a = m \ddot{x}(t) = -m \omega^2 x(t) \sim x(t).$$

Dies ist das *Hooksche Gesetz*, welches besagt, dass die Rückstellkraft proportional zur Auslenkung $x(t)$ ist.

7.3.2 Induktionsgesetz

Das *Induktionsgesetz* aus der Physik lautet: Eine zeitliche Änderung des magnetischen Flusses ϕ induziert in einem Leiter mit Windungszahl n eine elektrische Spannung U_i gemäß

$$U_i(t) = -n \frac{d}{dt} \phi(t).$$

Dabei ist der magnetische Fluss $\phi = B \cdot A_{eff}$, B das angelegte Magnetfeld und A_{eff} die vom Magnetfeld durchdrungene Fläche.

(1) Wenden wir das Induktionsgesetz auf eine in einem konstanten Magnetfeld rotierende Spule an, so ist die effektiv vom Magnetfeld durchdrungene Fläche

$$A_{eff} = A \cos \varphi(t) = A \cos(\omega t).$$

Nach dem Induktionsgesetz wird die Wechselspannung

$$\begin{aligned} U_i &= -n \frac{d}{dt} \phi = -n \frac{d}{dt} B A \cos(\omega t) \\ &= n B A \omega \sin(\omega t) \end{aligned}$$

mit der Scheitelspannung $U_0 = n B A \omega$ induziert.

(2) Ist die Leiterschleife fest und variiert das Magnetfeld senkrecht zur Leiterschleife gemäß $B = B_0 \cos(\omega t)$ mit Amplitude B_0 und Frequenz ω , so wird in der Leiterschleife (Querschnittsfläche A) die Spannung U_i induziert gemäß der Formel

$$U_i = -n \frac{d}{dt} \phi = -n \frac{d}{dt} A B_0 \cos(\omega t) = n A B_0 \omega \sin(\omega t).$$

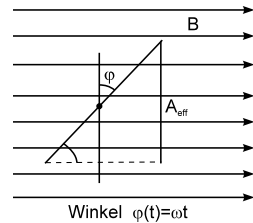
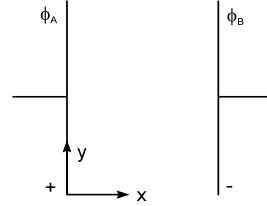


Abb. 7.5. Drehender Leiter

7.3.3 Elektrostatik

(1) In einem *Plattenkondensator* mit Anodenspannung ϕ_A , Kathodenspannung ϕ_K und Spaltabstand d ist das Potenzial $\phi(x)$ gegeben durch

$$\begin{aligned}\phi(x) &= \left(1 - \frac{x}{d}\right) \phi_A + \frac{x}{d} \phi_K \\ &= \phi_A + \left(\frac{x}{d}\right) (\phi_K - \phi_A).\end{aligned}$$



Das zugehörige *elektrische Feld* E ist definiert als

$$E(x) := -\frac{d}{dx} \phi(x) \quad (\text{elektrisches Feld}).$$

Dann ist

$$E(x) = -\frac{d}{dx} \phi(x) = -\frac{1}{d} (\phi_K - \phi_A) = \frac{\phi_A - \phi_K}{d} = \text{const.}$$

(2) **Kondensatormikrophon.** An den Platten eines Kondensatormikrophons liegt eine konstante Spannung U_0 an. Der Druck der Schallwellen (Frequenz ω) ändert den Plattenabstand nach der Formel

$$d = d_0 + a \sin(\omega t).$$

Damit variiert die Ladung Q am Kondensator

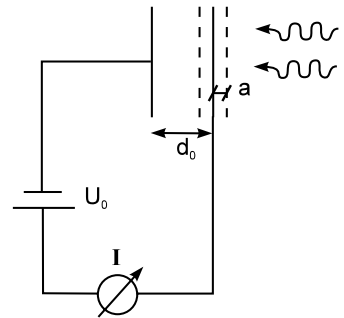
$$Q(t) = C(t) \cdot U_0 = \frac{\varepsilon_0 A}{d(t)} \cdot U_0,$$

da sich die Kapazität C zeitlich ändert. Der zeitliche Verlauf des Stromes

$$I(t) = \frac{d}{dt} Q(t)$$

ist demnach gegeben durch

$$\begin{aligned}I(t) &= \frac{d}{dt} Q(t) = \varepsilon_0 A U_0 \frac{d}{dt} \frac{1}{d_0 + a \sin(\omega t)} \\ &= -\varepsilon_0 A U_0 \frac{\omega a \cos(\omega t)}{[d_0 + a \sin(\omega t)]^2}.\end{aligned}$$



7.4 Differenzial einer Funktion

In diesem Abschnitt untersuchen wir das Verhalten von Funktionen, indem wir den Zuwachs der Funktion mit dem Zuwachs der Tangente vergleichen. Die Formeln bei der Linearisierung von Funktionen und der Fehlerrechnung basieren auf diesem Vergleich.

Um den Zuwachs einer differenzierbaren Funktion f in unmittelbarer Umgebung eines Punktes x_0 zu bestimmen, berechnen wir den Funktionswert an der Stelle x_0 und bei $x_0 + \Delta x$. Ändert sich der Abszissenwert um Δx , so ändert sich der Funktionswert um Δy . Für den **Zuwachs** Δy der Funktion f gilt

$$\Delta y = f(x_0 + \Delta x) - f(x_0).$$

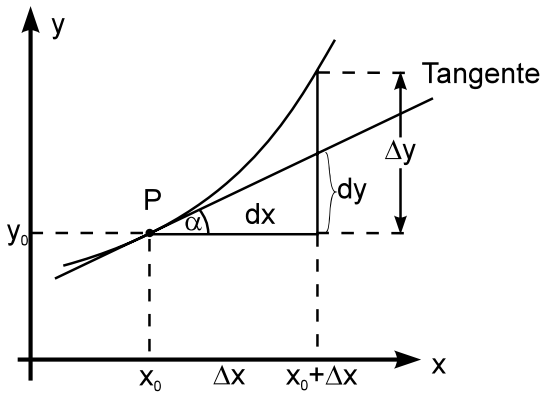


Abb. 7.6. Differenzial einer Funktion

Wir vergleichen den Zuwachs der Funktion mit der Änderung der Tangente dy . Man nennt

dx : unabhängiges Differenzial

dy : abhängiges Differenzial (= Änderung der Kurventangente)

Es gilt nach Abb. 7.6

$$\tan \alpha = f'(x_0) = \frac{dy}{dx} \Rightarrow dy = f'(x_0) \cdot dx.$$

Definition: (Differenzial einer Funktion). Das Differenzial

$$dy = df = f'(x_0) \cdot dx$$

einer Funktion $y = f(x)$ beschreibt die **Änderung längs der Tangente** im Punkte x_0 , wenn sich die Abszisse um dx ändert. df wird auch das Differenzial der Funktion f im Punkte x_0 bezeichnet.

Δy : **Änderung der Funktion** bei Änderung des x -Wertes um Δx .
 dy : **Änderung der Tangente** bei Änderung des x -Wertes um dx .

Beispiele 7.17:

- ① Gesucht ist das Differenzial der Funktion $f(x) = \sqrt{x+1}$ bei $x_0 = 0$:

$$f(x) = (x+1)^{\frac{1}{2}} \hookrightarrow f'(x) = \frac{1}{2}(x+1)^{-\frac{1}{2}} \hookrightarrow f'(x_0 = 0) = \frac{1}{2}.$$

Damit ist

$$df = f'(x_0) dx = \frac{1}{2} dx.$$

- ② Gesucht ist das Differenzial von $f(x) = \arctan x$ an der Stelle $x_0 = 0$:

$$f(x) = \arctan(x) \hookrightarrow f'(x) = \frac{1}{1+x^2} \hookrightarrow f'(x_0 = 0) = 1.$$

Damit ist

$$df = 1 \cdot dx.$$

□

7.4.1 Linearisierung von Funktionen

Aus dem Differenzial df einer Funktion werden wir nun noch eine für die Anwendungen wichtige Folgerung ziehen: Für kleine $\Delta x = dx$ gilt näherungsweise

$$\Delta y \approx dy.$$

Denn für kleine $\Delta x = dx$ ist die Änderung der Tangente vergleichbar mit der Änderung der Funktion. Setzt man noch die Formel für dy ein, erhält man

$$\Rightarrow \Delta y \approx dy = f'(x_0) \cdot \Delta x.$$

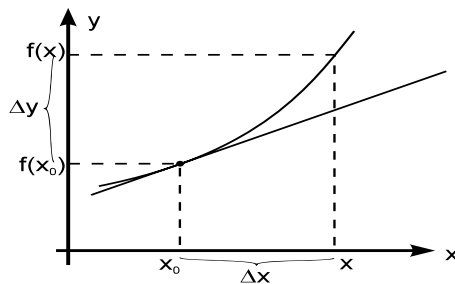


Abb. 7.7. Linearisierung einer Funktion

Für kleine Δx kann die Funktion f in unmittelbarer Umgebung des Punktes x_0 durch die Kurventangente ersetzt werden. Man nennt dieses Vorgehen die

Linearisierung der Funktion f . Wegen

$$\Delta y = f(x) - f(x_0) \approx dy = f'(x_0) (x - x_0)$$

folgt für die **Linearisierung** der Funktion f an der Stelle x_0 :

$$f(x) \approx f(x_0) + f'(x_0) \cdot (x - x_0).$$

Beispiele 7.18:

- ① Gegeben ist die Funktion $f(x) = \frac{4x}{x(x+1)}$. Gesucht ist die Linearisierung der Funktion bei $x_0 = 1$:

$$\text{Wegen } f'(x) = \frac{4x(x+1) - 4x(2x+1)}{x^2(x+1)^2} \quad \text{ist} \quad f'(x_0 = 1) = -1$$

$$\Rightarrow f(x) \approx f(x_0) + f'(x_0) (x - x_0) = 2 + (-1) (x - 1) = 3 - x.$$

- ② Gegeben ist die Funktion $f(\varphi) = \sin \varphi$. Gesucht ist die Linearisierung der Funktion bei $\varphi_0 = 0$:

$$\text{Wegen } f'(\varphi) = \cos \varphi \quad \text{ist} \quad f'(\varphi_0 = 0) = 1$$

$$\Rightarrow f(\varphi) \approx f(\varphi_0) + f'(\varphi_0) (\varphi - \varphi_0) = 0 + 1 \cdot (\varphi - 0) = \varphi$$

$$\Rightarrow \boxed{\sin \varphi \approx \varphi} \quad (\text{für kleine Winkel}).$$

□

Anwendungsbeispiel 7.19 (Harmonisches Pendel).

An einem Faden der Länge l sei eine Masse m befestigt. Gesucht ist der Winkel $\varphi(t)$ als Funktion der Zeit, wenn die Masse um einen kleinen Winkel φ_0 ausgelenkt wird. Wir bestimmen alle auf die Masse wirkenden Kräfte, denn nach dem *Newtonschen Bewegungsgesetz* ist die Beschleunigungskraft gegeben durch die Summe aller Kräfte.

Gehen wir von einer reibungsfreien Bewegung aus, wirkt auf den Massenpunkt nur die Gewichtskraft $F_G = m \cdot g$. Da die Beschleunigung der Masse senkrecht zum Faden erfolgt, wirkt als Beschleunigungskraft

$$F_B = -F_G \cdot \sin \varphi = -m \cdot g \sin \varphi.$$

Bei einer Kreisbewegung ist die Geschwindigkeit $v(t) = l\omega = l \cdot \dot{\varphi}(t)$ und damit die Beschleunigung $a(t) = \dot{v}(t) = l\ddot{\varphi}(t)$. Insgesamt gilt nach dem

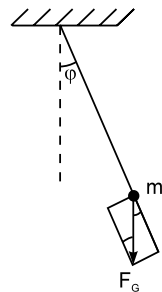


Abb. 7.8.
Fadenpendel

Newtonschen Gesetz

$$m l \ddot{\varphi}(t) = -m g \sin(\varphi(t))$$

bzw.

$$\ddot{\varphi}(t) + \frac{g}{l} \sin \varphi(t) = 0.$$

Diese Gleichung nennt man *Differenzialgleichung*, da neben der gesuchten Funktion $\varphi(t)$ (= Winkelauslenkung zu einem beliebigen Zeitpunkt t) auch deren (zweite) Ableitung vorkommt. Für diese Gleichung lässt sich keine geschlossene Lösung $\varphi(t)$ finden. Für kleine Winkel $\varphi(t) \approx 0$ gilt aber nach Beispiel 7.18 ②

$$\sin \varphi(t) \approx \varphi(t),$$

so dass die Bewegungsgleichung sich vereinfacht zu

$$\ddot{\varphi}(t) + \frac{g}{l} \varphi(t) = 0.$$

Die Lösung dieser Differenzialgleichung ist gegeben durch

$$\varphi(t) = \varphi_0 \cos(\omega t)$$

mit $\omega = \sqrt{\frac{g}{l}}$ und $\varphi_0 = \varphi(t=0)$ die Anfangsauslenkung. Dies prüft man durch Einsetzen der Funktion $\varphi(t)$ in die Differenzialgleichung direkt nach:

$$\begin{aligned} \ddot{\varphi}(t) &= -\omega^2 \varphi_0 \cos(\omega t) \\ \hookrightarrow \ddot{\varphi}(t) + \frac{g}{l} \varphi(t) &= -\omega^2 \varphi_0 \cos(\omega t) + \frac{g}{l} \varphi_0 \cos(\omega t) = 0. \end{aligned}$$

Zusammenfassung: Für kleine Winkelauslenkungen ist der Winkel $\varphi(t)$ eines Fadenpendels gegeben durch $\varphi(t) = \varphi_0 \cos(\omega t)$. Dies ist eine periodische Bewegung mit zeitlich konstanter Amplitude φ_0 , fester Kreisfrequenz $\omega = \sqrt{\frac{g}{l}}$ und zeitlich konstanter Schwingungsdauer $T = \frac{2\pi}{\omega} = 2\pi\sqrt{\frac{l}{g}}$. $\square$

➤ 7.4.2 Fehlerrechnung

Das Differenzial einer Funktion macht eine Aussage darüber, wie sich der Wert einer Funktion bei kleinen Änderungen der Argumente ändert. Eine Anwendung des Differenzials ist die Fehlerrechnung, da diese eine Aussage darüber trifft, in wieweit sich der Funktionswert in linearer Näherung maximal ändert, wenn das Argument nur bis auf $\pm dx$ bekannt ist.

Fehlerrechnung: Sei f eine physikalische Größe, die durch Messung einer Einzelgröße x berechnet wird. Sei die Einzelgröße x mit der Messungengenauigkeit dx behaftet. Dann ist der **maximale Fehler in linearer Näherung** gegeben durch den Betrag des Differenzials

$$|df| = |f'(x)| |dx|.$$

Anwendungsbeispiel 7.20 (Wheatstonesche Brückenschaltung).

Über die *Wheatstonesche Brückenschaltung* lässt sich ein unbekannter Ohmscher Widerstand R durch

$$R = R_0 \cdot \frac{x}{l - x}$$

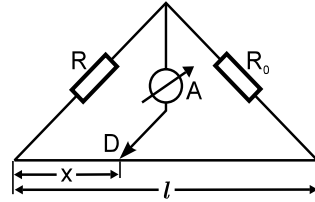


Abb. 7.9. Wheatstonesche Brücke

bestimmen, wenn der Abtastpunkt D auf dem Draht so verschoben wird, dass die Brücke stromlos ist. Aufgrund von Messungenauigkeiten ist x nur bis auf $\pm 1\text{mm}$ bekannt. Wie groß ist der maximale und relative Fehler in linearer Näherung, wenn $R_0 = 100\Omega$, $l = 1\text{m}$ und $x = 0.53\text{m}$?

Wir betrachten R als Funktion von x :

$$R(x) = R_0 \frac{x}{l - x}$$

und bestimmen das Differenzial

$$dR = R'(x) \cdot dx = R_0 \frac{1 \cdot (l - x) + x \cdot 1}{(l - x)^2} dx = -\frac{R_0 l}{(l - x)^2} dx.$$

Für $x_0 = 0.53\text{m}$ und $dx = 0.001\text{m}$ folgt für den maximalen Fehler in linearer Näherung

$$|dR| = 100 \frac{1}{0.47^2} \cdot 0.001\Omega = 0.452\Omega.$$

Der Widerstand ergibt sich also zu

$$R = R_0 \frac{x_0}{l - x_0} \pm |dR| = (112.76 \pm 0.45)\Omega$$

mit einem *relativen Fehler* von $\frac{|dR|}{R} = 0.4\%$. □

7.5 Anwendungen in der Mathematik

Aus den Ableitungen einer Funktion werden Rückschlüsse auf die Eigenschaften der Funktion gezogen. Aus der ersten Ableitung kann man die Monotonie und die Extremwerte ablesen; aus der zweiten Ableitung die Wendepunkte. Ziel dieses Abschnittes ist, eine vollständige Kurvendiskussion durchzuführen. Insbesondere die Charakterisierung der lokalen Extremwerte über die erste Ableitung werden wir im nächsten Abschnitt anwenden, um Optimierungsprobleme zu lösen.

7.5.1 Geometrische Bedeutung von f' und f'' .

Die Ableitung einer Funktion gibt die Steigung der Kurventangente an. Hat eine Funktion stets eine positive Steigung, $f'(x) > 0$, dann ist der Zuwachs $df = f'(x) \cdot dx > 0$ und die Funktion streng monoton wachsend (Abb. 7.10 (a)). Ist andererseits für eine Funktion $f'(x) < 0$, dann fällt die Funktion streng monoton (Abb. 7.10 (b)).

Satz: Monotonieverhalten einer Funktion.

Sei f eine in ihrem Definitionsbereich $\mathbb{D}$ differenzierbare Funktion.

- (1) Ist $f'(x) > 0$ für alle $x \in \mathbb{D} \Rightarrow f$ ist streng monoton wachsend.
- (2) Ist $f'(x) < 0$ für alle $x \in \mathbb{D} \Rightarrow f$ ist streng monoton fallend.

Die zweite Ableitung einer Funktion bestimmt das Krümmungsverhalten, denn f'' gibt die Steigung der Tangente an. Es gilt

Satz: Krümmungsverhalten einer Funktion.

Für eine in ihrem Definitionsbereich $\mathbb{D}$ zweimal differenzierbare Funktion gilt für $x_0 \in \mathbb{D}$:

- (1) Ist $f''(x_0) > 0$, dann nimmt die Steigung der Kurventangente beim Durchgang durch den Punkt $P(x_0, f(x_0))$ zu. Die Tangente dreht sich in positive Richtung (Gegenuhrzeigersinn). Die Kurve besitzt in P eine *Linkskrümmung* (Abb. 7.10 (b)).
- (2) Ist $f''(x_0) < 0$, dann nimmt die Steigung der Kurventangente beim Durchgang durch den Punkt $P(x_0, f(x_0))$ ab. Die Tangente dreht sich in negative Richtung (Uhrzeigersinn). Die Kurve besitzt in P eine *Rechtskrümmung* (Abb. 7.10 (a)).

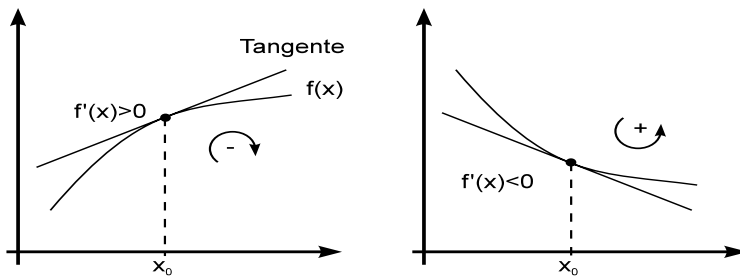


Abb. 7.10. (a) monoton wachsende Funktion mit Rechtskrümmung, (b) monoton fallende Funktion mit Linkskrümmung

7.5.2 Relative Extremalwerte

Die *relativen Extremalwerte* geben die Punkte einer Funktion an, bei denen die Funktionswerte in einer Umgebung relativ einen größten bzw. kleinsten Funktionswert besitzen.

Definition: Eine Funktion f besitzt an der Stelle $x_0 \in \mathbb{D}$ ein **relatives Maximum** bzw. **relatives Minimum**, wenn in einer Umgebung des Punktes x_0 gilt

$$f(x_0) > f(x) \quad \text{für alle } x \neq x_0 \quad (\text{rel. Maximum})$$

$$f(x_0) < f(x) \quad \text{für alle } x \neq x_0 \quad (\text{rel. Minimum}).$$

Eine differenzierbare Funktion besitzt in einem lokalen Extremum eine waagrechte Tangente. Somit ist eine notwendige Bedingung für einen relativen Extremwert x_0 , dass $f'(x_0) = 0$.

Achtung: Diese Bedingung ist jedoch nicht hinreichend, wie das folgende Beispiel zeigt:

Beispiel 7.21. Für die Funktion $f(x) = x^3$ gilt $f'(0) = 0$; aber $x_0 = 0$ ist kein lokales Extremum. Diese Funktion ändert bei $x_0 = 0$ gleichzeitig ihre Krümmung (von einer Rechts- in eine Linkskrümmung). Daher ist $f''(0) = 0$. Damit also zwingend eine waagrechte Tangente auch hinreichend für ein Extremum ist, darf die Kurve in x_0 die Krümmung nicht ändern: $f''(x_0)$ darf nicht verschwinden. $\square$

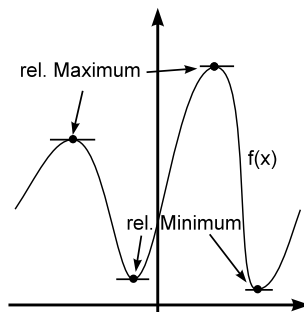


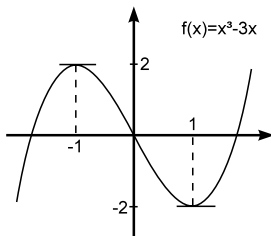
Abb. 7.11. Relative Extremwerte

Satz: Hinreichende Bedingung für ein relatives Extremum.

Sei f in einer Umgebung von $x_0 \in \mathbb{D}$ zweimal stetig differenzierbar. Ist $f'(x_0) = 0$ und $f''(x_0) \neq 0$, dann besitzt f in x_0 ein relatives Extremum.

kurz: $f'(x_0) = 0, f''(x_0) < 0 \Rightarrow x_0$ ist ein relatives Maximum.

$f'(x_0) = 0, f''(x_0) > 0 \Rightarrow x_0$ ist ein relatives Minimum.

Beispiele 7.22:

① Gesucht sind die relativen Extrema der Funktion

$$f(x) = x^3 - 3x.$$

Es ist $f'(x) = 3x^2 - 3$, $f''(x) = 6x$.

Aus $f'(x) = 0$, folgt $3x^2 - 3 = 0 \Leftrightarrow x_{1/2} = \pm 1$.

An diesen Stellen besitzt die Funktion eine waagrechte Tangente.

i) $f''(x_1) = f''(1) > 0 \Rightarrow$ bei $x_1 = 1$ ist ein

lokales Minimum,

ii) $f''(x_2) = f''(-1) < 0 \Rightarrow$ bei $x_2 = -1$ ist ein lokales Maximum.

② Für die Exponentialfunktion $f(x) = e^x$ gilt für alle $x \in \mathbb{R}$: $f'(x) = e^x > 0$. Die Exponentialfunktion besitzt also kein Extremum und ist auf ganz $\mathbb{R}$ streng monoton wachsend. $\square$

➤ 7.5.3 Wendepunkte und Sattelpunkte**Definition:**

- (1) Kurvenpunkte, in denen sich der Drehsinn der Tangente ändert, heißen **Wendepunkte** (Abb. 7.12 (a)).
- (2) Wendepunkte, die eine waagrechte Tangente besitzt, heißen **Sattelpunkte** (Abb. 7.12 (b)).

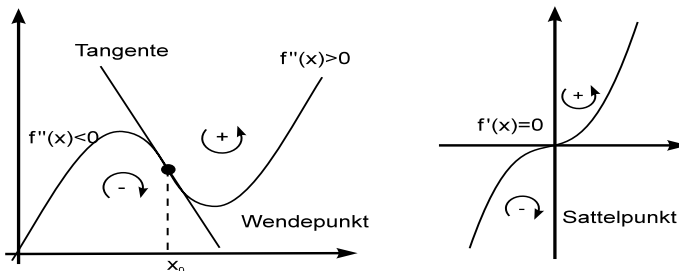


Abb. 7.12. Wende- und Sattelpunkte

In den Wendepunkten findet eine Änderung der Krümmung statt. Für eine zweimal stetig differenzierbare Funktion gilt in diesen Punkten $f''(x_0) = 0$. Diese Bedingung ist aber nicht hinreichend wie das Beispiel $f(x) = x^4$ zeigt: Bei $x_0 = 0$ gilt zwar $f''(x_0) = 0$, aber auch $f'''(x_0) = 0$. Folgende Bedingungen sind hinreichend:

Satz: Hinreichende Bedingung für Wende- und Sattelpunkte.

Sei die Funktion f in ihrem Definitionsbereich 3-mal stetig differenzierbar und $x_0 \in \mathbb{D}$.

- (1) Ist $f''(x_0) = 0$ und $f'''(x_0) \neq 0 \Rightarrow x_0$ ist **Wendepunkt**.
- (2) Ist $f''(x_0) = 0$, $f'''(x_0) \neq 0$ und $f'(x_0) = 0 \Rightarrow x_0$ ist **Sattelpunkt**.

Beispiele 7.23:

- ① Die Funktion $f(x) = x^3$ hat in $x_0 = 0$ einen Sattelpunkt:

$$f'(x) = 3x^2, \quad f''(x) = 6x, \quad f'''(x) = 6.$$

Wegen $f'(0) = f''(0) = 0$ und $f'''(0) \neq 0$ ist $x_0 = 0$ ein Sattelpunkt.

- ② Die Funktion $f(x) = x^3 - 3x$ hat in $x_0 = 0$ einen Wendepunkt:

$$f'(x) = 3x^2 - 3, \quad f''(x) = 6x, \quad f'''(x) = 6.$$

Aus $f''(x) = 0 \hookrightarrow x_0 = 0$. Wegen $f'''(0) = 6 \neq 0$ ist $x_0 = 0$ ein Wendepunkt.

- ③ Die Wendepunkte der Sinusfunktion fallen mit den Nullstellen zusammen:

$$f(x) = \sin x, \quad f'(x) = \cos x, \quad f''(x) = -\sin x, \quad f'''(x) = -\cos x.$$

Aus $f''(x) = -\sin x = 0$ folgt $x_k = k \cdot \pi$. Wegen

$$f'''(x_k) = -\cos(k\pi) = -(-1)^k \neq 0$$

sind damit alle Nullstellen auch Wendepunkte. □

7.5.4 Kurvendiskussion

Charakteristische Eigenschaften einer Funktion werden mit Hilfe der Differenzialrechnung erfasst. Charakteristische Punkte sind dabei Nullstellen und Polstellen aber auch die relativen Extrema, Wende- und Sattelpunkte. Eine vollständige *Kurvendiskussion* orientiert sich an den folgenden 10 Punkten:

1. Definitionsbereich
2. Symmetrieverhalten
3. Nullstellen
4. Polstellen
5. Asymptoten (Verhalten der Funktion für $x \rightarrow \pm \infty$)
6. Ableitungen der Funktion (bis zur Ordnung 3)
7. Relative Extrema
8. Wende- und Sattelpunkte
9. Wertebereich
10. Funktionsgraph

Beispiel 7.24. Gebrochenrationale Funktion $y = \frac{-5x^2 + 5}{x^3}$.

Definitionsbereich: Der Nenner wird bei $x_0 = 0$ Null $\Rightarrow \mathbb{D} = \mathbb{R} \setminus \{0\}$. $x_0 = 0$ ist Definitionslücke.

Symmetrieverhalten: Die Funktion ist punktsymmetrisch, da sie der Quotient einer geraden und ungeraden Funktion und damit insgesamt ungerade ist.

Nullstellen: Zähler und Nenner werden in Linearfaktoren zerlegt und gemeinsame Faktoren gekürzt:

$$y = \frac{-5x^2 + 5}{x^3} = \frac{-5(x+1)(x-1)}{x^3}$$

$\Rightarrow$ Nullstellen sind $x_1 = -1$ und $x_2 = 1$.

Polstellen: $x_3 = 0$.

Asymptoten: Der Grad des Zählers ist kleiner als der Grad des Nenners
 $\Rightarrow \lim_{x \rightarrow \pm \infty} y = 0 \hookrightarrow x$ -Achse ist Asymptote.

Ableitungen:

$$y' = \frac{5(x^2 - 3)}{x^4},$$

$$y'' = \frac{-10(x^2 - 6)}{x^5},$$

$$y''' = \frac{30(x^2 - 10)}{x^6}.$$

Relative Extrema: $y' = 0 \hookrightarrow x^2 - 3 = 0 \Rightarrow x_{4/5} = \pm\sqrt{3}$.

$y''(\sqrt{3}) = \frac{10}{9}\sqrt{3} > 0 \Rightarrow (\sqrt{3}, -\frac{10}{9}\sqrt{3})$ ist relatives Minimum.

$y''(-\sqrt{3}) = -\frac{10}{9}\sqrt{3} < 0 \Rightarrow (-\sqrt{3}, \frac{10}{9}\sqrt{3})$ ist relatives Maximum.

Wendepunkte: $y'' = 0 \hookrightarrow x^2 - 6 = 0 \Rightarrow x_{6/7} = \pm\sqrt{6}$.

$y'''(\pm\sqrt{6}) \neq 0 \Rightarrow x_{6/7}$ sind Wendepunkte.

Funktionsgraph:

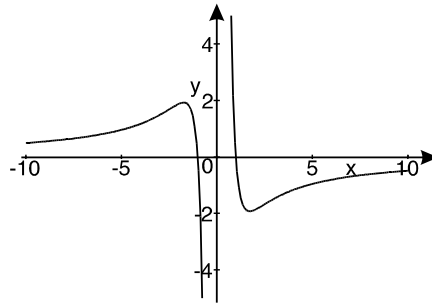


Abb. 7.13. Funktionsgraph von $y = \frac{-5x^2 + 5}{x^3}$

Wertebereich: $\mathbb{W} = \mathbb{R}$.

□

Anwendungsbeispiel 7.25 (Gedämpfte freie Schwingung).

Gesucht ist der Funktionsverlauf der gedämpften freien Schwingung

$$x(t) = 4e^{-0.1t} \cdot \cos(2t) \quad t \geq 0.$$

Definitionsbereich: An die Funktion ergeben sich keine Einschränkungen. Aufgrund von physikalischen Gründen ist jedoch $\text{ID} = \mathbb{R}_{\geq 0}$.

Symmetrie: - (da für negative t keine Kurvendiskussion)

Nullstellen: $x(t) = 0 \hookrightarrow e^{0.1t} \cos(2t) = 0 \hookrightarrow \cos(2t) = 0$
 $\Rightarrow 2t = \frac{\pi}{2} + k\pi \hookrightarrow t_k = \frac{\pi}{4} + k\frac{\pi}{2}, \quad k = 0, 1, 2, 3, \dots$

Polstellen: -

Asymptoten: Für $t \rightarrow \infty$ ist $x(t) \rightarrow 0$. Die t -Achse ist Asymptote.

Ableitungen:

$$\begin{aligned}\dot{x}(t) &= (-0.4 \cos(2t) - 8 \sin(2t)) e^{-0.1t}, \\ \ddot{x}(t) &= (-15.96 \cos(2t) + 1.6 \sin(2t)) e^{-0.1t}, \\ \dddot{x}(t) &= (4.796 \cos(2t) + 31.76 \sin(2t)) e^{-0.1t}.\end{aligned}$$

Relative Extrema: $\dot{x}(t) = 0 \hookrightarrow -0.4 \cos(2t) - 8 \sin(2t) = 0$

$$\hookrightarrow \tan(2t) = -\frac{0.4}{8} = -0.05$$

 $\hookrightarrow 2t = \arctan(-0.05) + k \cdot \pi$, da der Arkustangens mehrdeutig

$$\hookrightarrow \boxed{2t_k = 3.091 + k \cdot \pi} \quad k = 0, 1, 2, 3, \dots$$

(i) Für gerade $k = 2n$ ist $\ddot{x}(t_k)$ positiv:

$$\begin{aligned}\ddot{x}(t_k) &= \ddot{x}(t_{2n}) \\ &= (-15.96 \cos(3.091 + 2n\pi) + 1.6 \sin(3.091 + 2n\pi)) e^{-0.1t_{2n}} \\ &= 16.020 e^{-0.1t_{2n}} > 0 \Rightarrow \text{rel. Minima:}\end{aligned}$$

$$\text{Min}_1 = (1.545, -3.422); \text{Min}_2 = (4.687, -2.500); \text{Min}_3 = (7.828, -1.826); \dots$$

(ii) Für ungerade $k = 2n + 1$ ist $\ddot{x}(t_k)$ negativ:

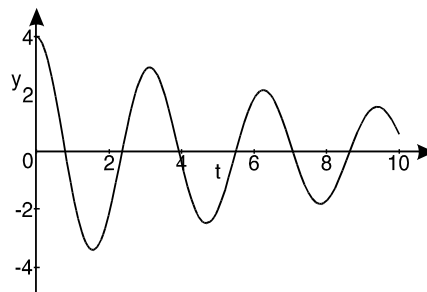
$$\begin{aligned}\ddot{x}(t_k) &= \ddot{x}(t_{2n+1}) \\ &= (-15.96 \cos(3.091 + \pi) + 1.6 \sin(3.091 + \pi)) e^{-0.1t_{2n+1}} \\ &= -16.020 e^{-0.1t_{2n+1}} < 0 \Rightarrow \text{rel. Maximum:}\end{aligned}$$

$$\text{Max}_1 = (3.116, 2.925); \text{Max}_2 = (6.257, 2.136); \text{Max}_3 = (9.399, 1.560); \dots$$

Der Abstand von Minima und Maxima beträgt $\frac{\pi}{2}$.**Wendepunkte:** $\ddot{x}(t) = 0 \hookrightarrow -15.96 \cos(2t) + 1.6 \sin(2t) = 0$

$$\Rightarrow \tan(2t) = \frac{15.96}{1.6} \Rightarrow \boxed{2t_k = 4.612 + k \cdot \pi}.$$

Da die dritte Ableitung an diesen Stellen abwechselnd positiv und negativ ist, liegen tatsächlich Wendepunkte vor.

Funktionsgraph:Abb. 7.14. Funktionsgraph von $x(t) = 4 \exp(-0.1t) \cos(2t)$.**Wertebereich:** $\text{WV} = [-3.422, 4]$; der größte Wert wird bei $t = 0$, der kleinste im 1. Minimum angenommen. $\square$

7.6 Extremwertaufgaben (Optimierungsprobleme)

Die Differenzialrechnung ist ein wichtiges Hilfsmittel zur Lösung von Optimierungsproblemen in der Technik. Bei den Optimierungsproblemen sucht man das absolute Maximum oder Minimum eines Problems unter Nebenbedingungen.

In der Regel werden technische Sachverhalte in ein mathematisches Modell (Zielfunktion) umgesetzt und die Lösung des Problems auf die Bestimmung von Extremwerten dieser Funktion zurückgeführt. Das gesuchte absolute Maximum (oder Minimum) kann aber auch in einem Randpunkt des Intervalls liegen. Für Extremwertaufgaben wählt man folgendes Schema:

- (1) Ansatz: Die zu optimierende Größe wird mathematisch beschrieben. Aufgrund von Randbedingungen werden alle Variablen bis auf eine eliminiert. Die zu optimierende Größe ist dann eine Funktion (*Zielfunktion*) von einer Variablen.
- (2) Die Extrema der Zielfunktion werden ermittelt und deren Verträglichkeit mit dem gegebenen Problem überprüft.
- (3) Durch Vergleich der Extremwerte mit den Funktionswerten in den Randpunkten des Lösungsintervalls erhält man den größten (oder kleinsten) Wert der Funktion im Intervall.

Anwendungsbeispiel 7.26 (RCL-Wechselstromkreis).

In einem RCL-Wechselstromkreis fließt beim Anlegen einer Wechselspannung

$$U(t) = U_0 \sin(\omega t)$$

ein Wechselstrom

$$I(t) = I_0 \sin(\omega t + \varphi)$$

mit frequenzabhängiger Amplitude

$$I_0(\omega) = U_0 / \sqrt{R^2 + \left(\omega L - \frac{1}{\omega C}\right)^2}$$

Bei welcher Frequenz ω besitzt I_0 seinen größten Wert?

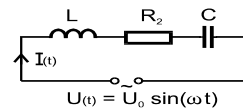


Abb. 7.15. RCL-Kreis

$I_0(\omega)$ ist maximal, wenn $\sqrt{R^2 + (\omega L - \frac{1}{\omega C})}$ minimal, bzw. wenn die folgende Funktion f minimal wird:

$$\begin{aligned} f(\omega) &= R^2 + (\omega L - \frac{1}{\omega C})^2 \\ f'(\omega) &= 2(\omega L - \frac{1}{\omega C}) \cdot (L + \frac{1}{\omega^2 C}) \stackrel{!}{=} 0 \quad \Rightarrow \omega_0 = \frac{1}{\sqrt{LC}} \\ f''(\omega) &= 2(L + \frac{1}{\omega^2 C})^2 - \frac{4}{\omega^3 C}(\omega L - \frac{1}{\omega C}) \quad \Rightarrow f''\left(\frac{1}{\sqrt{LC}}\right) = 8L^2 > 0. \end{aligned}$$

f nimmt in $\omega_0 = \frac{1}{\sqrt{LC}}$ sein Minimum an, und $I_0(\omega)$ hat den Maximalwert $I_0(\omega_0) = \frac{U_0}{R}$.

Wegen $I_0(\omega) \xrightarrow{\omega \rightarrow 0} 0$ und $I_0(\omega) \xrightarrow{\omega \rightarrow \infty} 0$, stellt das relative Maximum auch das absolute Maximum dar. I_0 hat also sein Maximum bei der Resonanzfrequenz $\omega_0 = \frac{1}{\sqrt{LC}}$. Der Scheinwiderstand ist dann gleich dem Ohmschen Widerstand R und es gilt

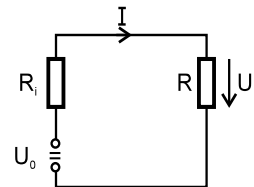
$$I_{0,\max} = \frac{U_0}{R}.$$

□

Anwendungsbeispiel 7.27 (Leistungsanpassung des Widerstandes).

An einer Gleichstromquelle U_0 mit Innenwiderstand R_i ist ein Außenwiderstand R (Lastwiderstand) angeschlossen. Wie groß muss der äußere Widerstand gewählt werden, damit die Nutzleistung maximal wird?

$$\left. \begin{aligned} \text{Leistung am Widerstand } R : \\ P &= U \cdot I = R \cdot I \cdot I \\ \text{Maschensatz (Nebenbedingung):} \\ U_0 &= U_i + U = R_i I + R I = (R_i + R) I \end{aligned} \right\}$$



$$\Rightarrow I = \frac{U_0}{R_i + R}.$$

Damit ergibt sich die Nutzleistung als Funktion von R :

$$P(R) = R \left(\frac{U_0}{R_i + R} \right)^2 = U_0^2 \frac{R}{(R_i + R)^2} \quad (\text{Nutzleistung am Widerstand})$$

Für das Extremum gilt

$$P'(R) = U_0^2 \frac{1 \cdot (R_i + R)^2 - R \cdot 2(R_i + R)}{(R_i + R)^4} = U_0^2 \frac{R_i - R}{(R_i + R)^3} \stackrel{!}{=} 0$$

$$\Rightarrow \boxed{R = R_i}.$$

Für die zweite Ableitung von $P(R)$ folgt

$$\begin{aligned} P''(R) &= U_0^2 \frac{-1(R_i + R)^3 - (R_i - R) 3(R_i + R)^2}{(R_i + R)^6} \\ &= U_0^2 \frac{-(R_i + R) - 3(R_i - R)}{(R_i + R)^4}. \end{aligned}$$

Damit ist

$$P''(R_i) = U_0^2 \frac{-2R_i}{16} = -\frac{1}{8} \frac{U_0^2}{R_i^3} < 0.$$

Für $\boxed{R = R_i}$ liegt ein relatives Extremum vor. Da $P(R) \xrightarrow{R \rightarrow \infty} 0$ und $P(R) \xrightarrow{R \rightarrow 0} 0$ ist das relative auch das absolute Maximum. Die Nutzleistung wird maximal, wenn der Lastwiderstand gleich dem Innenwiderstand gewählt wird. Die maximale Nutzleistung bei diesem äußeren Widerstand beträgt dann

$$\boxed{P_{\max} = \frac{1}{4} \frac{U_0^2}{R_i}}. \quad (1)$$

Man nennt dies auch die Spannungsteilerregel: Wenn der Innenwiderstand R_i gleich dem äußeren Widerstand R_a (=Lastwiderstand) gewählt wird, dann ist der Spannungsabfall bei R_i gleich dem Spannungsabfall bei R_a , nämlich $\frac{U_0}{2}$. $\square$

Anwendungsbeispiel 7.28 (Balkenbiegung).

Die Biegelinie eines einseitig eingespannten Balkens der Länge L unter dem Einfluss einer Kraft F lautet näherungsweise

$$\boxed{y(x) = \frac{F}{2EI} \left(Lx^2 - \frac{1}{3}x^3 \right) \quad (0 \leq x \leq L),}$$



Abb. 7.16. Balkenbiegung

wenn E der Elastizitätsmodul und I das Flächenträgheitsmoment ist. An welcher Stelle ist die Balkenbiegung am größten?

Wir bestimmen die relativen Extrema der Biegung $y(x)$:

$$\begin{aligned} y'(x) &= \frac{F}{2EI} (2Lx - x^2) = 0 \Rightarrow x_1 = 0 \quad \text{oder} \quad x_2 = 2L \\ y''(x) &= \frac{F}{2EI} (2L - 2x) \end{aligned}$$

und setzen $x_1 = 0$, $x_2 = 2L$ in die 2. Ableitung ein: $y''(0) = \frac{FL}{EI} > 0 \Rightarrow$ rel. Minimum; $y''(2L) = -\frac{FL}{EI} < 0 \Rightarrow$ rel. Maximum. x_2 liegt aber außerhalb des physikalischen Bereichs.

△ Die **maximale** Durchbiegung des Balkens (absolutes Maximum) findet allerdings am freien Ende ($x = L$) statt:

$$y_{\max} = y(L) = \frac{FL^3}{3EI}. \quad \square$$

Anwendungsbeispiel 7.29 (Magnetfeld von Leiterschleifen, mit [MAPLE-Worksheet](#)).

Das durch eine stromdurchflossene Leiterschleife erzeugte Magnetfeld ist auf der Achse der Leiterschleife gegeben durch die Formel

$$B(z) = \frac{\mu_0 I R^2}{2 (R^2 + z^2)^{\frac{3}{2}}},$$

wenn $R (= 0.1m)$ der Radius der Leiterschleife, $I (= 1A)$ der Strom und $\mu_0 (= 4\pi \cdot 10^{-7} H/m)$ die Permeabilität von Vakuum. Die Effekte der Stromzuleitung werden vernachlässigt.

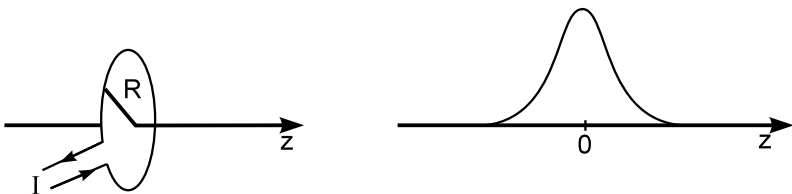


Abb. 7.17. (a) Stromdurchflossene Leiterschleife und (b) Magnetfeld auf der Achse

Das Magnetfeld von zwei stromdurchflossenen Leiterschleifen ist die Überlagerung der Einzelmagnetfelder. Gesucht ist der Abstand d der Leiterschleifen, so dass das Magnetfeld zwischen den Einzelschleifen möglichst homogen (= gleichförmig) wird.

Das Gesamtmagnetfeld zweier stromdurchflossener Leiter ergibt sich aus der Superposition der Einzelmagnetfelder. Beträgt der Abstand der Leiterschleifen

d , dann berechnen sich das Gesamtmagnetfeld auf der Achse durch

$$B = \frac{1}{2} \mu_0 I R^2 \left(\frac{1}{\left(R^2 + \left(z - \frac{1}{2}d\right)^2\right)^{3/2}} + \frac{1}{\left(R^2 + \left(z + \frac{1}{2}d\right)^2\right)^{3/2}} \right),$$

wenn die erste Leiterschleife bei $z = \frac{d}{2}$ und die zweite bei $z = -\frac{d}{2}$ liegt.

Die mathematische Bedingung für die Homogenität des Magnetfeldes bei $z = 0$ ist, dass $B''(z = 0) = 0$. Wir setzen daher $z = 0$ und differenzieren die Funktion $B(d)$ zweimal nach d :

$$\begin{aligned} B''(d) &= \frac{15}{16} \mu_0 I R^2 \frac{d^2}{\left(R^2 + \frac{1}{4}d^2\right)^{7/2}} - \frac{3}{4} \mu_0 I R^2 \frac{1}{\left(R^2 + \frac{1}{4}d^2\right)^{5/2}} \\ &= 96 \mu_0 I R^2 \frac{d^2 - R^2}{(4R^2 + d^2)^{7/2}} \\ &= 0. \end{aligned}$$

Lösen wir obigen Ausdruck nach d auf, erhalten wir $d = R$: **Ist der Spulenabstand d gleich dem Radius R der Spulen, so ist das Magnetfeld auf der Achse homogen (Helmholtz-Spulen)!**



Im zugehörigen [MAPLE-Worksheet](#) ist eine Animation dargestellt, bei der die Abstände d variieren. Es zeigt sich, dass bei $d = R$ das Magnetfeld relativ homogen verläuft. $\square$

7.7 Sätze der Differenzialrechnung

7.7

In diesem Kapitel werden einige wichtige Sätze der Differenzialrechnung zusammengefasst. Als zentraler Satz der Differenzialrechnung steht der Mittelwertsatz. Er findet u.a. bei den Regeln von *l'Hospital* seine Anwendung.

Der folgende Satz besagt, dass die Differenzierbarkeit einer Funktion eine stärkere Eigenschaft ist als deren Stetigkeit:

Satz: Ist eine Funktion $f : \text{ID} \rightarrow \mathbb{R}$ an der Stelle $x_0 \in \text{ID}$ differenzierbar, dann ist f an dieser Stelle auch stetig.

Beweis: Sei f eine in x_0 differenzierbare Funktion, dann gilt mit den Limesrechenregeln:

$$\begin{aligned}\lim_{h \rightarrow 0} (f(x_0 + h) - f(x_0)) &= \lim_{h \rightarrow 0} [h \cdot (f(x_0 + h) - f(x_0)) / h] \\ &= \lim_{h \rightarrow 0} h \cdot \lim_{h \rightarrow 0} \frac{1}{h} (f(x_0 + h) - f(x_0)) \\ &= 0 \cdot f'(x_0) = 0 \\ \Rightarrow \lim_{h \rightarrow 0} f(x_0 + h) &= f(x_0). \quad \square\end{aligned}$$

⚠ Die Umkehrung dieses Satzes gilt nicht, da z.B. die Betragsfunktion $\text{abs}(x) = |x|$ an der Stelle $x_0 = 0$ stetig, aber nicht differenzierbar ist!

➤ 7.7.1 Satz über die Exponentialfunktion

Eine für die Anwendungen bei den Differenzialgleichungen wichtige Eigenschaft der Exponentialfunktion ist, dass die Ableitung wieder e^x ergibt:

$$(e^x)' = e^x.$$

Die Exponentialfunktion ist die einzige Funktion, welche diese Eigenschaft besitzt:

Satz über die Exponentialfunktion:

Ist a eine Konstante und $f : \mathbb{R} \rightarrow \mathbb{R}$ eine differenzierbare Funktion mit $f'(x) = a f(x)$ für alle $x \in \mathbb{R}$.

Dann gilt:

$$f(x) = f(0) e^{ax}$$

für alle $x \in \mathbb{R}$.

Beweis: Definieren wir die Funktion $F(x) = f(x) \cdot e^{-ax}$. Dann gilt für die Ableitung von $F(x)$ nach der Produktregel:

$$\begin{aligned}\Rightarrow F'(x) &= f'(x) e^{-ax} + f(x) e^{-ax} (-a) \\ &= (f'(x) - a f(x)) e^{-ax} = 0 \quad \text{für alle } x \in \mathbb{R}.\end{aligned}$$

Ist die Ableitung einer Funktion gleich Null, so ist die Funktion konstant. Folglich ist $F(x)$ die konstante Funktion mit $F(0) = f(0) = \text{const}$ und damit $F(x) = f(0)$ für alle $x \in \mathbb{R}$.

$$\begin{aligned}\Rightarrow f(x) &= F(x) \cdot e^{ax} = f(0) \cdot e^{ax} \quad \text{für alle } x \in \mathbb{R}. \\ \Rightarrow f(x) &= f(0) \cdot e^{ax} \quad \text{für alle } x \in \mathbb{R}. \quad \square\end{aligned}$$

Folgerungen:

(1) $\exp : \mathbb{R} \rightarrow \mathbb{R}$ mit $x \mapsto e^x$ ist die einzige differenzierbare Funktion mit $f'(x) = f(x)$ und $f(0) = 1$.

(2) Die Lösung der Differenzialgleichung

$$y'(x) = a y(x) \quad \text{mit } y(0) = y_0$$

ist gegeben durch $y(x) = y_0 e^{ax}$.

(3) Die Lösung der Differenzialgleichung

$$y'(x) = a y(x) + f(x) \quad \text{mit } y(0) = y_0$$

ist eindeutig.

► 7.7.2 Mittelwertsatz

Ein anschaulich plausibler Satz besagt, dass jede Funktion $f : [a, b] \rightarrow \mathbb{R}$ mit $f(a) = f(b)$ im Innern des Intervalls eine Stelle mit waagrechtener Tangente besitzt. Dies ist die Aussage des *Satzes von Rolle* (1652 - 1719).

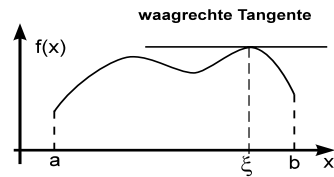


Abb. 7.18. Satz von Rolle

Satz von Rolle: Ist $f : [a, b] \rightarrow \mathbb{R}$ eine differenzierbare Funktion mit $f(a) = f(b)$. Dann existiert eine Zwischenstelle $\xi \in (a, b)$ mit

$$f'(\xi) = 0.$$

Eine Erweiterung dieser Aussage bringt der folgende Mittelwertsatz, den wir auf den Satz von Rolle zurückführen:

Mittelwertsatz: Ist $f : [a, b] \rightarrow \mathbb{R}$ eine differenzierbare Funktion. Dann existiert eine Zwischenstelle $\xi \in (a, b)$ mit

$$f'(\xi) = \frac{f(b) - f(a)}{b - a}.$$

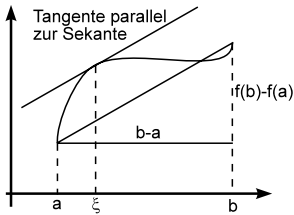


Abb. 7.19. Mittelwertsatz

Begründung: Wir definieren die Funktion $F : [a, b] \rightarrow \mathbb{R}$ mit

$$F(x) = f(x) - \frac{f(b) - f(a)}{b - a} (x - a).$$

F ist stetig und differenzierbar mit $F(a) = f(a) = F(b)$. Nach dem Satz von Rolle existiert ein $\xi \in (a, b)$ mit

$$F'(\xi) = f'(\xi) - \frac{f(b) - f(a)}{b - a} = 0.$$

Daraus folgt die Behauptung. $\square$

Geometrisch besagt der Mittelwertsatz, dass der Graph der Funktion an mindestens einer Stelle eine Tangente besitzt, die parallel zur Sekante durch die Punkte $(a, f(a))$ und $(b, f(b))$ verläuft.

7.7.3 Die Regeln von l'Hospital

Wir wenden uns nun wieder dem Berechnen von Funktionsgrenzwerten zu. Eine der Regeln lautet $\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{f(x_0)}{g(x_0)}$, wenn $g(x_0) \neq 0$. Die Regel von l'Hospital (1661 - 1704) liefert eine Methode auch den Grenzwert zu berechnen, wenn $g(x_0) = f(x_0) = 0$.

Satz: Regel von l'Hospital. Sind f und g in x_0 stetig differenzierbar und ist $f(x_0) = g(x_0) = 0$, so gilt

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)}.$$

Beweis: Nach dem Mittelwertsatz ist für $a = x_0$ und $b = x_0 + h$

$$f(x_0 + h) = f(x_0) + h f'(\xi), \quad \text{wenn } \xi \in (x_0, x_0 + h).$$

Mit $\delta \in (0, 1)$ ist $\xi = x_0 + \delta h$ und

$$f(x_0 + h) = f(x_0) + h f'(x_0 + \delta \cdot h).$$

Wenden wir diese Form des Mittelwertsatzes auf f und g an, folgt mit $f(x_0) = g(x_0) = 0$

$$\frac{f(x_0 + h)}{g(x_0 + h)} = \frac{f(x_0) + h f'(x_0 + \delta_1 h)}{g(x_0) + h g'(x_0 + \delta_2 h)} = \frac{f'(x_0 + \delta_1 h)}{g'(x_0 + \delta_2 h)}.$$

Durch Grenzübergang $h \rightarrow 0$ folgt die Behauptung. $\square$

Bemerkungen:

- (1) Die Regel von l'Hospital gilt auch für Grenzübergänge der Form " $\frac{\infty}{\infty}$ ":
 Wenn $\lim_{x \rightarrow x_0} f(x) = \infty$ und $\lim_{x \rightarrow x_0} g(x) = \infty$, dann ist

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)}.$$

- (2) Die Regeln von l'Hospital gelten auch für Grenzübergänge $x \rightarrow \infty$:

$$\lim_{x \rightarrow \infty} \frac{f(x)}{g(x)} = \lim_{x \rightarrow \infty} \frac{f'(x)}{g'(x)},$$

wenn $\left(\lim_{x \rightarrow \infty} f(x) = \lim_{x \rightarrow \infty} g(x) = 0 \right)$
 oder $\left(\lim_{x \rightarrow \infty} f(x) = \infty \text{ und } \lim_{x \rightarrow \infty} g(x) = \infty \right)$.

- (3) Die l'Hospital'schen Regeln setzen immer voraus, dass die Funktionen in einer Umgebung von x_0 differenzierbar sind.
- (4) Unter Umständen müssen die l'Hospital'schen Regeln mehrfach angewendet werden, um zum Ziel zu führen. Es gibt aber auch Fälle, bei denen die mehrmalige Anwendung der Regeln versagt.

Beispiele 7.30:

$$\textcircled{1} \lim_{x \rightarrow 0} \frac{\sin x}{x} \stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow 0} \frac{\cos x}{1} = 1.$$

$$\textcircled{2} \lim_{x \rightarrow 0} \frac{e^x - 1}{2x} \stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow 0} \frac{e^x}{2} = \frac{1}{2}.$$

$$\textcircled{3} \lim_{x \rightarrow \infty} \frac{3x^2 - 5}{x + 4x^2} \stackrel{\frac{\infty}{\infty}}{=} \lim_{x \rightarrow \infty} \frac{6x}{8x + 1} \stackrel{\frac{\infty}{\infty}}{=} \lim_{x \rightarrow \infty} \frac{6}{8} = \frac{3}{4}.$$

$$\textcircled{4} \lim_{x \rightarrow \infty} \frac{\ln(2x - 1)}{e^x} \stackrel{\frac{\infty}{\infty}}{=} \lim_{x \rightarrow \infty} \frac{\frac{2}{2x-1}}{e^x} = \lim_{x \rightarrow \infty} \frac{2}{(2x - 1) e^x} = 0. \quad \square$$

Ausdrücke der Form $0 \cdot \infty$, $\infty - \infty$, 0^0 , ∞^0 , 1^∞ : Die Regeln von l'Hospital gelten nur für unbestimmte Ausdrücke der Form " $\frac{0}{0}$ " oder " $\frac{\infty}{\infty}$ ". Andere unbestimmte Ausdrücke wie z.B. $0 \cdot \infty$, $\infty - \infty$, 0^0 , ∞^0 , 1^∞ lassen sich durch die folgenden elementaren Umformungen auf obige Fälle zurückführen:

	Funktion $\varphi(x)$	$\lim_{x \rightarrow x_0} \varphi(x)$	Umformung
(A)	$u(x) \cdot v(x)$	$0 \cdot \infty$	$\frac{u(x)}{1/v(x)}$ oder $\frac{v(x)}{1/u(x)}$
(B)	$u(x) - v(x)$	$\infty - \infty$	$\frac{1/v(x) - 1/u(x)}{1/(u(x) \cdot v(x))}$
(C)	$u(x)^{v(x)}$	$0^0, \infty^0, 1^\infty$	$\exp(v(x) \ln(u(x)))$

Beispiele 7.31:

$$\textcircled{1} \lim_{x \rightarrow \infty} x \cdot \ln\left(1 + \frac{1}{x}\right) \stackrel{(A)}{=} \lim_{x \rightarrow \infty} \frac{\ln\left(1 + \frac{1}{x}\right)}{\frac{1}{x}} \stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow \infty} \frac{\frac{1}{1+1/x} \cdot \left(-\frac{1}{x^2}\right)}{\left(-\frac{1}{x^2}\right)}$$

$$= \lim_{x \rightarrow \infty} \frac{1}{1 + \frac{1}{x}} = 1.$$

$$\textcircled{2} \lim_{x \rightarrow 0} \left(\frac{1}{x} - \frac{1}{\sin x}\right) \stackrel{(B)}{=} \lim_{x \rightarrow 0} \frac{\sin x - x}{x \cdot \sin x} \stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow 0} \frac{\cos x - 1}{\sin x + x \cdot \cos x}$$

$$\stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow 0} \frac{-\sin x}{2 \cos x - x \cdot \sin x} = \frac{0}{2} = 0.$$

$$\textcircled{3} \lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x \stackrel{(C)}{=} \lim_{x \rightarrow \infty} e^{x \ln\left(1 + \frac{1}{x}\right)}.$$

Nach ① ist $\lim_{x \rightarrow \infty} x \ln\left(1 + \frac{1}{x}\right) = 1$ und damit

$$\lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x = \exp\left(\lim_{x \rightarrow \infty} x \ln\left(1 + \frac{1}{x}\right)\right) = e^1 = e.$$

$$\textcircled{4} \lim_{x \rightarrow 0} (1 + tx)^{\frac{1}{x}} \stackrel{(C)}{=} \lim_{x \rightarrow 0} \exp\left(\frac{1}{x} \ln(1 + tx)\right).$$

Wegen $\lim_{x \rightarrow 0} \frac{\ln(1 + tx)}{x} \stackrel{\frac{0}{0}}{=} \lim_{x \rightarrow 0} \frac{\frac{1}{1+tx} \cdot t}{1} = t$ folgt

$$\lim_{x \rightarrow 0} (1 + tx)^{\frac{1}{x}} = \exp\left(\lim_{x \rightarrow 0} \frac{1}{x} \ln(1 + tx)\right) = \exp(t) = e^t. \quad \square$$

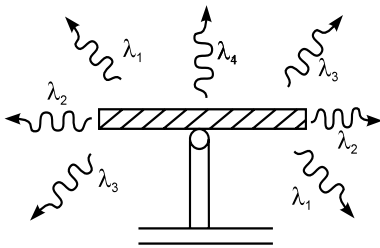
7.8 Spektrum eines strahlenden schwarzen Körpers

Bei der Diskussion der Strahlungsformel eines schwarzen Körpers werden die Rechentechniken aus der Differenzialrechnung angewendet, um den Verlauf des Spektrums zu erhalten. Wir benötigen dazu sowohl die Regeln von l'Hospital, die Extremwertbestimmung über die erste Ableitung als auch das logarithmische Differenzieren zur Berechnung von f' .

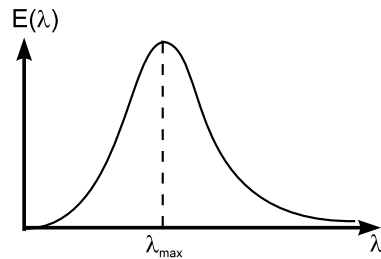
Ein Körper mit nicht spiegelnder Oberfläche (*schwarzer Körper*) sendet bei der absoluten Temperatur T Strahlen aus. Für die spektrale Strahlungsdichte $E(\lambda)$ gilt nach dem **Planckschen Strahlungsgesetz** im Raumwinkel Ω_0

$$E(\lambda) = \frac{c_1}{\lambda^5 (\exp(\frac{c_2}{T\lambda}) - 1)} \cdot \frac{1}{\Omega_0}$$

mit $c_1 = 2hc^2 = 1.191 \cdot 10^{-16} \text{ Wm}$, $c_2 = h \frac{c}{k} = 1.439 \cdot 10^{-2} \text{ mK}$ (c ist die Lichtgeschwindigkeit, k die Boltzmann-Konstante, h das Plancksche Wirkungsquantum).



Schwarzer Körper



Spektrale Strahlungsdichte

Verhalten für $\lambda \rightarrow 0$ und $\lambda \rightarrow \infty$: Für $\lambda \rightarrow 0$ und $\lambda \rightarrow \infty$ gilt $E(\lambda) \rightarrow 0$, wie man mit den Regeln von l'Hospital nachrechnet, denn für den Nenner der Strahlungsformel berechnet man

(i) durch 5-maliges Anwenden der Regeln

$$\begin{aligned} \lim_{\lambda \rightarrow 0} \lambda^5 (\exp(\frac{c_2}{T\lambda}) - 1) &\stackrel{0:\infty}{=} \lim_{\lambda \rightarrow 0} \frac{\exp(\frac{c_2}{T\lambda}) - 1}{\frac{1}{\lambda^5}} \stackrel{\infty}{=} \frac{1}{5} \frac{c_2}{T} \lim_{\lambda \rightarrow 0} \frac{\exp(\frac{c_2}{T\lambda})}{\frac{1}{\lambda^4}} \\ &\stackrel{\infty}{=} \frac{1}{5 \cdot 4} \left(\frac{c_2}{T}\right)^2 \lim_{\lambda \rightarrow 0} \frac{\exp(\frac{c_2}{T\lambda})}{\frac{1}{\lambda^3}} = \dots = \frac{1}{5!} \left(\frac{c_2}{T}\right)^5 \lim_{\lambda \rightarrow 0} \exp(\frac{c_2}{T\lambda}) = \infty \end{aligned}$$

(ii) durch einmaliges Anwenden der Regeln

$$\begin{aligned}\lim_{\lambda \rightarrow \infty} \lambda^5 \left(\exp\left(\frac{c_2}{T\lambda}\right) - 1 \right) &\stackrel{\infty \cdot 0}{=} \lim_{\lambda \rightarrow \infty} \frac{1}{5} \frac{c_2}{T} \frac{\exp\left(\frac{c_2}{T\lambda}\right)}{\frac{1}{\lambda^4}} = \\ &= \lim_{\lambda \rightarrow \infty} \frac{1}{5} \frac{c_2}{T} \lambda^4 \exp\left(\frac{c_2}{T\lambda}\right) = \infty \cdot 1 = \infty.\end{aligned}$$

Maximum. Wir bestimmen die Wellenlänge λ , bei welcher die Strahlungsdichte $E(\lambda)$ bei festem T ein Maximum besitzt. Durch logarithmische Differenziation des Strahlungsgesetzes folgt

$$\begin{aligned}\ln E(\lambda) &= \ln c_1 - 5 \ln \lambda - \ln \left(\exp\left(\frac{c_2}{T\lambda}\right) - 1 \right) - \ln \Omega_0 \\ \hookrightarrow \frac{1}{E(\lambda)} \cdot E'(\lambda) &= -\frac{5}{\lambda} - \frac{1}{\exp\left(\frac{c_2}{T\lambda}\right) - 1} \cdot \exp\left(\frac{c_2}{T\lambda}\right) \cdot \left(-\frac{c_2}{\lambda^2 T} \right).\end{aligned}$$

Setzt man $z = \frac{c_2}{\lambda T}$ gilt für das Extremum von $E(\lambda)$

$$E'(\lambda) = 0 \Rightarrow \frac{e^z}{e^z - 1} \cdot z = 5$$

(denn $E(\lambda) \neq 0$ für alle $\lambda > 0$). Folglich gilt für z die nichtlineare Gleichung

$$1 - \frac{1}{5}z = e^{-z}$$

($\rightarrow$ §7.9 Beispiel 7.32). Durch numerisches Lösen dieser Gleichung erhält man $z \approx 4.965$, also ist

$$\lambda_{\max} \cdot T = c_2 / 4.965 = 2898 \mu m K.$$

Dieses Ergebnis heißt das **Wiensche Verschiebungsgesetz**. Es besagt, dass $\lambda_{\max} \cdot T = \text{const.}$ (dass dies auch das Maximum darstellt, zeigt man, indem $E''(\lambda_{\max}) < 0$ nachgeprüft wird.)

Diskussion: Für steigende Temperaturen verschiebt sich das Maximum der Strahlung zu kleineren Wellenlängen hin. Die Strahlung eines Körpers wird sichtbar, wenn die Temperatur etwa $600^\circ C$ erreicht (Rotglut). Mit steigender Temperatur verschiebt sich die Glühfarbe von $850^\circ C$ hellrot, $1000^\circ C$ gelb, hin zu weiß bei $1300^\circ C$.

Bemerkung: Das Wiensche Verschiebungsgesetz wird in Temperatursensoren herangezogen, um die Temperatur eines Körpers kontaktfrei zu messen: Aus der Analyse des Strahlungsmaximums erhält man unter der Annahme eines idealen schwarzen Strahlers die Körpertemperatur

$$T = \frac{2898 \mu m K}{\lambda_{\max} [\mu m]}.$$

7.9 Newton-Verfahren

Das *Newton-Verfahren* ist ein schnelles, numerisches Verfahren, um Nullstellen von Funktionen näherungsweise zu bestimmen. Die Beschreibung weiterer Verfahren befinden sich auf der CD-Rom im Kapitel über [das Lösen von nichtlinearen Gleichungen](#) der Form $f(x) = 0$.

Wir suchen eine Methode, um von einer Funktion f eine einfache Nullstelle

$$f(x) = 0$$

im Intervall $[a, b]$ näherungsweise zu bestimmen.

Das Verfahren. Das Newton-Verfahren ist ein iteratives Verfahren, d.h. man startet mit einer Anfangsschätzung x_0 für die Nullstelle und verbessert in weiteren Iterationsschritten diesen Wert: Berechnen wir dazu im Punkt x_0 die Tangente an f und bestimmen den Schnittpunkt der Tangente mit der x -Achse. Dieser Wert sei x_1 . x_1 liegt in der Regel näher an der Nullstelle als x_0 . Nun berechnet man in x_1 die Tangente der Funktion und bestimmt den Achsenschnittpunkt x_2 . Durch Fortführung des Verfahrens nähert man sich der Nullstelle an (siehe Abb. 7.20).

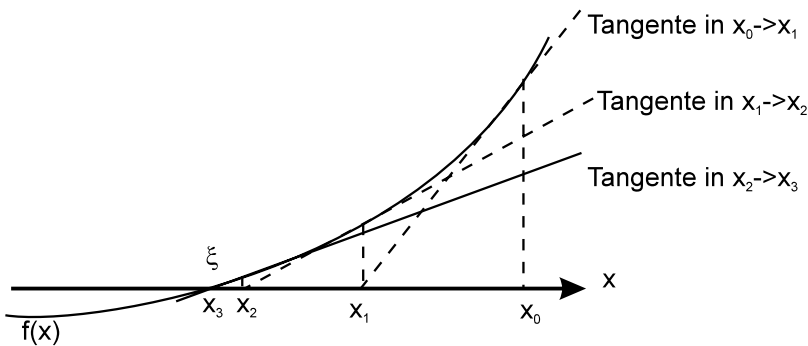


Abb. 7.20. Geometrische Interpretation des Newton-Verfahrens

Aufstellen der Formeln: Die Tangentengleichung in x_0 hat die Form

$$y = f(x_0) + f'(x_0)(x - x_0)$$

und der Schnittpunkt x_1 mit der x -Achse ist definiert durch $y = 0$:

$$0 = f(x_0) + f'(x_0)(x_1 - x_0).$$

Damit ist

$$x_1 = x_0 - \frac{f(x_0)}{f'(x_0)}.$$

Durch Iteration erhält man das folgende Verfahren:

Algorithmus (Newton-Verfahren)

- (1) **Initialisierung:** Wähle Startwert x_0 ; $\delta := 10^{-5}$.
- (2) **Iteration:** $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} \quad n = 0, 1, 2, 3, \dots$
- (3) **Abbruchbedingung:**
 - Falls $|x_{n+1} - x_n| < \delta$, dann $\xi = x_{n+1}$. Stop.
 - Falls $|x_{n+1} - x_n| \geq \delta$, dann weiter mit (2).

Eigenschaften des Verfahrens

- (1) Sehr schnelle Konvergenz (nur wenige Iterationsschritte sind nötig).
- (2) ⚠ Das Verfahren kann divergieren, falls der Startwert nicht nahe genug an der Lösung liegt.
- (3) Für das Newton-Verfahren gilt die Fehlerabschätzung

$$|x_{n+1} - x^*| \leq L^n |x_{n+1} - x_1| \leq L^n (b - a),$$

wenn $|f'(x)| \leq L$ für alle $x \in I = [a, b]$.

- (4) Eine analytische Formel der Ableitung muss vorliegen.
- (5) Das Newton-Verfahren konvergiert immer für eine auf $\mathbb{R}$ konvexe oder konkave Funktion, die eine Nullstelle mit $f'(x_0) \neq 0$ besitzt.
- (6) Die Konvergenz des Newton-Verfahrens kann man auch sicherstellen, wenn f eine 3-mal stetig differenzierbare Funktion ist mit der Eigenschaft

$$f'(x) \neq 0 \text{ für alle } x \in I \text{ und } \left| \frac{f(x) f''(x)}{f'(x)^2} \right| \leq K < 1 \text{ für alle } x \in I.$$

Je kleiner die Konstante K , desto besser ist die Konvergenz.

Falls man eine schlechte Schätzung für den Startwert hat, sollte man sich zunächst mit dem Bisektionsverfahren ($\rightarrow$ 6.4) eine bessere Startnäherung verschaffen und anschließend das Newton-Verfahren verwenden. Oftmals kann man durch Zeichnen des Funktionsgraphen einen guten Startwert x_0 finden.

Anwendungsbeispiel 7.32 (Mit MAPLE-Worksheet). Wir wenden das Newton-Verfahren auf das Problem aus 7.8 an: Gesucht ist eine Lösung der Gleichung $1 - \frac{1}{5}z = e^{-z}$, d.h. gesucht ist eine positive Nullstelle der Funktion

$$f(z) = 1 - \frac{1}{5}z - e^{-z}.$$

Mit der Ableitung

$$f'(z) = -\frac{1}{5} + e^{-z}$$

stellt man die entsprechende Newtonfolge

$$x_{n+1} = x_n - \frac{1 - \frac{1}{5}x_n - e^{-x_n}}{-\frac{1}{5} + e^{-x_n}} = x_n - \frac{(5 - x_n)e^{x_n} - 5}{-e^{x_n} + 5}$$

auf. Mit dem Startwert $x_0 = 2$ erhält man das Ergebnis:

n	x_n	$f(x_n)$
0	2	0.4646
1	9.1857	-.8372
2	4.9973	-.00622
3	4.9651	$-0.3 \cdot 10^{-5}$

Nach 3 Iterationen erhält man für die Lösung der Gleichung $z = 4.9651$ mit einer Genauigkeit von 4 Dezimalstellen. $\square$



Animation: Der Algorithmus des Newton-Verfahrens wird direkt in die Prozedur **newton** übernommen. Im zugehörigen MAPLE-Worksheet ist eine Animation dargestellt, die den Konvergenzprozess der Newton-Folge an die Nullstelle visualisiert.

⊙ Anwendung des Newton-Verfahrens: Berechnung von Wurzeln.

Gesucht ist die Quadratwurzel $\sqrt{a}$ einer positiven Zahl a . Wir interpretieren $\sqrt{a}$ als die einzige positive Nullstelle der Funktion

$$f(x) = x^2 - a.$$

Zur numerischen Berechnung wenden wir das Newton-Verfahren auf diese Funktion an. Mit

$$f'(x) = 2x$$

erhält man die Newtonfolge

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} = x_n - \frac{x_n^2 - a}{2x_n} = \frac{2x_n^2 - x_n^2 + a}{2x_n} = \frac{1}{2} \left(x_n + \frac{a}{x_n} \right).$$

Setzt man $x_0 := a$ und iteriert gemäß

$$x_{n+1} := \frac{1}{2} \left(x_n + \frac{a}{x_n} \right)$$

für $n = 0, 1, 2, \dots$, hat man eine schnell konvergierende Folge (vgl. Beispiel 6.4: Babylonisches Wurzelziehen).

Beispiel 7.33. Berechnung von $\sqrt{3}$ mit obigem Schema

n	0	1	2	3	4	5
x_n	3.00000	2.00000	1.75000	1.73214	1.73205	1.73205

Die angegebene Methode ist eine der besten zur Berechnung von Quadratwurzeln. Die meisten Computerprogramme beruhen darauf. $\square$

⊗ Berechnung von k -ten Wurzeln

Dieses Verfahren ist nicht nur auf die Berechnung von Quadratwurzeln beschränkt, sondern kann auch zur Berechnung von k -ten Wurzeln herangezogen werden. Denn $\sqrt[k]{a}$ ist die positive Nullstelle von

$$f(x) = x^k - a.$$

Mit $f'(x) = k x^{k-1}$ lautet die Newtonfolge

$$x_{n+1} = \frac{1}{k} \left[(k-1) x_n - \frac{a}{x_n^{k-1}} \right] \quad n = 0, 1, 2, 3, \dots$$

MAPLE-Worksheets zu Kapitel 7



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 7 mit MAPLE zur Verfügung.

- Zum Ableitungsbegriff
- Differenzieren mit MAPLE
- Kurvendiskussion mit MAPLE
- Helmholtz-Spulen mit MAPLE
- Newton-Verfahren mit MAPLE
- MAPLE-Lösungen zu den Aufgaben

7.10 Aufgaben zur Differenzialrechnung

7.1 Bestimmen Sie die erste Ableitung der folgenden Funktionen unter Verwendung der Potenzregel:

$$\begin{array}{ll} \text{a) } y = 8x^7 - 10x^3 + \frac{10}{x^3} - \frac{8}{x^7} & \text{b) } y = 12\sqrt[4]{x^3} - 7\sqrt[7]{x^4} + 11x - \frac{8}{\sqrt{x^3}} \\ \text{c) } y(l) = 2\sqrt[4]{15\sqrt{l}} + 3\sqrt[5]{12\sqrt{l}} - 3\sqrt[3]{20\sqrt{l}} - 3\sqrt[6]{10\sqrt{l}} & \text{d) } y(a) = \frac{9}{a^5\sqrt{a^3}} \\ \text{e) } y(x) = (a + bx^2)(c + ex)^3 & \text{f) } y(x) = (x^3 + x^2)\sqrt{x} \quad \text{g) } y(x) = x^\alpha x^\beta \end{array}$$

7.2 Bestimmen Sie mit der Produkt- und Quotientenregel die Ableitung der Funktionen:

$$\begin{array}{lll} \text{a) } \sin(x) \cdot \frac{10}{x^3} & \text{b) } \sin(x) \cdot \cos(x) & \text{c) } x^n e^x \\ \text{d) } \frac{x^2 - 5x + 6}{x^2 - 12x + 20} & \text{e) } \frac{\sin(\varphi)}{(1 - \cos(\varphi))} & \text{f) } \frac{4t}{(t^2 - 1)(t + 1)} \\ \text{g) } \frac{x}{x^2 + 2} & \text{h) } \frac{x^2 e^x}{e^x - 1} & \text{i) } \frac{x \cdot \ln(x)}{(x - 1)^2} \end{array}$$

7.3 Man bestimme die erste Ableitung der folgenden Funktionen unter Zuhilfenahme der Kettenregel

$$\begin{array}{lll} \text{a) } y(x) = \cos(3x + 2) & \text{b) } y(x) = (3x - 2)^3 & \text{c) } y(x) = 3 \cdot \sin(5x) \\ \text{d) } y(x) = e^{4x^2 - 3x + 2} & \text{e) } y(x) = 10 \cdot \ln(1 + x^2) & \\ \text{f) } x(t) = A \cdot \sin(\omega t + \varphi) & \text{g) } y(x) = \ln(\sin(2x - 3)) & \text{h) } y(x) = \sqrt{\ln(x^2 - 1)} \end{array}$$

7.4 Man berechne durch logarithmische Differenzierung die Ableitung von

$$\text{a) } y(x) = x^x \quad \text{b) } y(x) = x^{\sin x}$$

7.5 Wie lautet die Ableitung von

$$\begin{array}{lll} \text{a) } f_1(x) = x^{(x^x)} & \text{b) } f_2(x) = (x^x)^x & \text{c) } f_3(x) = x^{(x^a)} \\ \text{d) } f_4(x) = x^{(a^x)} & \text{e) } f_5(x) = a^{(x^x)} & \end{array}$$

7.6 Bestimmen Sie die erste Ableitung von

$$\begin{array}{lll} \text{a) } y(t) = \ln \sqrt{a^2 - t^2} & \text{b) } y = \ln \sqrt{\frac{1-x^2}{1+x^2}} & \text{c) } y = \ln \frac{(x-5)^3}{(x+1)^2} \\ \text{d) } y(x) = a^{\ln(x-3)} & \text{e) } y = e^x \cdot \sqrt{\frac{1+x}{1-x}} & \text{f) } y = e^{\ln x} \end{array}$$

7.7 Gegeben seien die Funktionen

$$\begin{array}{l} \sinh : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \sinh(x) := \frac{1}{2}(e^x - e^{-x}) \quad (\text{Sinushyperbolikus}) \\ \cosh : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \cosh(x) := \frac{1}{2}(e^x + e^{-x}) \quad (\text{Kosinushyperbolikus}) \\ \tanh : \mathbb{R} \rightarrow \mathbb{R} \text{ mit } \tanh(x) := \frac{\sinh(x)}{\cosh(x)} \quad (\text{Tangenshyperbolikus}) \end{array}$$

i) Man zeichne den Graphen der 3 Hyperbolikusfunktionen.

ii) Man berechne die Ableitung der Funktionen.

iii) Man zeige, dass $\cosh^2(x) - \sinh^2(x) = 1$.

7.8 Man berechne die Ableitungen der Arkusfunktionen

$$\arcsin(x), \arccos(x), \arctan(x), \operatorname{arccot}(x)$$

als Ableitung der Umkehrfunktion der trigonometrischen Funktionen.

7.9 Berechnen Sie die Ableitung der Areafunktionen

$$\operatorname{ar} \sinh(x) \quad \text{und} \quad \operatorname{ar} \cosh(x)$$

als Ableitung der Umkehrfunktion von $\sinh$ und $\cosh$.

7.10 Beweisen Sie die Potenzregel $y(x) = x^n \Rightarrow y'(x) = n x^{n-1}$ mit Hilfe der logarithmischen Differenziation.

7.11 Bilden Sie die erste Ableitung der implizit gegebenen Funktionen

$$\begin{array}{ll} \text{a) } e^{x \cdot y(x)} + y^3(x) \ln x = \cos(2x) & \text{b) } y(e^{-x y(x)}) = \left(\sqrt[x]{y(x) \sqrt{a^{2x+2} a^{2x}}} \right)^{-\frac{x}{y(x)}} \\ \text{c) } \ln y(x) - \sqrt{y(x)} - x = 0 & \text{d) } \sin y(x) = y(x) \cdot x^2 \end{array}$$

7.12 Bestimmen Sie durch implizite Differenziation den Anstieg der Kreistangente im Punkte $P_0 = (4, y_0 > 0)$ des Kreises $(x-2)^2 + (y-1)^2 = 25$.

7.13 Gegeben seien die Funktionen

$$\begin{array}{ll} \text{a) } f_1(x) = \sqrt{1+x^4}; x_0 = 1 & \text{b) } f_2(x) = 3 \ln(1+3x^5); x_0 = 3 \\ \text{c) } y(x) = 2 \cos x; x_0 = \frac{\pi}{4}. \end{array}$$

Man berechne für die Funktionen i) das totale Differenzial ii) das totale Differenzial am Punkte x_0 . Man bestimme außerdem iii) die Tangente im Punkte x_0 und iv) die Linearisierung am Punkte x_0 . v) Man gebe einen Näherungswert für $f(x_0 + 0.01)$ an und vergleiche diesen mit dem exakten Wert.

7.14 Ein gedämpftes Feder-Masse-System hat ein Weg-Zeit-Gesetz der Form

$$x(t) = A e^{-\gamma t} \cos(\omega t).$$

i) Man berechne die Geschwindigkeit und Beschleunigung zu jedem Zeitpunkt.
ii) Man gebe eine Bedingung für die Nebenmaxima an.

7.15 Die potentielle Energie für ein Ion in einem Kristallgitter lautet näherungsweise

$$V(r) = -D \left(\frac{2a}{r} - \frac{a^2}{r^2} \right) \quad (D > 0).$$

Man zeige, dass $V(r)$ an der Stelle $r_0 = a$ ein relatives Minimum besitzt.

7.16 Bei der Spiegelabmessung mit Skala und Fernrohr wird bei festem Skalenabstand s der Ausschlag x gemessen. Wie beeinflusst ein kleiner Messfehler von x den Wert des Ergebnisses α , wenn $\alpha = \arctan \frac{x}{s}$? ($s = 2m$, $x = 250mm$, $dx = 1mm$.) Welches ist der relative Fehler?

7.17 Wo besitzen die folgenden Funktionen relative Extremwerte?

$$\begin{array}{ll} \text{a) } y(x) = -8x^3 + 12x^2 + 18x & \text{b) } z(t) = t^4 - 8t^2 + 16 \\ \text{c) } u(z) = \sqrt{1+z} + \sqrt{1-z} & \text{d) } y(x) = x e^{-x} \\ \text{e) } y(x) = \sin x \cdot \cos x & \text{f) } y(x) = \frac{2x-2x^2}{x^2-x-6} \end{array}$$

7.18 Man diskutierte den Verlauf der folgenden Funktionen:

$$\text{a) } y = \frac{x^2 + 1}{x - 3} \quad \text{b) } y = \frac{(x-1)^2}{x+1} \quad \text{c) } y = \frac{\ln x}{x} \quad \text{d) } y = \sin^2 x$$

7.19 Bestimmen Sie die folgenden Funktionswerte mit den Regeln von l'Hospital

$$\begin{array}{lll} \text{a) } \lim_{x \rightarrow a} \frac{x^2 - a^2}{x - a} & \text{b) } \lim_{x \rightarrow 0} \frac{\sin(2x)}{\sin(x)} & \text{c) } \lim_{x \rightarrow 0} \frac{\sin^2 x}{1 - \cos x} \\ \text{d) } \lim_{x \rightarrow 0} \frac{x^2 - 2 + 2 \cos x}{x^4} & \text{e) } \lim_{x \rightarrow 0} \left(\frac{1}{x} - \frac{1}{\sin x} \right) & \text{f) } \lim_{x \rightarrow 1} \frac{\ln x - x + 1}{(x-1)^2} \\ \text{g) } \lim_{x \rightarrow 0} x^x & \text{h) } \lim_{x \rightarrow \infty} \left(1 + \frac{a}{x} \right)^x & \end{array}$$

Kapitel 8
Integralrechnung

8

8	Integralrechnung	295
8.1	Das Riemann-Integral	297
8.2	Fundamentalsatz der Differenzial- und Integralrechnung...	302
8.3	Grundlegende Regeln der Integralrechnung	311
8.4	Integrationsmethoden	313
8.4.1	Partielle Integration	313
8.4.2	Integration durch Substitution	315
8.4.3	Partialbruchzerlegung	321
8.5	Uneigentliche Integrale	327
8.6	Anwendungen der Integralrechnung	329
8.6.1	Flächenberechnungen	329
8.6.2	Kinematik	330
8.6.3	Elektrodynamik	332
8.6.4	Energieintegrale	333
8.6.5	Lineare und quadratische Mittelwerte	334
8.6.6	Schwerpunkt einer ebenen Fläche	336
8.7	Aufgaben zur Integralrechnung	339

Zusätzliche Abschnitte auf der CD-Rom

8.8	Weitere Anwendungen	cd
8.8.1	Mittelungseigenschaft	cd
8.8.2	Bogenlänge	cd
8.8.3	Krümmung	cd
8.8.4	Volumen und Mantelflächen von Rotationskörpern	cd
8.9	MAPLE: Integralrechnung	cd
8.10	Zusätzliche Aufgaben zur Integralrechnung	cd

8 Integralrechnung

Der Integralbegriff ist wie der Ableitungsbegriff motiviert durch die physikalische Beschreibung von Bewegungsabläufen (Geschwindigkeit, Beschleunigung). Er ist u.a. auch von Bedeutung bei der Berechnung von Flächen, Volumeninhalten von Körpern, Schwerpunktsberechnungen usw.

Hinweis: Auf der CD-Rom befindet sich ein zusätzliches Kapitel über das [numerische Integrieren](#) sowie einige weitere [Anwendungen der Integralrechnung](#) wie z.B. die Mittelungseigenschaft, die Bogenlänge und das Krümmungsverhalten sowie die Berechnung von Rotationskörpern zusammen mit der Visualisierung in MAPLE.

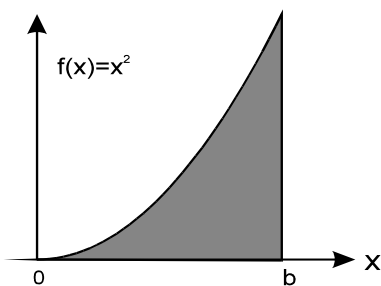
8.1 Das Riemann-Integral

8.1

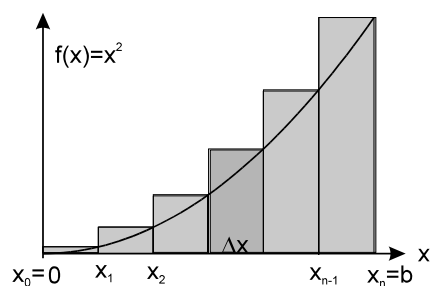
Wir beginnen mit der geometrischen Fragestellung: Gegeben ist eine Funktion $f(x)$, wie groß ist die Fläche, welche die Kurve mit der x -Achse in einem Intervall $[a, b]$ einschließt?

Anwendungsbeispiel 8.1. Zur Bestimmung des Flächeninhalts A_0^b unter dem Graphen der Funktion $f(x) = x^2$ in $[0, b]$ unterteilen wir das Intervall $[0, b]$ durch eine gleichmäßige Zerlegung Z_n in n Teilintervalle

$$x_0 = 0, x_1 = \frac{b}{n}, x_2 = 2\frac{b}{n}, \dots, x_{n-1} = (n-1)\frac{b}{n}, x_n = b.$$



Fläche unter Kurve



Näherung durch Rechtecke

Für jedes Teilintervall mit Intervall-Länge $\Delta x = x_k - x_{k-1} = \frac{b}{n}$ wählen wir den rechten Eckpunkt $x_k = k \cdot \Delta x$ und werten die Funktion darauf aus

$$f(x_k) = x_k^2 = (k \cdot \Delta x)^2.$$

Der Flächeninhalt des zugehörigen Rechtecks ist

$$\Delta x \cdot f(x_k) = \Delta x \cdot k^2 \cdot \Delta x.$$

Anschließend summieren wir alle Rechtecksflächen auf

$$\begin{aligned} S_n &= \Delta x f(x_1) + \Delta x f(x_2) + \dots + \Delta x f(x_n) \\ &= \sum_{k=1}^n \Delta x f(x_k) = (\Delta x)^3 \sum_{k=1}^n k^2. \end{aligned}$$

Nach Abschnitt 1.2.2 gilt die Formel

$$\sum_{k=1}^n k^2 = \frac{1}{6} n(n+1)(2n+1)$$

$$\Rightarrow S_n = (\Delta x)^3 \frac{1}{6} n(n+1)(2n+1) = \frac{b^3}{n^3} \frac{1}{6} n(n+1)(2n+1) \xrightarrow{n \rightarrow \infty} \frac{b^3}{3}.$$

Mit Hilfe einer Verfeinerung der Zerlegung Z_n des Intervalls $[0, b]$ wird die Fläche unter $f(x) = x^2$ beliebig genau angenähert. Für den Grenzübergang $n \rightarrow \infty$ geht die sog. *Zwischensumme* S_n in die Fläche unterhalb des Graphen von x^2 über: $A_0^b = \frac{b^3}{3}$. $\square$

Die Vorgehensweise aus dem Einführungsbeispiel wird verallgemeinert, indem man zur Berechnung der Fläche unterhalb einer beliebigen Kurve $f(x)$ mit der x -Achse die folgende Konstruktion anwendet: Zunächst wird die Kurve durch eine stückweise konstante Funktion (Treppenfunktion) angenähert. Man summiert alle Rechteckflächen auf und erhält so eine Näherung für die Fläche unterhalb der Kurve. Man erhöht die Anzahl der Unterteilungen und erhält eine immer feinere Annäherung der Funktion durch die Treppenfunktionen bzw. Annäherung der Fläche unter der Kurve durch die Summe der Rechteckflächen. Genauer erhält man die folgende Definition für das bestimmte Integral:

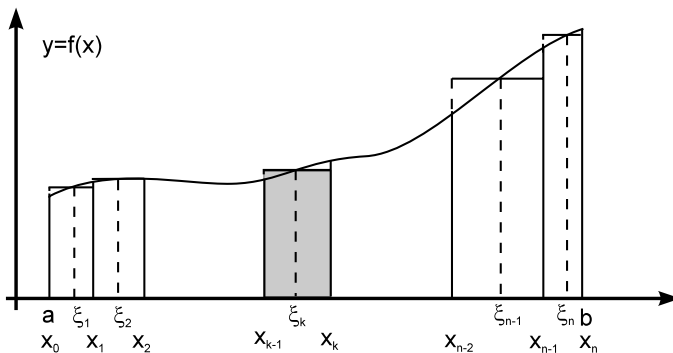


Abb. 8.1. Riemannsche Zwischensumme

Definition: (Bestimmtes Integral; Riemann-Integral)

Gegeben ist eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ mit $y = f(x)$.

- (1) Z_n sei eine Unterteilung des Intervalls $a \leq x \leq b$ in n Teilintervalle

$$a = x_0 < x_1 < x_2 < \dots < x_{n-1} < x_n = b$$

der Längen $\Delta x_k = x_k - x_{k-1}$. Es sei $\xi_k \in [x_{k-1}, x_k]$ ein beliebiger Zwischenwert aus dem Intervall. Dann heißt

$$S_n = \sum_{k=1}^n \Delta x_k f(\xi_k)$$

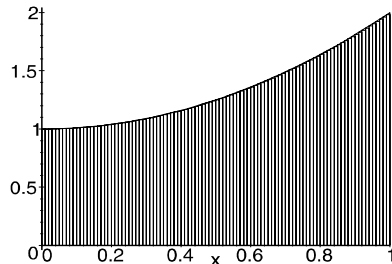
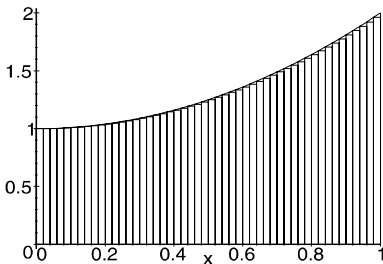
die Riemannsche Zwischensumme bezüglich der Zerlegung Z_n .

- (2) Unter dem bestimmten Integral (Riemann-Integral) der stetigen Funktion f in den Grenzen von $x = a$ bis $x = b$ wird der Grenzwert der Riemannschen Zwischensumme S_n für $n \rightarrow \infty$ verstanden:

$$\int_a^b f(x) dx := \lim_{n \rightarrow \infty} \sum_{k=1}^n \Delta x_k f(\xi_k).$$

**Visualisierung mit MAPLE:** Illustrativer als jede präzise mathe-

matische Definition ist die anschauliche Interpretation. MAPLE liefert im die Möglichkeit, den Übergang von der Zwischensumme zum Integral durchzuführen, indem die Anzahl der Unterteilungen des Intervalls $[a, b]$ immer größer gewählt wird. Die im zugehörigen [MAPLE-Worksheet](#) vorhandene Animation suggeriert den Übergang von der diskreten Zwischensumme zum bestimmten Integral. Dargestellt sind die Werte am Beispiel des Integrals $\int_0^1 (x^2 + 1) dx$ für die Unterteilungen $N = 50$ (links) und $N = 100$ (rechts).

**Bemerkungen:**

- (1) Der Integralbegriff in der obigen Definition wird zur Unterscheidung von anderen Integralbegriffen nach dem Mathematiker *Riemann* (1826 - 1866)

benannt. Da wir uns ausschließlich mit diesem Integral beschäftigen, sprechen wir kurz vom Integral.

- (2) Ist f stetig, so konvergiert die Zwischensumme für jede beliebige Unterteilung Z_n und jede beliebige Wahl von $\xi_k \in [x_{k-1}, x_k]$ gegen den gleichen Wert. Man sagt, das Integral ist *wohldefiniert*.
- (3) Allgemeiner bezeichnet man eine Funktion als *integrierbar*, wenn die Zwischensumme für jede beliebige Unterteilung Z_n und jede beliebige Wahl von $\xi_k \in [x_{k-1}, x_k]$ gegen den gleichen Wert konvergiert. So sind z.B. stückweise stetige Funktionen integrierbar.
- (4) Diese algebraische Definition des Integrals entspricht genau dem Vorgehen bei der Flächenberechnung aus dem Eingangsbeispiel. Bei der geometrischen Motivation ist die Funktion f so gewählt, dass sie im Intervall $[a, b]$ nur positive Werte besitzt. Die algebraische Definition ist jedoch allgemeiner und geht somit über die Flächenberechnung hinaus.
- (5) Allgemein übliche Bezeichnungen für die im bestimmten Integral $\int_a^b f(x) dx$ auftretenden Symbole sind:
 x : Integrationsvariable; $f(x)$: Integrand;
 a : untere Grenze; b : obere Grenze.

Beispiel 8.2. Sei $v(t)$ die Geschwindigkeit eines Massenpunktes als Funktion der Zeit t , der sich entlang der x -Achse bewegt. Zur Zeit $t = 0$ befindet er sich an der Stelle $x = 0$. Gesucht ist der zurückgelegte Weg $x(T)$ zum Zeitpunkt $t = T$.

Ist die Geschwindigkeit konstant, $v(t) = v_0$, so ist der zurückgelegte Weg $x(T) = v_0 T$. Bei variabler Geschwindigkeit $v(t)$ zerlegt man das Zeitintervall $[0, T]$ in Teilintervalle

$$0 = t_0 < t_1 < t_2 < \dots < t_n = T,$$

so dass $v(t)$ sich in jedem Teilintervall annähernd konstant verhält:

$$v(t) \approx v(t_{k-1}) \quad \text{für } t \in [t_{k-1}, t_k] \quad k = 1, \dots, n.$$

Dann berechnet sich der zurückgelegte Weg $x(t_k)$ zum Zeitpunkt $t = t_k$ ($k = 1, 2, \dots, n$) näherungsweise durch

$$\begin{aligned} x(t_1) &\approx (t_1 - t_0) \cdot v(t_0) = \Delta t_1 v(t_0) \\ x(t_2) &\approx x(t_1) + (t_2 - t_1) \cdot v(t_1) = x(t_1) + \Delta t_2 v(t_1) \\ x(t_3) &\approx x(t_2) + (t_3 - t_2) \cdot v(t_2) = \Delta t_1 v(t_0) + \Delta t_2 v(t_1) + \Delta t_3 v(t_2) \\ &\vdots \end{aligned}$$

$$x(t_n) = x(T) \approx \sum_{k=1}^n \Delta t_k v(t_{k-1}).$$

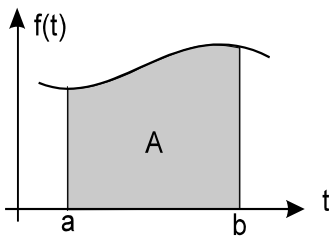
Der erhaltene Näherungswert für $x(T)$ ist somit die Riemannsche Zwischen-summe S_n . Der exakte Wert des zurückgelegten Weges ist

$$x(T) = \int_0^T v(t) dt.$$

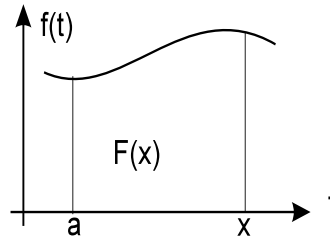
□

➤ Das unbestimmte Integral

Das bestimmte Integral $\int_a^b f(t) dt$ repräsentiert für eine positive Funktion den Flächeninhalt zwischen der Kurve $f(t)$ und der Zeitachse. Betrachtet man die untere Integrationsgrenze als fest, die obere als variabel, so hängt der Integralwert nur noch von der oberen Grenze ab.



Bestimmtes Integral



Integralfunktion

Um die Abhängigkeit von der oberen Grenze zu symbolisieren, ersetzt man b durch x und erhält die Funktion

$$F(x) = \int_a^x f(t) dt.$$

Definition: Unbestimmtes Integral, Integralfunktion.

Unter dem unbestimmten Integral

$$F(x) := \int_a^x f(t) dt$$

versteht man die Integralfunktion $F(x)$, für welche die obere Grenze des Integrals variabel gewählt wird.

Das unbestimmte Integral $F(x) = \int_a^x f(t) dt$ repräsentiert also den Flächeninhalt zwischen der Funktion $f(t)$ und der t -Achse in Abhängigkeit der oberen Grenze.

Beispiel 8.3. Für die Funktion $f(t) = t^2$ ist nach Beispiel 8.1 die zugehörige Integralfunktion für $a = 0$ die Funktion $F : \mathbb{R} \rightarrow \mathbb{R}$ mit $F(x) = \frac{x^3}{3}$. Man beachte, dass hierbei ein Zusammenhang zwischen Integralfunktion und Integrand besteht: $F'(x) = f(x)$. Dieser Zusammenhang gilt ganz allgemein, wie in Abschnitt 8.2 gezeigt wird. $\square$

⑦ Numerische Integration

Schon verhältnismäßig einfache Funktionen lassen sich nicht mehr elementar integrieren. Beispiele sind z.B. e^{-x^2} oder $\frac{\sin x}{x}$. Man ist in diesen Fällen auf numerische Methoden angewiesen. Wir zerlegen dazu das Intervall $[a, b]$ in n Teilintervalle $[x_i, x_{i+1}]$ mit der Intervall-Länge $h := \frac{b-a}{n}$ und setzen

$$x_0 = a; \quad x_{i+1} = x_i + h \quad (i = 0, \dots, n-1); \quad x_n = b.$$

Ersetzt man die zu integrierende Funktion $f(x)$ in jedem Intervall $[x_i, x_{i+1}]$ durch eine konstante $f(\xi_i)$, $\xi_i \in [x_i, x_{i+1}]$, so wird das Integral durch die Zwischensumme

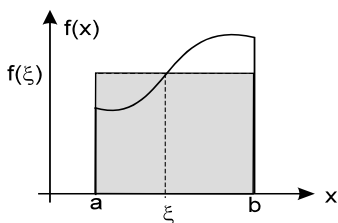
$$I \approx \sum_{i=0}^{n-1} A_i = \sum_{i=0}^{n-1} f(\xi_i) (x_{i+1} - x_i) = h \sum_{i=0}^{n-1} f(\xi_i)$$

approximiert. Somit erhält man als einfachste Näherung

$$\int_a^b f(x) dx \approx h (f(\xi_0) + f(\xi_1) + \dots + f(\xi_{n-1})).$$

Hinweis: Auf der CD-Rom befindet sich ein eigenes Kapitel über das [numerische Integrieren](#) mit der Diskussion der Eigenschaften der unterschiedlichen Näherungsverfahren.

8.2 Fundamentalsatz der Differenzial- und Integralrechnung



So kompliziert die Konstruktion des bestimmten Integrals auch aussieht; es zeigt sich, dass die Berechnung in vielen Fällen sehr einfach wird. Diese Tatsache verdankt man dem Zusammenhang zwischen der Ableitung der Integralfunktion und dem Integranden, der nun hergeleitet wird. Dazu stellen wir zunächst eine Verallgemeinerung des Mittelwertsatzes der

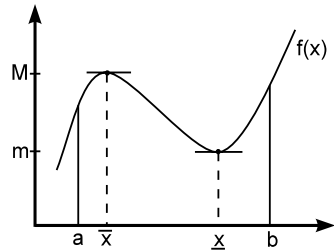
Differenzialrechnung vor, der besagt, dass die Fläche unterhalb einer Kurve

$f(x)$ ersetzt werden kann durch eine flächengleiche Rechtecksfläche mit gleicher Grundseite und mit Höhe $f(\xi)$. Dabei heißt $f(\xi)$ *integraler Mittelwert* der Funktion f im Intervall $[a, b]$:

Mittelwertsatz der Integralrechnung: Ist $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann gibt es ein $\xi \in (a, b)$ mit der Eigenschaft, dass

$$f(\xi) \cdot (b - a) = \int_a^b f(x) dx.$$

Beweis: Zunächst ist aufgrund der Definition des bestimmten Integrals, das Integral über eine konstante Funktion $\int_a^b c dx = c \cdot (b - a)$. Setzen wir $m := \min_{x \in [a, b]} f(x)$ als Minimum und $M := \max_{x \in [a, b]} f(x)$ als das Maximum der Funktion f im Intervall $[a, b]$, so gibt es nach dem Zwischenwertsatz ein $\underline{x}$ mit $f(\underline{x}) = m$ und ein $\bar{x}$ mit $f(\bar{x}) = M$. Damit ist



$$f(\underline{x}) \leq f(x) \leq f(\bar{x}).$$

Da $f(\underline{x})$ und $f(\bar{x})$ konstante Zahlen sind, gilt

$$\begin{aligned} f(\underline{x}) (b - a) &= \int_a^b f(\underline{x}) dx \leq \int_a^b f(x) dx \leq \int_a^b f(\bar{x}) dx = f(\bar{x}) (b - a) \\ \Rightarrow f(\underline{x}) &\leq \frac{1}{b - a} \int_a^b f(x) dx \leq f(\bar{x}). \end{aligned}$$

Nach dem Zwischenwertsatz gibt es dann wiederum ein $\xi \in (a, b)$ mit dem Funktionswert

$$f(\xi) = \frac{1}{b - a} \int_a^b f(x) dx.$$

Bei obiger Betrachtung verwendeten wir die Monotonieeigenschaft des bestimmten Integrals. Sie besagt, dass aus $g(x) \leq f(x) \leq h(x)$ folgt:

$$\int_a^b g(x) dx \leq \int_a^b f(x) dx \leq \int_a^b h(x) dx.$$

Diese Eigenschaft rechnet man aufgrund der algebraischen Definition der Integrale direkt nach. $\square$

Eine allgemeinere Formulierung des Mittelwertsatzes lautet:

Satz: (Allgemeiner Mittelwertsatz der Integralrechnung).

Seien $f, \varphi : [a, b] \rightarrow \mathbb{R}$ stetige Funktionen und $\varphi \geq 0$. Dann gibt es ein $\xi \in (a, b)$ mit der Eigenschaft, dass

$$\int_a^b f(x) \varphi(x) dx = f(\xi) \int_a^b \varphi(x) dx.$$

Wir stellen nun den Zusammenhang zwischen Differenzial- und Integralrechnung her. Dieser Zusammenhang ist nicht nur theoretisch von Bedeutung, er liefert auch eine praktische Methode zur Berechnung von Integralen.

Satz über Integralfunktionen: Ist $f : [a, b] \rightarrow \mathbb{R}$ stetig und $F(x) := \int_a^x f(t) dt$ eine Integralfunktion zu f . Dann ist F differenzierbar und es gilt:

$$F'(x) = f(x).$$

Beweis: Wir betrachten die Differenz der Flächen

$$\Delta F = F(x+h) - F(x) = \int_a^{x+h} f(x) dx - \int_a^x f(x) dx = \int_x^{x+h} f(x) dx$$

und wenden auf das Integral der rechten Seite den Mittelwertsatz der Integralrechnung an: $\int_x^{x+h} f(x) dx = h f(\xi_h)$ mit $\xi_h \in (x, x+h)$. Anschließend bilden wir den Differenzenquotienten

$$\frac{F(x+h) - F(x)}{h} = \frac{1}{h} \Delta F = f(\xi_h).$$

Durch Grenzübergang $h \rightarrow 0$ geht die linke Seite gegen $F'(x)$ und die rechte Seite gegen $f(x)$. Denn $\xi_h \xrightarrow{h \rightarrow 0} x$ und da f stetig ist, gilt: $f(\xi_h) \xrightarrow{h \rightarrow 0} f(x)$. $\square$

Beispiele 8.4:

- ① Für $f(x) = 1$ ist $F(x) = \int_0^x 1 dt = x$, denn $F'(x) = 1$;
- ② Für $g(x) = x$ ist $G(x) = \int_0^x t dt = \frac{x^2}{2}$, denn $G'(x) = x$;
- ③ Für $h(x) = x^2$ ist $H(x) = \int_0^x t^2 dt = \frac{x^3}{3}$, denn $H'(x) = x^2$. $\square$

② **Stammfunktionen**

Wir haben allgemein gezeigt, dass die erste Ableitung der Integralfunktion $F(x) = \int_a^x f(t) dt$ als Ergebnis den Integranden $f(x)$ liefert. Wir nennen solche Funktionen $F(x)$ mit $F'(x) = f(x)$ *Stammfunktionen*:

Definition: Jede Funktion $F(x)$ mit $F'(x) = f(x)$ heißt **Stammfunktion** von $f(x)$.

Mit diesem Sprachgebrauch kann man den Satz über die Integralfunktionen umformulieren:

Jedes unbestimmte Integral

$$I(x) = \int_a^x f(t) dt$$

ist eine Stammfunktion von $f(x)$, d.h. $I'(x) = f(x)$.

Beispiel 8.5. In folgender Tabelle sind für einige Funktionen $f(x)$ eine Stammfunktion $F(x)$ angegeben. Die Eigenschaft $F' = f$ ist direkt nachzurechnen.

$f(x)$	$x^n, n \neq -1$	x^{-1}	$\sqrt[n]{x}$	e^x	$\sin x$	$\cos x$
$F(x)$	$\frac{x^{n+1}}{n+1}$	$\ln x$	$\frac{x^{\frac{1}{n}+1}}{\frac{1}{n}+1}$	e^x	$-\cos x$	$\sin x$

Tabelle 8.2: Elementare Stammfunktionen

□

Bemerkung: Zu jeder stetigen Funktion gibt es unendlich viele Stammfunktionen, denn z.B. zu x^n ist sowohl $\frac{1}{n+1} x^{n+1}$, als auch $\frac{1}{n+1} x^{n+1} + 2$, als auch $\frac{1}{n+1} x^{n+1} + C$ eine Stammfunktion. Allerdings unterscheiden sich zwei Stammfunktionen zu einer Funktion f immer nur durch eine Konstante:

Satz: Sind F_1 und F_2 zwei Stammfunktionen von f , so stimmen sie bis auf eine additive Konstante $C \in \mathbb{R}$ überein:

$$F_1(x) = F_2(x) + C.$$

Beweis: Da F_1 und F_2 Stammfunktionen zu f , folgt $F_1'(x) = f(x) = F_2'(x)$.
 $\Rightarrow (F_1(x) - F_2(x))' = 0 \Rightarrow F_1(x) - F_2(x) = \text{const.}$ □

Folglich lässt sich jedes unbestimmte Integral schreiben in der Form

$$\int_a^x f(t) dt = F(x) + C,$$

wobei $F(x)$ irgendeine Stammfunktion und C eine Konstante ist. Zu jeder stetigen Funktion $f(x)$ gibt es also unendlich viele unbestimmte Integrale. Daher kennzeichnet man diese Funktionenschar durch das Weglassen der Integrationsgrenzen

$$\int f(x) dx = \{ \text{Menge aller unbestimmten Integrale von } f(x) \}.$$

Da sich alle Stammfunktionen nur durch eine Konstante unterscheiden, schreibt man **kurz**

$$\int f(x) dx = F(x) + C,$$

und nennt C die *Integrationskonstante*.

Beispiele 8.6:

- ① $\int e^x dx = e^x + C.$
- ② $\int x^k dx = \frac{1}{k+1} x^{k+1} + C \quad (k \neq -1).$
- ③ $\int \cos x dx = \sin x + C.$

⚠ Bemerkung: Ein Grund, warum die Integralrechnung schwieriger als die Differenzialrechnung ist, liegt darin, dass sich nicht jede Stammfunktion durch elementare Funktionen darstellen lässt. Die Funktionen

$$f(x) = e^{x^2}, \quad f(x) = \frac{\sin x}{x}$$

besitzen keine elementar darstellbaren Stammfunktionen!

Tabelle von Stammfunktionen. In der folgenden Tabelle sind für viele elementare Funktionen die Menge aller Stammfunktionen $\int f(x) dx = F(x)$ zusammengestellt. Die Gültigkeit kann oftmals sehr einfach mit der Beziehung $F'(x) = f(x)$ bestätigt werden.

Tabelle 8.3: Stammfunktionen

$f(x) = F'(x)$	Stammfunktion F : $F(x) = \int f(x) dx + C$	Definitionsbereich $\mathbb{D}_f$
k ($k \in \mathbb{R}$)	$kx + C$	$\mathbb{R}$
x^α ($\alpha \neq -1$)	$\frac{1}{\alpha+1} x^{\alpha+1} + C$	$\mathbb{R}_{>0}$
x^{-1}	$\ln x + C$	$\mathbb{R} \setminus \{0\}$
$\sin x$	$-\cos x + C$	$\mathbb{R}$
$\cos x$	$\sin x + C$	$\mathbb{R}$
$\tan x$	$-\ln \cos x + C$	$\mathbb{R} \setminus \{x = \frac{\pi}{2} + k\pi, k \in \mathbb{Z}\}$
$\cot x$	$\ln \sin x + C$	$\mathbb{R} \setminus \{x = k\pi, k \in \mathbb{Z}\}$
a^x ($a > 0, \neq 1$)	$\frac{a^x}{\ln a} + C$	$\mathbb{R}$
e^x	$e^x + C$	$\mathbb{R}$
e^{ax} ($a \neq 0$)	$\frac{1}{a} e^{ax} + C$	$\mathbb{R}$
$\ln x$	$x \cdot \ln x - x + C$	$\mathbb{R}_{>0}$
$\frac{1}{\cos^2 x}$	$\tan x + C$	$\mathbb{R} \setminus \{x = \frac{\pi}{2} + k\pi, k \in \mathbb{Z}\}$
$\frac{1}{\sin^2 x}$	$-\cot x + C$	$\mathbb{R} \setminus \{x = k\pi, k \in \mathbb{Z}\}$
$\sin^2 x$	$\frac{1}{2} (x - \sin x \cdot \cos x) + C$	$\mathbb{R}$
$\cos^2 x$	$\frac{1}{2} (x + \sin x \cdot \cos x) + C$	$\mathbb{R}$
$\tan^2 x$	$\tan x - x + C$	$\mathbb{R} \setminus \{x = \frac{\pi}{2} + k\pi, k \in \mathbb{Z}\}$
$\cot^2 x$	$-\cot x - x + C$	$\mathbb{R} \setminus \{x = k\pi, k \in \mathbb{Z}\}$

$f(x) = F'(x)$	Stammfunktion F : $F(x) = \int f(x) dx + C$	Definitionsbereich $\mathbb{D}_f$
$\arcsin x$	$x \cdot \arcsin x + \sqrt{1-x^2} + C$	$(-1, 1)$
$\arccos x$	$x \cdot \arccos x - \sqrt{1-x^2} + C$	$(-1, 1)$
$\arctan x$	$x \cdot \arctan x - \frac{1}{2} \ln(x^2 + 1) + C$	$\mathbb{R}$
$\operatorname{arccot} x$	$x \cdot \operatorname{arccot} x + \frac{1}{2} \ln(x^2 + 1) + C$	$\mathbb{R}$
$\frac{1}{\sqrt{1-x^2}}$	$\arcsin x + C$	$(-1, 1)$
$\frac{-1}{\sqrt{1-x^2}}$	$\arccos x + C$	$(-1, 1)$
$\frac{1}{1+x^2}$	$\arctan x + C$	$\mathbb{R}$
$\frac{-1}{1+x^2}$	$\operatorname{arccot} x + C$	$\mathbb{R}$

$f(x) = F'(x)$	Stammfunktion $F(x) = \int f(x) dx + C$	Definitionsbereich $\mathbb{D}_f$
$\sinh x$	$\cosh x + C$	$\mathbb{R}$
$\cosh x$	$\sinh x + C$	$\mathbb{R}$
$\tanh x$	$\ln(\cosh x) + C$	$\mathbb{R}$
$\coth x$	$\ln \sinh x + C$	$\mathbb{R} \setminus \{0\}$
$\frac{1}{\cosh^2 x}$	$\tanh x + C$	$\mathbb{R}$
$\frac{1}{\sinh^2 x}$	$-\coth x + C$	$\mathbb{R} \setminus \{0\}$
$\frac{1}{\sqrt{1+x^2}}$	$\operatorname{ar} \sinh x + C$	$\mathbb{R}$
$\frac{1}{\sqrt{x^2-1}}$	$\operatorname{ar} \cosh x + C$	$(1, \infty)$
$\frac{1}{1-x^2}$	$\operatorname{ar} \tanh x + C$	$(-1, 1)$
$\frac{1}{1-x^2}$	$\operatorname{ar} \coth x + C$	$\mathbb{R} \setminus [-1, 1]$

Stammfunktionen sind aber weder geometrisch noch physikalisch so bedeutsam wie die bestimmten Integrale. Bisher haben wir erst ein einziges explizit berechnet, nämlich $\int_0^b x^2 dx = \frac{b^3}{3}$. Durch den folgenden Hauptsatz der Differenzial- und Integralrechnung wird die schwierige Aufgabe, der Berechnung von bestimmten Integralen auf eine einfachere Aufgabe, nämlich das Aufsuchen von Stammfunktionen, zurückgeführt:

Fundamentalsatz der Differenzial- und Integralrechnung

Ist $f : [a, b] \rightarrow \mathbb{R}$ stetig und F eine Stammfunktion von f . Dann gilt

$$\int_a^b f(x) dx = F(b) - F(a).$$

Beweis: Sei $x \in [a, b]$ und $F_0(x) := \int_a^x f(t) dt$. Dann ist $F_0(x)$ eine Stammfunktion von f mit $F_0(a) = 0$ und $F_0(b) = \int_a^b f(t) dt$. Ist $F(x)$ eine beliebige Stammfunktion von f , so folgt $F - F_0 = \text{const} = c$ und

$$\begin{aligned} F(b) - F(a) &= F_0(b) + c - (F_0(a) + c) \\ &= F_0(b) - F_0(a) = F_0(b) = \int_a^b f(t) dt. \end{aligned} \quad \square$$

Damit erfolgt die Berechnung von bestimmten Integralen in zwei Schritten:

Berechnung von bestimmten Integralen

- (1) Man bestimme eine Stammfunktion $F(x)$ zum Integranden $f(x)$.
- (2) Mit dieser Stammfunktion berechnet man die Differenz $F(b) - F(a)$:

$$\int_a^b f(x) dx = [F(x)]_a^b = F(x)|_a^b = F(b) - F(a).$$

Hierbei ist $[F(x)]_a^b = F(x)|_a^b$ eine abkürzende Schreibweise für die Differenz $F(b) - F(a)$.

Beispiele 8.7 (Berechnung bestimmter Integrale):

- ① $\int_a^b x^3 dx = ?$: Eine Stammfunktion von x^3 ist nach Tab. 8.3 $\frac{1}{4}x^4$, so dass

$$\int_a^b x^3 dx = \left. \frac{1}{4} x^4 \right|_a^b = \frac{1}{4} (b^4 - a^4).$$

Für $0 \leq a < b$ ist dies der Flächeninhalt der Kurve $y = x^3$ mit der x -Achse im Bereich von $a \leq x \leq b$.

- ② $\int_0^\pi \sin x \, dx = ?$: Eine Stammfunktion von $\sin x$ ist nach Tab. 8.3 $F(x) = -\cos x$, so dass

$$\int_0^\pi \sin x \, dx = -\cos x \Big|_0^\pi = -\cos \pi - (-\cos(0)) = 1 - (-1) = 2.$$

Dies ist der Flächeninhalt unter der Sinuskurve in der ersten Halbperiode.

Anwendungsbeispiel 8.8 (Ausdehnungsarbeit eines Gases).

In einem Zylinder der Grundfläche $F \, [\text{cm}^2]$ befinde sich ein durch einen beweglichen Kolben komprimiertes Gas. Wenn der Kolben den Abstand $x \, [\text{cm}]$ vom Zylinderboden hat, sei der Gasdruck im Zylinder $p(x) \, [\frac{\text{g}}{\text{cm} \cdot \text{s}^2}]$. Bei Verschiebung des Kolbens von $x = a$ nach $x = b$ wird vom Gas *Arbeit* geleistet, die gegeben ist durch

$$A = \int_a^b F p(x) \, dx.$$

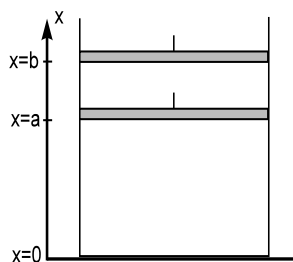
Als einfachen Sonderfall betrachten wir die *isotherme Ausdehnung* eines idealen Gases mit der Zustandsgleichung

$$p(x) \cdot V(x) = p(a) \cdot V(a) = \text{const} \\ (\text{Boyle-Mariottesches Gesetz})$$

Mit dem Volumen $V(x) = F \cdot x$, folgt

$$p(x) = \frac{p(a) \cdot V(a)}{V(x)} = \frac{p(a) \cdot V(a)}{F \cdot x}.$$

$$A = F \int_a^b \frac{p(a) \cdot V(a)}{F \cdot x} \, dx = p(a) \cdot V(a) \int_a^b \frac{1}{x} \, dx.$$



Nach Tab. 8.3 ist dann

$$\begin{aligned} A &= p(a) \cdot V(a) \cdot [\ln x]_a^b = p(a) \cdot V(a) \cdot [\ln(b) - \ln(a)] \\ &= p(a) \cdot V(a) \cdot \ln \frac{b}{a}. \end{aligned}$$

□

8.3 Grundlegende Regeln der Integralrechnung

Die Berechnung von bestimmten Integralen vereinfacht sich mit Hilfe von Integrationsregeln. Sie ergeben sich unmittelbar aus der Definition des bestimmten Integrals als Grenzwert der Zwischensumme. Die auftretenden Funktionen werden als stetig vorausgesetzt.

Faktorregel: Ein konstanter Faktor c darf vor das Integral gezogen werden:

$$\int_a^b c f(x) dx = c \int_a^b f(x) dx.$$

Beispiel 8.9.
$$\int_0^{\pi/2} 4 \cos x dx = 4 \int_0^{\pi/2} \cos x dx$$

$$= 4 [\sin x]_0^{\pi/2} = 4 \left(\sin\left(\frac{\pi}{2}\right) - \sin(0) \right) = 4. \quad \square$$

Summenregel: Eine Summe von Funktionen darf gliedweise integriert werden:

$$\int_a^b (f_1(x) + f_2(x)) dx = \int_a^b f_1(x) dx + \int_a^b f_2(x) dx.$$

Beispiel 8.10.
$$\int_0^1 (-3x^2 + x) dx = -3 \int_0^1 x^2 dx + \int_0^1 x dx$$

$$= -3 \left[\frac{x^3}{3} \right]_0^1 + \left[\frac{x^2}{2} \right]_0^1 = -\frac{1}{2}. \quad \square$$

Bemerkungen:

- (1) Die Faktor- und Summenregel gelten sinngemäß auch für unbestimmte Integrale.
- (2) Bisher war stets $a < b$ vorausgesetzt. Die Faktor- und Summenregel bleiben gültig für beliebige reelle Zahlen a, b aus dem Definitionsbereich von f , wenn man folgende Definition hinzunimmt:

Definition:

(1) Zusammenfallen der Integrationsgrenzen: $\int_a^a f(x) dx := 0.$

(2) Vertauschen der Integrationsgrenzen: $\int_b^a f(x) dx := -\int_a^b f(x) dx.$

Beispiele 8.11:

$$\textcircled{1} \quad \int_2^2 \frac{1}{x} dx = \ln x \Big|_2^2 = \ln(2) - \ln(2) = 0.$$

$$\textcircled{2} \quad \int_{\pi/2}^0 \sin x dx = -\int_0^{\pi/2} \sin x dx = -[-\cos x]_0^{\pi/2} = -1. \quad \square$$

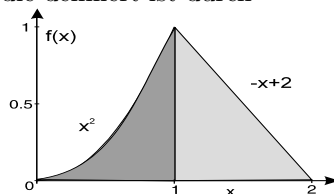
Additivität des Integrals: Für jede beliebige Stelle c aus dem Integrationsbereich $a \leq x \leq b$ von f gilt

$$\int_a^b f(x) dx = \int_a^c f(x) dx + \int_c^b f(x) dx.$$

Die Additivität des Integrals nutzt man aus, wenn eine Funktion auf Teilintervallen unterschiedliche Funktionsvorschriften besitzt.

Beispiel 8.12. Gegeben ist die Funktion $f(x)$, die definiert ist durch

$$f(x) = \begin{cases} x^2 & \text{für } 0 \leq x \leq 1 \\ -x+2 & \text{für } 1 \leq x \leq 2 \end{cases}$$



Um die Fläche unterhalb der Kurve f im Bereich $[0, 2]$ zu berechnen, muss man das Integral aufsplitten in ein Integral über $[0, 1]$ und $[1, 2]$, da die Funktion in beiden Teilintervallen eine unterschiedliche Funktionsvorschrift hat. Daher ist

$$A = \int_0^2 f(x) dx = \int_0^1 x^2 dx + \int_1^2 (-x+2) dx$$

$$= \left[\frac{x^3}{3} \right]_0^1 + \left[-\frac{x^2}{2} + 2x \right]_1^2 = \frac{1}{3} + \frac{1}{2} = \frac{5}{6}. \quad \square$$

8.4 Integrationsmethoden

Die Integration von Funktionen erweist sich in praktischen Fällen oftmals schwieriger als die Differenziation. Während sich das Differenzieren durch Anwendung einfacher Regeln (Produkt-, Quotienten-, Kettenregel) erledigen lässt, ist das Integrieren mit größeren Schwierigkeiten verbunden. Trotzdem kann in vielen Fällen durch eine der folgenden Integrationsmethoden eine Stammfunktion gefunden werden.

► 8.4.1 Partielle Integration

Die *partielle Integration* ist das Pendant zur Produktregel der Differenziation, welche besagt, dass

$$(u(x) \cdot v(x))' = u'(x) v(x) + u(x) v'(x).$$

Wir lösen diese Gleichung nach $u(x) v'(x)$ auf und integrieren anschließend

$$u(x) v'(x) = (u(x) v(x))' - u'(x) v(x)$$

$$\int_a^b u(x) v'(x) dx = \int_a^b (u(x) v(x))' dx - \int_a^b u'(x) v(x) dx.$$

Nach dem Fundamentalsatz der Differenzial- und Integralrechnung ist

$$\int_a^b (u(x) v(x))' dx = [u(x) \cdot v(x)]_a^b,$$

so dass gilt

Partielle Integration:

$$\int_a^b u(x) v'(x) dx = [u(x) v(x)]_a^b - \int_a^b u'(x) v(x) dx.$$

Bemerkungen:

- (1) Ob die Integration nach der Methode der partiellen Integration gelingt, hängt von der "richtigen" (geeigneten) Wahl von $u(x)$ und $v'(x)$ ab.
- (2) In manchen Fällen muss das Integrationsverfahren mehrmals angewendet werden, ehe man auf ein Grundintegral stößt.
- (3) Insbesondere bei der Integration von Funktionen, die als einen Faktor eine trigonometrische Funktion enthalten, tritt nach ein- bzw. mehrmaliger partieller Integration der Fall auf, dass das zu berechnende Integral, mit

einem Faktor versehen, auf der rechten Seite wieder auftritt. In diesem Fall löst man die Gleichung nach dem gesuchten Integral auf.

- (4) Die Formel der partiellen Integration gilt auch für unbestimmte Integrale

$$\int u(x) v'(x) dx = u(x) v(x) - \int u'(x) v(x) dx.$$

Beispiele 8.13 (Partiellen Integration):

- ① Gesucht ist $\int_1^2 x e^x dx$.

Wir setzen

$$\begin{array}{ll} u(x) = x & \Rightarrow u'(x) = 1 \\ v'(x) = e^x & \Rightarrow v(x) = e^x \end{array}$$

und erhalten

$$\begin{aligned} \int_1^2 x e^x dx &= [x e^x]_1^2 - \int_1^2 1 \cdot e^x dx = [x e^x]_1^2 - [e^x]_1^2 \\ &= 2e^2 - e^1 - e^2 + e^1 = e^2. \end{aligned}$$

Ferner gilt $\int x e^x dx = e^x (x - 1) + C$.

- ② Gesucht ist $\int x^2 \cos x dx$.

Wir setzen

$$\begin{array}{ll} u(x) = x^2 & \Rightarrow u'(x) = 2x \\ v'(x) = \cos x & \Rightarrow v(x) = \sin x \end{array}$$

und erhalten

$$\int x^2 \cos x dx = x^2 \sin x - \int 2x \sin x dx.$$

Nochmalige partielle Integration von $\int 2x \sin x dx$ liefert mit

$$\begin{array}{ll} u(x) = 2x & \Rightarrow u'(x) = 2 \\ v'(x) = \sin x & \Rightarrow v(x) = -\cos x \end{array}$$

$$\begin{aligned} \int x^2 \cos x dx &= x^2 \sin x - \left[2x (-\cos x) - \int 2 (-\cos x) dx \right] \\ &= x^2 \sin x + 2x \cos x - 2 \sin x + C. \end{aligned}$$

In der Regel setzt man $u(x)$ gleich dem Potenzfaktor, um so durch mehrmalige partielle Integration diesen Term zum Verschwinden zu bringen. In manchen Fällen führt aber $v'(x) = 1 \Rightarrow v(x) = x$ zum Ziel:

③ Gesucht ist $\int \ln x \, dx$.

Mit

$$\begin{array}{ll} u(x) = \ln x & \Rightarrow \quad u'(x) = \frac{1}{x} \\ v'(x) = 1 & \Rightarrow \quad v(x) = x \end{array}$$

folgt

$$\begin{aligned} \int \ln x \, dx &= x \ln x - \int \frac{1}{x} \cdot x \, dx = x \ln x - x + C \\ &= x (\ln x - 1) + C. \end{aligned}$$

④ Gesucht ist $\int \cos^2 x \, dx = \int \cos x \cdot \cos x \, dx$.

Mit

$$\begin{array}{ll} u(x) = \cos x & \Rightarrow \quad u'(x) = -\sin x \\ v'(x) = \cos x & \Rightarrow \quad v(x) = \sin x \end{array}$$

ist

$$\int \cos^2 x \, dx = \cos x \sin x - \int -\sin x \sin x \, dx = \cos x \sin x + \int \sin^2 x \, dx.$$

Wir ersetzen $\sin^2 x = 1 - \cos^2 x$.

$$\Rightarrow \int \cos^2 x \, dx = \cos x \sin x + x - \int \cos^2 x \, dx.$$

Addieren wir $\int \cos^2 x \, dx$ auf beiden Seiten und dividieren anschließend durch den Faktor 2, folgt

$$\int \cos^2 x \, dx = \frac{1}{2} (\sin x \cos x + x) + C.$$

► 8.4.2 Integration durch Substitution

Ähnlich wie die partielle Integration auf der Produktregel basiert, lässt sich aus der Kettenregel die *Integralsubstitutions-Methode* herleiten. Mit $y = f(x)$ folgt für die Ableitung der Funktion $g(y) = g(f(x))$ nach x :

$$\frac{d}{dx} g(f(x)) = g'(f(x)) \cdot f'(x).$$

Hieraus ergibt sich dann durch Integration:

Substitutionsregel für unbestimmte Integrale:

$$\int g(f(x)) f'(x) dx = G(f(x)) + C,$$

wenn G eine Stammfunktion von g ist.

In der folgenden Tabelle sind einfache Spezialfälle dieser allgemeinen Substitutionsregel angegeben. Sie lassen sich in konkreten Beispielen einfacher anwenden. Man beachte, dass die Fälle (A), (B) und (C) als Spezialfälle in (D) enthalten sind.

	Integraltyp	Substitution	Lösung
(A)	$\int g(ax+b) dx$	$y = ax+b$	$\frac{1}{a} G(ax+b) + C$
(B)	$\int f(x) f'(x) dx$	$y = f(x)$	$\frac{1}{2} f^2(x) + C$
(C)	$\int \frac{f'(x)}{f(x)} dx$	$y = f(x)$	$\ln f(x) + C$
(D)	$\int g(f(x)) f'(x) dx$	$y = f(x)$	$G(f(x)) + C$

Tabelle 8.4: Einfache Integralsubstitutionen

Beispiele 8.14 (Integralsubstitutionen nach Tabelle 8.4):

$$(A1) \int_2^3 (2x-3)^4 dx = ?$$

Wir bestimmen zuerst eine Stammfunktion und setzen dann zur Berechnung des bestimmten Integrals die obere und untere Integrationsgrenze ein. Dazu substituiert man $y = 2x - 3$ und ersetzt jeden Term des Integrals, der die Integrationsvariable x enthält, durch einen entsprechenden Term mit y . Insbesondere muss auch das Differenzial dx durch einen entsprechenden Term mit dy ersetzt werden.

Aus $y = 2x - 3 \hookrightarrow y' = \frac{dy}{dx} = 2 \hookrightarrow dx = \frac{1}{2} dy$. Somit ist

$$\int (2x-3)^4 dx = \int y^4 \frac{1}{2} dy = \frac{1}{2} \int y^4 dy = \frac{1}{2} \frac{1}{5} y^5 + C.$$

Nach Berechnung des Integrals wird durch Rücksubstitution y wieder durch $2x - 3$ ersetzt:

$$\int (2x-3)^4 dx = \frac{1}{10} (2x-3)^5 + C.$$

Das bestimmte Integral ist daher

$$\int_2^3 (2x-3)^4 dx = \left[\frac{1}{10} (2x-3)^5 \right]_2^3 = \frac{1}{10} [243 - 1] = 24.2.$$

(A1') Führt man alternativ die Substitutionsmethode direkt beim bestimmten Integral

$$\int_2^3 (2x-3)^4 dx$$

durch, müssen auch die Integrationsgrenzen ersetzt werden! Es erfolgt dann nach der Berechnung des substituierten Integrals keine Rücksubstitution mehr. Aus $y = 2x - 3$ folgt für die untere Grenze $x_u = 2 \hookrightarrow y_u = 1$ und für die obere Grenze $x_o = 3 \hookrightarrow y_o = 3$:

$$\int_2^3 (2x-3)^4 dx = \int_1^3 y^4 \frac{1}{2} dy = \left[\frac{1}{10} y^5 \right]_1^3 = \frac{1}{10} [243 - 1] = 24.2.$$

(A2) $\int_0^1 \frac{1}{1+4x} dx = ?$

Man substituiere $\boxed{y = 1 + 4x} \hookrightarrow y' = \frac{dy}{dx} = 4 \hookrightarrow dx = \frac{1}{4} dy.$

$$\hookrightarrow \int \frac{1}{1+4x} dx = \int \frac{1}{y} \cdot \frac{1}{4} dy = \frac{1}{4} \int \frac{1}{y} dy = \frac{1}{4} \ln |y| + C.$$

Durch Rücksubstitution $y = 1 + 4x$ folgt

$$\int \frac{1}{1+4x} dx = \frac{1}{4} \ln |1 + 4x| + C.$$

Für das bestimmte Integral gilt

$$\int_0^1 \frac{1}{1+4x} dx = \left[\frac{1}{4} \ln |1 + 4x| \right]_0^1 = \frac{1}{4} \ln 5 - \frac{1}{4} \ln 1 = \frac{1}{4} \ln 5.$$

(B1) $\int \sin x \cos x dx = ?$

Man substituiere $\boxed{y = \sin x} \hookrightarrow y' = \frac{dy}{dx} = \cos x \hookrightarrow dx = \frac{1}{\cos x} dy.$

$$\hookrightarrow \int \sin x \cos x dx = \int y \cos x \frac{1}{\cos x} dy = \int y dy = \frac{1}{2} y^2 + C.$$

Durch Rücksubstitution $y = \sin x$ folgt

$$\int \sin x \cos x dx = \frac{1}{2} \sin^2 x + C.$$

$$(B2) \quad \int_1^2 \frac{\ln x}{x} dx = ?$$

Substitution: $\boxed{y = \ln x} \quad \hookrightarrow y' = \frac{dy}{dx} = \frac{1}{x} \quad \hookrightarrow dx = x dy.$

$$\hookrightarrow \int \frac{\ln x}{x} dx = \int \frac{y}{x} \cdot x dy = \int y dy = \frac{1}{2} y^2 + C.$$

Rücksubstitution: $\int \frac{\ln x}{x} dx = \frac{1}{2} \ln^2 x + C.$

Berechnung des bestimmten Integrals:

$$\int_1^2 \frac{\ln x}{x} dx = \frac{1}{2} [\ln^2 x]_1^2 = \frac{1}{2} \ln^2 2.$$

$$(C1) \quad \int \frac{2x-3}{x^2-3x+1} dx = \int \frac{2x-3}{y} \cdot \frac{dy}{2x-3} \\ = \int \frac{1}{y} dy = \ln |y| + C = \ln |x^2 - 3x + 1| + C$$

mit der Substitution $\boxed{y = x^2 - 3x + 1}$ und $dx = \frac{1}{2x-3} dy.$

$$(C2) \quad \int_0^1 \frac{e^x}{2e^x+5} dx = ?$$

Substitution: $\boxed{y = 2e^x + 5} \quad \hookrightarrow y' = \frac{dy}{dx} = 2e^x \quad \hookrightarrow dx = \frac{1}{2e^x} dy.$

obere Grenze $x_o = 1 \hookrightarrow y_o = 2e + 5$

untere Grenze $x_u = 0 \hookrightarrow y_u = 7.$

$$\begin{aligned} \int_0^1 \frac{e^x}{2e^x+5} dx &= \int_7^{2e+5} \frac{e^x}{y} \cdot \frac{1}{2e^x} dy = \frac{1}{2} \int_7^{2e+5} \frac{1}{y} dy \\ &= \frac{1}{2} \ln y \Big|_7^{2e+5} = \frac{1}{2} [\ln(2e+5) - \ln 7] = 0.1997. \end{aligned}$$

$$(D1) \quad \int (x^3 + 2)^{\frac{1}{2}} x^2 dx = ?$$

Substitution: $\boxed{y = x^3 + 2} \quad \hookrightarrow \frac{dy}{dx} = 3x^2 \quad \hookrightarrow dx = \frac{1}{3x^2} dy.$

$$\int (x^3 + 2)^{\frac{1}{2}} x^2 dx = \int y^{\frac{1}{2}} \cdot x^2 \cdot \frac{1}{3x^2} dy = \frac{1}{3} \int y^{\frac{1}{2}} dy = \frac{1}{3} \cdot \frac{2}{3} y^{\frac{3}{2}} + C.$$

Rücksubstitution:

$$\int (x^3 + 2)^{\frac{1}{2}} x^2 dx = \frac{2}{9} (x^3 + 2)^{\frac{3}{2}} + C.$$

$$(D2) \int \frac{e^x + x e^x}{(x e^x)^3} dx = ?$$

Substitution: $\boxed{y = x e^x} \quad \hookrightarrow \frac{dy}{dx} = e^x + x e^x \quad \hookrightarrow dx = \frac{1}{e^x + x e^x} dy.$

$$\int \frac{e^x + x e^x}{(x e^x)^3} dx = \int \frac{e^x + x e^x}{y^3} \cdot \frac{1}{e^x + x e^x} dy = \int y^{-3} dy = -\frac{1}{2} y^{-2} + C.$$

Rücksubstitution: $\int \frac{e^x + x e^x}{(x e^x)^3} dx = -\frac{1}{2} (x e^x)^{-2} + C. \quad \square$

Man beachte: Wenn bei einem bestimmten Integral eine Substitution durchgeführt wird, müssen auch die Integrationsgrenzen ersetzt werden. Dafür erspart man sich zum Schluss die Rücksubstitution. Außer den angegebenen Substitutionsregeln gibt es noch viele andere.

Tabelle 8.5: Weitere Integralsubstitutionen

	Integraltyp	Substitution
(E)	$\int g(x, \sqrt{a^2 - x^2}) dx$	$x = a \cdot \sin y$
(F)	$\int g(x, \sqrt{x^2 + a^2}) dx$	$x = a \cdot \sinh y$
(G)	$\int g(x, \sqrt{x^2 - a^2}) dx$	$x = a \cdot \cosh y$

Beispiele 8.15 (Integralsubstitutionen nach Tabelle 8.5):

$$(E1) \int \frac{dx}{\sqrt{4 - x^2}} = \int \frac{2 \cos(y)}{2 \cos(y)} dy = \int dy = y + C = \arcsin\left(\frac{1}{2}x\right) + C$$

mit der Substitution $\boxed{x = 2 \sin(y)} \quad \hookrightarrow \frac{dx}{dy} = 2 \cos(y) \hookrightarrow dx = 2 \cos(y) dy$

und $\sqrt{4 - x^2} = \sqrt{4 - 4 \sin^2(y)} = 2 \cos(y)$, da $\cos^2(y) + \sin^2(y) = 1$.

$$\begin{aligned}
 \text{(E2)} \quad \int \frac{x}{\sqrt{4-x^2}} dx &= \int \frac{2 \sin(y) \cdot 2 \cos(y)}{2 \cos(y)} dy = 2 \int \sin(y) dy \\
 &= -2 \cos(y) + C = -2 \sqrt{1 - \sin^2(y)} + C \\
 &= -2 \sqrt{1 - \frac{x^2}{4}} + C = -\sqrt{4-x^2} + C
 \end{aligned}$$

mit der gleichen Substitution wie unter (E1): $x = 2 \sin(y)$
 und $y = \arcsin(\frac{x}{2})$.

$$\begin{aligned}
 \text{(F)} \quad \int \frac{dx}{\sqrt{1+x^2}} &= \int \frac{\cosh(y)}{\cosh(y)} dy = \int dy = \\
 y + C &= \operatorname{ar} \sinh(x) + C = \ln(x + \sqrt{1+x^2}) + C
 \end{aligned}$$

mit der Substitution $x = \sinh(y)$ $\hookrightarrow \frac{dx}{dy} = \cosh(y) \hookrightarrow dx = \cosh(y) dy$

und $\sqrt{1+x^2} = \sqrt{1+\sinh^2(y)} = \cosh(y)$, da $\cosh^2(y) - \sinh^2(y) = 1$.

$$\begin{aligned}
 \text{(G)} \quad \int \frac{dx}{\sqrt{x^2-25}} &= \int \frac{5 \sinh(y)}{5 \sinh(y)} dy = \int dy = \\
 y + C &= \operatorname{ar} \cosh\left(\frac{x}{5}\right) + C = \ln\left(\frac{x}{5} + \sqrt{\left(\frac{x}{5}\right)^2 + 1}\right) + C
 \end{aligned}$$

mit der Substitution $x = 5 \cosh(y)$ $\hookrightarrow \frac{dx}{dy} = 5 \sinh(y) \hookrightarrow dx = 5 \sinh(y) dy$

und $\sqrt{x^2-25} = \sqrt{25 \cosh^2(y) - 25} = 5 \sqrt{\cosh^2(y) - 1} = 5 \sinh(y)$,

da $\cosh^2(y) - \sinh^2(y) = 1$. □

Trotz der Vielfalt der Substitutionen gibt es bei der Berechnung von Integralen keine allgemeinen Rezepte, die stets zum Ziel führen!

8.4.3 Partialbruchzerlegung

Für rationale Funktionen $f(x) = \frac{Z(x)}{N(x)}$ ($Z(x)$, $N(x)$ Polynome) gibt es eine spezielle Integrationstechnik, die sog. *Partialbruchzerlegung*. Durch diese Methode lassen sich rationale Funktionen in geschlossener Form integrieren. Zur Durchführung müssen sie in einer echt gebrochenrationalen Darstellung vorliegen. Eine rationale Funktion heißt *echt gebrochenrational*, wenn der Grad des Zählerpolynoms kleiner dem Grad des Nennerpolynoms ist, sonst heißt die Funktion *unecht gebrochen*. Bei einer unecht gebrochenrationalen Funktion sorgt man durch *Polynomdivision* dafür, dass anschließend der Grad des Zählers kleiner als der Grad des Nenners ist:

Beispiel 8.16. $\frac{2x^3 - 2x^2 - 5x + 7}{x^2 - 3x + 2} :$

$$\begin{array}{r}
 (2x^3 \quad -2x^2 \quad -5x \quad +7) : (x^2 - 3x + 2) = 2x + 4 + \frac{3x - 1}{x^2 - 3x + 2} \\
 \underline{-(2x^3 \quad -6x^2 \quad +4x)} \\
 \quad 4x^2 \quad -9x \quad +7 \\
 \quad \underline{-(4x^2 \quad -12x \quad +8)} \\
 \qquad \qquad 3x \quad -1
 \end{array}$$

□

Eine echt gebrochenrationale Funktion lässt sich eindeutig in Partialbrüche zerlegen. Wir gehen im Folgenden immer davon aus, dass

$$f(x) = \frac{p(x)}{q(x)}$$

eine echt gebrochenrationale Funktion mit $\text{Grad}(p) < \text{Grad}(q) = n$ ist.

Satz: Hat $q(x) = a_n (x - x_1)(x - x_2) \cdots (x - x_n)$ n **einfache reelle Nullstellen**, dann führt der Ansatz

$$f(x) = \frac{A_1}{x - x_1} + \frac{A_2}{x - x_2} + \cdots + \frac{A_n}{x - x_n}$$

zu einer eindeutigen Zerlegung von $f(x)$ in **Partialbrüche** $\frac{A_i}{x - x_i}$.

Beispiel 8.17 (Musterbeispiel). $\int \frac{3x - 1}{x^2 - 3x + 2} dx = ?$

Um eine Partialbruchzerlegung durchführen zu können, müssen zuerst die Nullstellen des Nenners von $f(x) = \frac{3x - 1}{x^2 - 3x + 2}$ bestimmt werden. Aus

$$x^2 - 3x + 2 = 0$$

folgt durch die pq -Formel $x_{1/2} = \frac{3}{2} \pm \sqrt{\frac{9}{4} - 2}$. Daher sind $x_1 = 1$ und $x_2 = 2$ die Nullstellen des Nennerpolynoms und $f(x) = \frac{3x-1}{x^2-3x+2} = \frac{3x-1}{(x-1)(x-2)}$. Durch den Ansatz

$$f(x) = \frac{3x-1}{(x-1)(x-2)} = \frac{A_1}{x-1} + \frac{A_2}{x-2}$$

erhält man die Partialbruchzerlegung. Um A_1 und A_2 zu berechnen, bildet man den Hauptnenner

$$\frac{A_1}{x-1} + \frac{A_2}{x-2} = \frac{A_1(x-2) + A_2(x-1)}{(x-1)(x-2)} \stackrel{!}{=} \frac{3x-1}{x^2-3x+2}$$

und vergleicht den Zähler mit $3x-1$:

$$A_1(x-2) + A_2(x-1) = 3x-1 \quad \text{für alle } x.$$

Zur Bestimmung der Konstanten A_1 und A_2 führt man entweder einen Koeffizientenvergleich durch oder man setzt spezielle Werte für x ein:

$$\begin{array}{ll} x=1: & A_1(1-2) = 3-1 \\ x=2: & A_2(2-1) = 5 \end{array} \quad \begin{array}{l} \Rightarrow A_1 = -2 \\ \Rightarrow A_2 = 5. \end{array}$$

Also ist

$$f(x) = \frac{-2}{x-1} + \frac{5}{x-2}$$

und damit

$$\begin{aligned} \int f(x) dx &= -2 \int \frac{1}{x-1} dx + 5 \int \frac{1}{x-2} dx \\ &= -2 \ln|x-1| + 5 \ln|x-2| + C \\ &= \ln|x-1|^{-2} + \ln|x-2|^5 + C = \ln \left| \frac{(x-2)^5}{(x-1)^2} \right| + C. \quad \square \end{aligned}$$

Hat das Nennerpolynom eine Nullstelle, die doppelt oder mehrfach auftritt, dann muss die Partialbruchzerlegung für diese Nullstelle modifiziert werden:

Satz: $q(x)$ hat **mehrfache reelle Nullstellen**. Sei x_l eine k -fache Nullstelle, d.h. neben anderen Nullstellen tritt der Term $(x-x_l)$ mit der Potenz k in der Produktdarstellung von $q(x)$ auf. Dann ist diese k -fache Nullstelle neben den anderen folgendermaßen zu berücksichtigen:

$$f(x) = \dots + \frac{B_1}{x-x_l} + \frac{B_2}{(x-x_l)^2} + \dots + \frac{B_k}{(x-x_l)^k}.$$

Beispiel 8.18. $\int \frac{2x^2 + 3x + 1}{x^3 - 5x^2 + 8x - 4} dx = ?$

$x = 1$ ist eine einfache und $x = 2$ eine doppelte Nullstelle des Nennerpolynoms, da $x^3 - 5x^2 + 8x - 4 = (x - 1)(x - 2)^2$. Für die Partialbruchzerlegung des Integranden f wählen wir daher den Ansatz

$$\begin{aligned} f(x) &= \frac{A}{x-1} + \frac{B_1}{x-2} + \frac{B_2}{(x-2)^2} \\ &= \frac{A(x-2)^2 + B_1(x-2)(x-1) + B_2(x-1)}{(x-1)(x-2)^2}. \end{aligned}$$

Nach Multiplikation mit dem Hauptnenner folgt

$$2x^2 + 3x + 1 \stackrel{!}{=} A(x-2)^2 + B_1(x-2)(x-1) + B_2(x-1).$$

Wir setzen zur Bestimmung von A , B_1 und B_2 spezielle x -Werte ein:

$$x = 1 : \boxed{6 = A}$$

$$x = 2 : \boxed{15 = B_2}$$

$$x = 0 : 1 = 4A + 2B_1 - B_2 \Rightarrow 1 = 9 - 2B_1 \Rightarrow \boxed{B_1 = -4}$$

Folglich ist

$$\begin{aligned} &\int \frac{2x^2 + 3x + 1}{x^3 - 5x^2 + 8x - 4} dx \\ &= 6 \int \frac{1}{x-1} dx - 4 \int \frac{1}{x-2} dx + 15 \int \frac{1}{(x-2)^2} dx \\ &= 6 \ln|x-1| - 4 \ln|x-2| - 15 \frac{1}{x-2} + C. \end{aligned} \quad \square$$

Beispiel 8.19 (Musterbeispiel). $\int \frac{x^6 - 2x^5 + x^4 + 4x + 1}{x^4 - 2x^3 + 2x - 1} dx = ?$

- (i) Wir zerlegen durch Polynomdivision den Integrand in ein Polynom und eine echt gebrochenrationale Funktion:

$$\begin{aligned} \frac{x^6 - 2x^5 + x^4 + 4x + 1}{x^4 - 2x^3 + 2x - 1} &= (x^6 - 2x^5 + x^4 + 4x + 1) : (x^4 - 2x^3 + 2x - 1) \\ &= x^2 + 1 + \frac{x^2 + 2x + 2}{x^4 - 2x^3 + 2x - 1} \end{aligned}$$

- (ii) Da das Nennerpolynom vom Grad 4 ist, können die Nullstellen des Nennerpolynoms nicht durch eine Formel berechnet werden. Man muss daher eine Nullstelle x_0 erraten und dann den Grad entweder wieder durch Polynomdivision des Nenners durch $(x - x_0)$ oder durch Anwenden des Horner-Schemas reduzieren. Für das Nennerpolynom $x^4 - 2x^3 + 2x - 1$ erhält man:

$$\begin{aligned} x_1 &= 1 && \text{ist dreifache Nullstelle} \\ x_2 &= -1 && \text{ist einfache Nullstelle} \end{aligned}$$

Damit erhält man die Linearfaktorzerlegung

$$x^4 - 2x^3 + 2x - 1 = (x-1)^3 (x+1)$$

- (iii) Jeder Nullstelle werden gemäß ihrer Vielfachheit die Partialbrüche zugeordnet:

$$\begin{aligned} x_1 = 1 : & \quad \frac{A_1}{x-1} + \frac{A_2}{(x-1)^2} + \frac{A_3}{(x-1)^3} ; \\ x_2 = -1 : & \quad \frac{B}{x+1} . \end{aligned}$$

Darstellung von $\frac{p(x)}{q(x)}$ durch Partialbrüche:

$$\begin{aligned} \frac{x^2 + 2x + 2}{x^4 - 2x^3 + 2x - 1} &= \frac{A_1}{(x-1)} + \frac{A_2}{(x-1)^2} + \frac{A_3}{(x-1)^3} + \frac{B}{x+1} \\ &= \frac{A_1 (x-1)^2 (x+1) + A_2 (x-1) (x+1) + A_3 (x+1) + B (x-1)^3}{HN} \end{aligned}$$

- (iv) Die Bestimmung der Koeffizienten erfolgt, indem mit dem Hauptnenner (HN) multipliziert und spezielle x -Werte eingesetzt werden:

$$\begin{aligned} x = 1 : \quad 5 &= A_3 \cdot 2 & \Rightarrow \quad A_3 &= \frac{5}{2} \\ x = -1 : \quad 1 &= B \cdot (-2)^3 & \Rightarrow \quad B &= -\frac{1}{8} \\ x = 0 : \quad 2 &= A_1 - A_2 + \frac{5}{2} + \left(-\frac{1}{8}\right) (-1) & \Rightarrow \quad A_1 - A_2 &= -\frac{5}{8} \\ x = 2 : \quad 10 &= 3A_1 + 3A_2 + \frac{5}{2} \cdot 3 + \left(-\frac{1}{8}\right) & \Rightarrow \quad A_1 + A_2 &= \frac{7}{8} \end{aligned}$$

Durch Addition bzw. Subtraktion der beiden letzten Gleichungen folgt

$$A_1 = \frac{1}{8}, \quad A_2 = \frac{3}{4}.$$

- (v) Zum Schluss wird die Integration des Polynoms und der Partialbrüche durchgeführt:

$$\begin{aligned} \int \frac{x^6 - 2x^5 + x^4 + 4x + 1}{x^4 - 2x^3 + 2x - 1} dx &= \\ &= \int (x^2 + 1) dx + \frac{1}{8} \int \frac{1}{x-1} dx + \frac{3}{4} \int \frac{1}{(x-1)^2} dx \\ &\quad + \frac{5}{2} \int \frac{1}{(x-1)^3} dx - \frac{1}{8} \int \frac{1}{x+1} dx = \\ &= \frac{1}{3} x^3 + x + \frac{1}{8} \ln|x-1| - \frac{3}{4} \frac{1}{(x-1)} - \frac{5}{4} \frac{1}{(x-1)^2} - \frac{1}{8} \ln|x+1| + C. \square \end{aligned}$$

Zusammenfassung: Jede rationale Funktion $f(x) = \frac{Z(x)}{N(x)}$ ($Z(x)$, $N(x)$ Polynome) lässt sich mit Hilfe algebraischer Methoden integrieren, wenn $N(x)$ im Reellen in Linearfaktoren zerfällt:

- (1) Zerlegung der Funktion $f(x) = r(x) + \frac{p(x)}{q(x)}$ in ein Polynom

$r(x)$ und eine echt gebrochenrationale Funktion $\frac{p(x)}{q(x)}$. (Man beachte: $N(x) = q(x)$).

- (2) Bestimmung der reellen Nullstellen von $q(x)$ und deren Vielfachheit.

- (3) Jeder Nullstelle x_0 von $q(x)$ werden entsprechend ihrer Vielfachheit Partialbrüche zugeordnet

$$x_0 \text{ einfache Nullstelle} \rightarrow \frac{A}{x - x_0}$$

$$x_0 \text{ zweifache Nullstelle} \rightarrow \frac{A_1}{x - x_0} + \frac{A_2}{(x - x_0)^2}$$

$\vdots$

$$x_0 \text{ } k\text{-fache Nullstelle} \rightarrow \frac{A_1}{x - x_0} + \frac{A_2}{(x - x_0)^2} + \dots + \frac{A_k}{(x - x_0)^k}.$$

Die echt gebrochenrationale Funktion $\frac{p(x)}{q(x)}$ ist dann als Summe aller Partialbrüche darstellbar.

- (4) Bestimmung der in den Partialbrüchen auftretenden Konstanten (entweder durch Koeffizientenvergleich und Lösen des zugehörigen linearen Gleichungssystems oder durch Einsetzen spezieller x -Werte).

- (5) Integration von $r(x)$ und sämtlicher Partialbrüche mit

$$\int \frac{1}{x - x_0} dx = \ln|x - x_0| + C;$$

$$\int \frac{1}{(x - x_0)^k} dx = -\frac{1}{k-1} \frac{1}{(x - x_0)^{k-1}} + C.$$

Bemerkung: Nach dem Zusatz zum Fundamentalsatz der Algebra (vgl. Kap. 5.2.7) hat ein reelles Polynom genau n Nullstellen, die entweder reell oder paarweise komplex konjugiert auftreten. Für **komplexe Nullstellen** gelten die Partialbrüche:

- (1) Hat das Polynom $q(x)$ in $x_0 = a + ib$ eine komplexe Nullstelle, so ist auch $\bar{x}_0 = a - ib$ eine Nullstelle und das Produkt

$$(x - x_0)(x - \bar{x}_0) = (x - a)^2 + b^2$$

reell unzerlegbar. Alle derartigen einfachen komplexen Nullstellen sind im Ansatz neben den übrigen Nullstellen zu berücksichtigen durch

$$f(x) = \dots + \frac{Cx + D}{(x - a)^2 + b^2} + \dots$$

- (2) Liegen die komplexen Nullstellen k -fach vor, so wird der Ansatz modifiziert

$$f(x) = \dots + \frac{C_1 x + D_1}{(x - a)^2 + b^2} + \frac{C_2 x + D_2}{[(x - a)^2 + b^2]^2} + \dots + \frac{C_k x + D_k}{[(x - a)^2 + b^2]^k} + \dots$$



Die Durchführung dieser komplizierteren Partialbruchzerlegung ist in der Regel sehr aufwändig. Wir verweisen auf die Beispiele mit MAPLE.

Beispiele 8.20 (Mit MAPLE-Worksheet):

- ① $\int \frac{2x^3 + x^2 + 2x + 2}{x^4 + 2x^2 + 1} dx = ?$: Die Nullstellen des Nenners $q(x)$ sind $-i, -i, i, i$. D.h. $x = i$ und $x = -i$ sind jeweils doppelte Nullstellen. Die Zerlegung des Integranden in Partialbrüche lautet

$$\frac{2x}{x^2 + 1} + \frac{1}{x^2 + 1} + \frac{1}{(x^2 + 1)^2}$$

und die anschließende Integration liefert

$$\ln(x^2 + 1) + \frac{3}{2} \arctan(x) + \frac{x}{2(x^2 + 1)} + C$$

- ② $\int \frac{x^3 + 2x^2 - 1}{x^4 - 2x^3 + 2x^2 - 2x + 1} dx = ?$: Die Nullstellen des Nenners sind 1 (doppelt) und $\pm i$. Eine Partialbruchzerlegung liefert

$$\frac{1 - 3x}{2(x^2 + 1)} + \frac{5}{2(x - 1)} + \frac{1}{(x - 1)^2}$$

mit der anschließenden Integration

$$-\frac{3}{4} \ln(x^2 + 1) + \frac{1}{2} \arctan(x) + \frac{5}{2} \ln(x - 1) - \frac{1}{x - 1} + C. \quad \square$$

8.5 Uneigentliche Integrale

Bisher wurde von den bestimmten Integralen stets vorausgesetzt, dass die Integrationsgrenzen endlich sind. Es tritt in den Anwendungen aber auch der Fall ein, dass sie nicht beschränkt sind und die Integrale dennoch existieren. Man betrachtet diese als Grenzfall der eigentlichen Integrale.

Definition: Integrale, bei denen als Integrationsgrenzen $\pm \infty$ auftritt, bezeichnet man als **uneigentliche Integrale**:

$$\int_a^\infty f(x) dx, \quad \int_{-\infty}^b f(x) dx, \quad \int_{-\infty}^\infty f(x) dx.$$

Anwendungsbeispiel 8.21 (Austrittsarbeit).

Im Gravitationsfeld der Erde soll eine Masse m aus der Entfernung r_0 ins Unendliche ($r = \infty$) gebracht werden. Welche Arbeit W_∞ ist dazu aufzuwenden und welche Geschwindigkeit (= *Fluchtgeschwindigkeit*) benötigt die Masse dazu? Die Arbeit W_R , die aufgebracht werden muss, um die Masse von $r = r_0$ nach $r = R$ zu bringen, ist über das Gravitationsgesetz

$$F(r) = f \frac{mM}{r^2}$$

gegeben durch

$$\begin{aligned} W_R = \int_{r_0}^R F(r) dr &= \int_{r_0}^R f \frac{mM}{r^2} dr = f m M \int_{r_0}^R \frac{1}{r^2} dr = f m M \left[-\frac{1}{r} \right]_{r_0}^R \\ &= f m M \left(\frac{1}{r_0} - \frac{1}{R} \right). \end{aligned}$$

Dabei ist f die Gravitationskonstante und M die Masse der Erde. Für $R \rightarrow \infty$ gilt dann

$$W_\infty = \lim_{R \rightarrow \infty} W_R = \lim_{R \rightarrow \infty} f m M \left[\frac{1}{r_0} - \frac{1}{R} \right] = \frac{f m M}{r_0}.$$

Dies ist dann gleich der kinetischen Energie $\frac{1}{2} m v^2$, welche die Masse zu Beginn besitzen muss; also ist die Fluchtgeschwindigkeit

$$v_\infty = \sqrt{\frac{2 f M}{r_0}}.$$

□

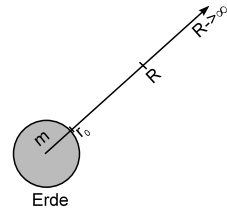


Abb. 8.2.
Austrittsarbeit

Das Vorgehen, welches in obigem Beispiel gewählt wurde, nämlich zunächst von r_0 bis R zu integrieren und dann $R \rightarrow \infty$ gehen zu lassen, ist die Berechnungsmethode von uneigentlichen Integralen:

Berechnung von uneigentlichen Integralen der Form $\int_a^\infty f(x) dx$:

(1) Bestimmung der Integralfunktion $I(\lambda)$ als Funktion der oberen Grenze $I(\lambda) = \int_a^\lambda f(x) dx$.

(2) Bestimmung des Grenzwertes der Integralfunktion für $\lambda \rightarrow \infty$:

$$\int_a^\infty f(x) dx = \lim_{\lambda \rightarrow \infty} I(\lambda) = \lim_{\lambda \rightarrow \infty} \int_a^\lambda f(x) dx.$$

Beispiele 8.22:

① $\int_1^\infty \frac{1}{x^3} dx = ?$

$$\int_1^\infty \frac{1}{x^3} dx = \lim_{\lambda \rightarrow \infty} \int_1^\lambda \frac{1}{x^3} dx = \lim_{\lambda \rightarrow \infty} \left[-\frac{1}{2} \frac{1}{x^2} \right]_1^\lambda = \lim_{\lambda \rightarrow \infty} \frac{1}{2} \left(-\frac{1}{\lambda^2} + 1 \right) = \frac{1}{2}.$$

② $\int_1^\infty \frac{1}{r} dr = ?$ Dieses uneigentliche Integral existiert **nicht**:

$$\int_1^\infty \frac{1}{r} dr = \lim_{\lambda \rightarrow \infty} \int_1^\lambda \frac{1}{r} dr = \lim_{\lambda \rightarrow \infty} \ln r \Big|_1^\lambda = \lim_{\lambda \rightarrow \infty} \ln(\lambda) = \infty. \quad \square$$

Anwendungsbeispiel 8.23 (RL-Kreis).

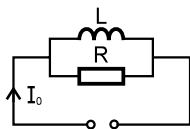


Abb. 8.3. RL-Kreis

Eine Spule (Induktivität L) und ein Ohmscher Widerstand R sind parallel geschaltet. Es fließt ein konstanter Strom I_0 . Zum Zeitpunkt $t_0 = 0$ wird die Stromquelle abgeschaltet und der Strom nimmt gemäß $I(t) = I_0 e^{-\frac{R}{L}t}$ ab. Die Energie, die in Form eines Magnetfeldes vorliegt, ist gegeben durch

$$\begin{aligned} E &= \int_0^\infty R I^2(t) dt = \lim_{T \rightarrow \infty} \int_0^T R I_0^2 e^{-2\frac{R}{L}t} dt \\ &= R I_0^2 \lim_{T \rightarrow \infty} \left[-\frac{L}{2R} e^{-2\frac{R}{L}t} \right]_0^T = \frac{1}{2} L I_0^2. \end{aligned} \quad \square$$

Bemerkungen:

- (1) Die für die Anwendungen wichtigsten *Transformationen*, die *Fourier-Transformation* ($\rightarrow$ Kapitel 15) und die *Laplace-Transformation* ($\rightarrow$ Kapitel 13), sind durch uneigentliche Integrale definiert.
- (2) Integrale mit *unbeschränktem* Integranden bezeichnet man ebenfalls als **uneigentliche Integrale**:

$$\int_0^1 \frac{1}{\sqrt{1-t}} dt$$

ist ein solches uneigentliches Integral, da der Integrand $\frac{1}{\sqrt{1-t}}$ nur für $0 \leq t < 1$ definiert ist. Dennoch hat das Integral einen endlichen Wert, da

$$\int_0^T \frac{1}{\sqrt{1-t}} dt = -2\sqrt{1-t} \Big|_0^T = 2 - 2\sqrt{1-T} \xrightarrow{T \rightarrow 1} 2.$$

- (3) Man unterscheidet also drei Formen von uneigentlichen Integralen:
1. Das Integrationsintervall ist unbeschränkt.
 2. Der Integrand ist unbeschränkt.
 3. Sowohl das Integrationsintervall als auch der Integrand sind unbeschränkt.

8.6 Anwendungen der Integralrechnung

8.6

In diesem Abschnitt werden einige wichtige Anwendungen der Integralrechnung angegeben. Weitere [Anwendungen der Integralrechnung](#) wie die Mittelungseigenschaft, die Bogenlänge und das Krümmungsverhalten sowie die Berechnung von Rotationskörpern zusammen mit der Visualisierung in MAPLE befinden sich auf der CD-Rom.

► 8.6.1 Flächenberechnungen

Aufgrund seiner Definition dient das Integral zunächst zur Berechnung von Flächeninhalten. Ein Flächenstück werde von $x = a$, $x = b$, der x -Achse und der Funktion $f(x)$ begrenzt. Dann ist der Inhalt der Fläche gegeben durch

$$\int_a^b f(x) dx.$$

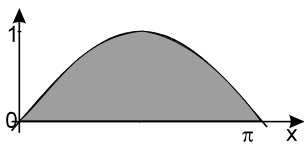


Abb. 8.4. Fläche unter Funktion

Anwendungsbeispiel 8.24: Gesucht ist das Flächenstück unter einer Sinushalbwellen (siehe Abb. 8.4):

$$A = \int_0^{\pi} \sin x \, dx = -\cos x \Big|_0^{\pi} = -(-1 - 1) = 2.$$

Die Kurve schließt mit der x -Achse eine Fläche mit dem Wert 2 ein. $\square$

Der **Flächeninhalt zwischen zwei Kurven** $y = f(x)$ und $y = g(x)$ berechnet sich aus der Differenz der Einzelintegralen:

$$A = \int_a^b (f(x) - g(x)) \, dx = \int_a^b f(x) \, dx - \int_a^b g(x) \, dx.$$

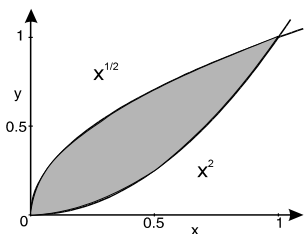


Abb. 8.5. Fläche zw. Funktionen

Anwendungsbeispiel 8.25 (Fläche zwischen zwei Kurven). Gesucht ist die grau unterlegte Fläche zwischen der Funktion $y = \sqrt{x}$ und $y = x^2$ (siehe Abb. 8.5).

Um die grau schraffierte Fläche berechnen zu können, müssen zuerst die Schnittpunkte der Kurven bestimmt werden, da diese die Integrationsgrenzen liefern:

$$\sqrt{x} = x^2 \quad \hookrightarrow x = 0 \text{ und } x = 1.$$

Die durch die Kurven eingeschlossene Fläche wird nun berechnet mit dem bestimmten Integral über die Differenz der Funktionen:

$$A = \int_0^1 (\sqrt{x} - x^2) \, dx = \int_0^1 x^{\frac{1}{2}} \, dx - \int_0^1 x^2 \, dx = \frac{2}{3} x^{\frac{3}{2}} \Big|_0^1 - \frac{1}{3} x^3 \Big|_0^1 = \frac{1}{3}. \quad \square$$

➤ 8.6.2 Kinematik

Für die Bewegung eines Massenpunktes gilt

$$v(t) = \dot{s}(t) = \frac{d}{dt} s(t) \quad (\text{Geschwindigkeit}),$$

$$a(t) = \frac{d}{dt} v(t) = \dot{v}(t) = \ddot{s}(t) \quad (\text{Beschleunigung}).$$

Ist die Beschleunigung als Funktion der Zeit bekannt (z.B. durch ein Kraftgesetz), so folgt durch Integration die Geschwindigkeit $v(t)$ und durch nochmalige Integration das Weg-Zeit-Gesetz $s(t)$:

$$\begin{aligned} v(t) &= \int a(t) dt, \\ s(t) &= \int v(t) dt. \end{aligned}$$

Anwendungsbeispiel 8.26 (Freier Fall ohne Luftreibung). Für den freien Fall ohne Luftwiderstand ist die Beschleunigungskraft

$$m \cdot a = F_G = m g \quad \Rightarrow \quad a(t) = g = \text{const.}$$

Damit folgt für die Geschwindigkeit

$$v(t) = \int a(t) dt = g t + C_1.$$

Die Integrationskonstante bestimmt sich aus der Anfangsgeschwindigkeit

$$v(0) = v_0 \quad \hookrightarrow \quad C_1 = v_0 \quad \Rightarrow \quad v(t) = g t + v_0.$$

Das Weg-Zeit-Gesetz folgt durch nochmalige Integration

$$s(t) = \int v(t) dt = \int (g t + v_0) dt = \frac{g}{2} t^2 + v_0 t + C_2.$$

Die Integrationskonstante bestimmt sich aus der Anfangsposition

$$s(0) = s_0 \quad \hookrightarrow \quad C_2 = s_0 \quad \Rightarrow \quad s(t) = \frac{1}{2} g t^2 + v_0 t + s_0. \quad \square$$

Anwendungsbeispiel 8.27 (Bewegungsgleichung einer Rakete). Eine Rakete steige senkrecht in die Luft auf und besitze eine konstante Schubkraft F_0 . Die Massenabnahme der Rakete aufgrund der Verbrennung des Brennstoffes sei linear, d.h.

$$m(t) = m_0 - q t = m_0 (1 - \alpha t) \quad \text{mit} \quad \alpha = \frac{q}{m_0},$$

wenn m_0 die Startmasse und q der Brennstoffverbrauch. Unter der Voraussetzung einer konstanten Erdbeschleunigung g und ohne Luftwiderstand ist die Beschleunigungskraft bzw. Beschleunigung

$$m a = F_0 - m g \quad \hookrightarrow \quad a = \frac{F_0}{m_0 (1 - \alpha t)} - g.$$

Die Geschwindigkeit ist

$$v(t) = \int a(t) dt = \frac{F_0}{m_0} \int \frac{dt}{1 - \alpha t} - g \int dt = -\frac{F_0}{m_0} \frac{1}{\alpha} \ln(1 - \alpha t) - g t + C.$$

Man beachte, dass das erste Integral mit der Substitution $y = 1 - \alpha t$ berechnet wird. Mit der Anfangsgeschwindigkeit $v(0) = 0$ wird $C = 0$.

Das Weg-Zeit-Gesetz erhält man durch nochmalige Integration

$$s(t) = \int v(t) dt = -\frac{F_0}{m_0 \alpha} \int \ln(1 - \alpha t) dt - g \int t dt.$$

Mit der Substitution $y = 1 - \alpha t$ und dem Ergebnis aus Beispiel 8.13 ③ $\int \ln x = x \cdot (\ln x - 1) + C$ ist

$$s(t) = \frac{F_0}{m_0 \alpha} \frac{1}{\alpha} [(1 - \alpha t) \ln(1 - \alpha t) - (1 - \alpha t)] - \frac{1}{2} g t^2 + C.$$

Mit der Anfangsbedingung $s(0) = 0$ folgt $C = \frac{F_0}{m_0 \alpha^2}$ und damit

$$s(t) = \frac{F_0}{m_0 \alpha^2} [(1 - \alpha t) \ln(1 - \alpha t) + \alpha t] - \frac{1}{2} g t^2.$$

□

8.6.3 Elektrodynamik

Eine Punktladung Q induziert ein radiales **elektrisches Feld** gemäß der Formel

$$E(r) = \frac{1}{4\pi \varepsilon_0} \frac{Q}{r^2}$$

mit der Dielektrizitätskonstanten $\varepsilon_0 = 8.8542 \cdot 10^{-12} \frac{F}{m}$. Die *Spannung* U_{12} zwischen zwei Punkten P_1 und P_2 mit Abständen r_1 und r_2 von Q ist gegeben durch

$$U_{12} = \int_{r_1}^{r_2} E(r) dr = \frac{Q}{4\pi \varepsilon_0} \int_{r_1}^{r_2} \frac{1}{r^2} dr = -\frac{Q}{4\pi \varepsilon_0} \frac{1}{r} \Big|_{r_1}^{r_2} = \frac{Q}{4\pi \varepsilon_0} \left(\frac{1}{r_1} - \frac{1}{r_2} \right).$$

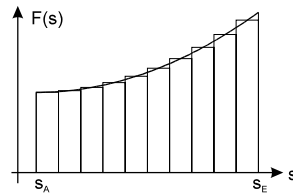
8.6.4 Energieintegrale

Wirkt auf einen Massenpunkt m eine **ortsunabhängige** Kraft F in Wegrichtung, so ist die verrichtete *Arbeit* definitionsgemäß

$$W := F \cdot \Delta s,$$

wenn $\Delta s = s_E - s_A$ die Strecke, um die der Massenpunkt verschoben wird.

Falls die Kraft jedoch **ortsabhängig** ist $F = F(s)$ (siehe Abb. 8.6), dann unterteilt man den Weg in kleine Intervalle Δs und nimmt für jedes Teilintervall eine konstante Kraft an. Im Intervall Δs_i wird näherungsweise die Arbeit



$$W_i = F(s_i) \cdot \Delta s_i \quad \text{Abb. 8.6. Ortsabhängige Kraft}$$

geleistet. Die Gesamtarbeit W ist die Summe aller Einzelbeiträge

$$W \approx \sum_{i=1}^n W_i = \sum_{i=1}^n F(s_i) \cdot \Delta s_i.$$

Den exakten Wert der geleisteten Arbeit erhält man, indem man zu einer beliebig feinen Unterteilung übergeht ($\Delta s_i \rightarrow 0$ bzw. $n \rightarrow \infty$):

$$W = \lim_{n \rightarrow \infty} \sum_{i=1}^n W_i = \lim_{n \rightarrow \infty} \sum_{i=1}^n F(s_i) \Delta s_i = \int_{s_A}^{s_E} F(s) ds.$$

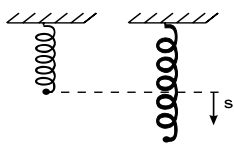
Zusammenfassung: Die **Arbeit** einer ortsabhängigen Kraft ist gegeben durch

$$W = \int_{s_A}^{s_E} F(s) ds,$$

wenn $F(s)$ die Kraftkomponente in Wegrichtung ist. Andernfalls muss $F(s) ds$ durch das Skalarprodukt

$$\vec{F} d\vec{s} = |\vec{F}| \cos \varphi ds$$

ersetzt werden, wenn φ der Winkel zwischen $\vec{F}$ und der Wegrichtung $\vec{s}$ ist.

Anwendungsbeispiel 8.28 (Spannarbeit einer elastischen Feder).

Wird eine elastische Feder aus ihrer Ruhelage um s ausgedehnt, so wirkt eine Rückstellkraft proportional zur Auslenkung:

$$F(s) = -D \cdot s \quad (\text{Hooksches Gesetz}).$$

Abb. 8.7. Feder-Masse-System

Um eine Feder aus der Ruhelage um die Strecke s_0 auszu lenken, wird die folgende Arbeit verrichtet:

$$W = \int_0^{s_0} F(s) \, ds = D \int_0^{s_0} s \, ds = \frac{1}{2} D s_0^2. \quad \square$$

Anwendungsbeispiel 8.29 (Arbeit im elektrostatischen Feld). Zwei Punktladungen q_1 und q_2 üben aufeinander eine Kraft aus, die umgekehrt proportional zum Quadrat ihrer Entfernung ist

$$F = \frac{1}{4\pi \varepsilon_0} \frac{q_1 \cdot q_2}{r^2}.$$

Ist q_1 im Ursprung und wird q_2 von r_1 nach r_2 verschoben, ist die folgende Arbeit zu verrichten:

$$W = \int_{r_1}^{r_2} F(r) \, dr = \frac{q_1 q_2}{4\pi \varepsilon_0} \int_{r_1}^{r_2} \frac{1}{r^2} \, dr = \frac{q_1 q_2}{4\pi \varepsilon_0} \left(\frac{1}{r_1} - \frac{1}{r_2} \right). \quad \square$$

8.6.5 Lineare und quadratische Mittelwerte

Problemstellung: Ein Zweiweggleichrichter erzeugt aus einem Sinuswechselstrom $i(t) = i_0 \sin(\omega t)$ mit $\omega = \frac{2\pi}{T}$ den in Abb. 8.8 gezeigten Verlauf. Gesucht ist der *lineare Mittelwert*.

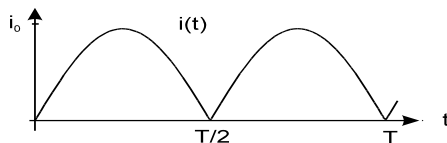


Abb. 8.8. Zweiweggleichrichter

Definition: Unter dem **linearen Mittelwert** einer Funktion $y = f(x)$ im Intervall $[a, b]$ versteht man die Größe

$$\bar{y} = \frac{1}{b-a} \int_a^b f(x) \, dx.$$

Bemerkung: Nach dem Mittelwertsatz der Integralrechnung gibt es dann immer einen Zwischenwert $\xi \in (a, b)$, so dass $\bar{y} = f(\xi)$. Geometrisch bedeutet der lineare Mittelwert, dass man die Fläche unter der Kurve durch ein flächengleiches Rechteck ersetzt.

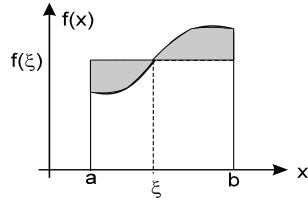


Abb. 8.9. Mittelwert $f(\xi)$

➤ Mittelwert eines Zweiweggleichrichters

Anwendungsbeispiel 8.30 (Linearer Mittelwert). Der *lineare* Mittelwert eines Zweiweggleichrichters berechnet man über das Integral

$$\begin{aligned}\bar{i} &= \frac{1}{T/2} \int_0^{T/2} i_0 \sin(\omega t) dt = -\frac{2}{T} i_0 \frac{1}{\omega} \cos(\omega t) \Big|_0^{T/2} \\ &= -\frac{2}{T} i_0 \frac{T}{2\pi} \left[\cos\left(\frac{2\pi}{T} \cdot \frac{T}{2}\right) - 1 \right] = \frac{2}{\pi} i_0.\end{aligned}$$

Der lineare Mittelwert eines sinusförmigen Wechselstroms ist 0. □

In elektrotechnischen Anwendungen verwendet man aber den sog. *quadratischen Mittelwert*.

Definition: Unter dem **quadratischen Mittelwert** einer Funktion $y = f(x)$ im Intervall $[a, b]$ versteht man die Größe

$$\bar{y}_q = \left(\frac{1}{b-a} \int_a^b f^2(x) dx \right)^{\frac{1}{2}}.$$

Anwendungsbeispiel 8.31 (Effektivwert des Wechselstroms). Der **Effektivwert** I_{eff} eines Wechselstroms ist der *quadratische* Mittelwert während einer Periode T :

$$\begin{aligned}I_{eff} &= \left(\frac{1}{T} \int_0^T i^2(t) dt \right)^{\frac{1}{2}} = \left(\frac{1}{T} \int_0^T i_0^2 \sin^2(\omega t) dt \right)^{\frac{1}{2}} \\ &= \left(\frac{i_0^2}{T} \int_0^T \sin^2(\omega t) dt \right)^{\frac{1}{2}} = \left(\frac{i_0^2}{T} \left[\frac{1}{2} t - \frac{1}{4\omega} \sin(2\omega t) \right]_0^T \right)^{\frac{1}{2}} \\ &= \left(\frac{i_0^2}{T} \left[\frac{1}{2} T - \frac{1}{4\omega} \sin(2\omega T) \right] \right)^{\frac{1}{2}} = \frac{i_0}{\sqrt{2}}.\end{aligned}$$

Analog ergibt sich der Effektivwert einer Wechselspannung durch

$$U_{eff} = \left(\frac{1}{T} \int_0^T u^2(t) dt \right)^{\frac{1}{2}} = \frac{U_0}{\sqrt{2}}. \quad \square$$

8.6.6 Schwerpunkt einer ebenen Fläche

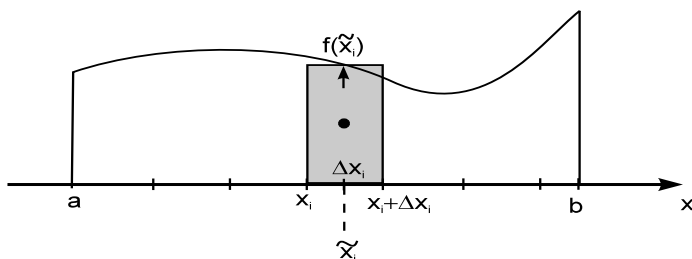


Abb. 8.10. Schwerpunktsberechnung einer ebenen Fläche

Sind $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ die Koordinaten von n Massenpunkten mit den Massen $m_1, m_2, \dots, m_n$, dann ergeben sich die Koordinaten des *Schwerpunktes* durch

$$x_s = \frac{\sum_{i=1}^n m_i x_i}{\sum_{i=1}^n m_i} \quad \text{und} \quad y_s = \frac{\sum_{i=1}^n m_i y_i}{\sum_{i=1}^n m_i}.$$

Zur Berechnung des Schwerpunktes einer ebenen Fläche gehen wir davon aus, dass die Fläche durch eine Kurve $y = f(x)$ und der x -Achse zwischen $x = a$ und $x = b$ begrenzt sei (Abb. 8.10). Wir belegen die Fläche homogen mit der Massendichte 1. Zur Berechnung des Schwerpunktes unterteilen wir das Intervall $[a, b]$ in n Teilintervalle, $a = x_0 < x_1 < \dots < x_n = b$, wählen für jedes Teilintervall Δx_i ein $\tilde{x}_i \in [x_i, x_i + \Delta x_i]$ und bestimmen $f(\tilde{x}_i)$. Der Schwerpunkt jedes der Rechtecke $A_i = \Delta x_i f(\tilde{x}_i)$ ist

$$x_{s,i} = x_i + \frac{1}{2} \Delta x_i \quad , \quad y_{s,i} = \frac{1}{2} f(\tilde{x}_i)$$

mit der Masse $m_i = \frac{A_i}{A} = A_i = \Delta x_i f(\tilde{x}_i)$ ($i = 1 \dots n$). Die Koordinaten des Schwerpunktes, der so gewonnenen n Massen mit Koordinaten $(x_{s,1}, y_{s,1}), \dots,$

$(x_{s,n}, y_{s,n})$, sind nach obigen Formeln gegeben durch

$$x_s = \frac{\sum_{i=1}^n m_i x_{s,i}}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n \Delta x_i f(\tilde{x}_i) \cdot (x_i + \frac{1}{2} \Delta x_i)}{\sum_{i=1}^n \Delta x_i f(\tilde{x}_i)}$$

$$y_s = \frac{\sum_{i=1}^n m_i y_{s,i}}{\sum_{i=1}^n m_i} = \frac{\sum_{i=1}^n \Delta x_i f(\tilde{x}_i) \cdot \frac{1}{2} f(\tilde{x}_i)}{\sum_{i=1}^n \Delta x_i f(\tilde{x}_i)}$$

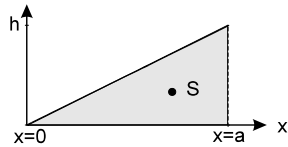
Für den Grenzübergang $n \rightarrow \infty$ gehen die Zwischensummen in die zugehörigen Integrale über.

Satz: Die Koordinaten des **Schwerpunktes** $S = (x_s, y_s)$ der Fläche unter dem Graphen der Funktion $y = f(x)$ zwischen $x = a$ und $x = b$ sind

$$x_s = \frac{\int_a^b x f(x) dx}{\int_a^b f(x) dx} \quad \text{und} \quad y_s = \frac{\frac{1}{2} \int_a^b f^2(x) dx}{\int_a^b f(x) dx}.$$

Beispiele 8.32:

① Die Schwerpunktskoordinaten des nebenstehenden Dreiecks unter der Geraden $y = f(x) = \frac{h}{a} x$ sind gegeben durch



$$x_s = \frac{1}{A} \int_0^a x \frac{h}{a} x dx = \frac{1}{A} \frac{h}{a} \left[\frac{x^3}{3} \right]_0^a = \frac{1}{A} \frac{h}{3} a^2,$$

$$y_s = \frac{1}{2A} \int_0^a \left(\frac{h}{a} x \right)^2 dx = \frac{1}{2A} \frac{h^2}{a^2} \left[\frac{x^3}{3} \right]_0^a = \frac{1}{2A} \frac{h^2}{3} a.$$

Mit

$$A = \int_0^a f(x) dx = \int_0^a \frac{h}{a} x dx = \frac{h}{a} \frac{x^2}{2} \Big|_0^a = \frac{1}{2} h a$$

folgt

$$x_s = \frac{2}{3} a$$

und

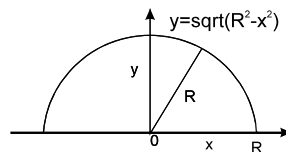
$$y_s = \frac{1}{3} h.$$

② Wir berechnen die Schwerpunktskoordinaten des nebenstehenden Halbkreises mit Radius R : Aus Symmetriegründen liegt der Schwerpunkt auf der y -Achse, so dass $x_s = 0$. Für die y -Koordinate des Schwerpunktes gilt

$$\begin{aligned} y_s &= \frac{1}{2A} \int_{-R}^R y^2 dx \\ &= \frac{1}{2A} \int_{-R}^R (R^2 - x^2) dx \\ &= \frac{1}{2A} \left[R^2 x - \frac{1}{3} x^3 \right]_{-R}^R \\ &= \frac{1}{2A} \frac{4}{3} R^3. \end{aligned}$$

Mit $A = \frac{\pi}{2} R^2$ folgt insgesamt

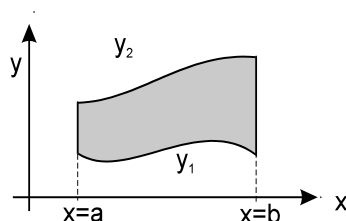
$$y_s = \frac{4}{3\pi} R.$$



③ Die Koordinaten des Schwerpunktes der Fläche A , die durch zwei Funktionen $y_2 = f(x)$ und $y_1 = g(x)$ mit $f \geq g$ sowie den Geraden $x = a$ und $x = b$ begrenzt ist, sind gegeben durch die Differenz der Einzelschwerpunkte:

$$x_s = \frac{1}{A} \int_a^b x [f(x) - g(x)] dx$$

$$y_s = \frac{1}{2A} \int_a^b [f^2(x) - g^2(x)] dx$$



mit

$$A = \int_a^b [f(x) - g(x)] dx.$$

MAPLE-Worksheets zu Kapitel 8



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 8 mit MAPLE zur Verfügung.

- Visualisierung des Integralbegriffs
- Integration mit MAPLE
- Integrationsmethoden mit MAPLE
- Uneigentliche Integrale mit MAPLE
- Anwendungen mit MAPLE: Bogenlänge und Krümmung
- Anwendungen mit MAPLE: Mittelungseigenschaft
- Anwendungen mit MAPLE: Rotationskörper
- MAPLE-Lösungen zu den Aufgaben

8.7 Aufgaben zur Integralrechnung

8.1 Stellen Sie mit MAPLE zu der Funktion $f(x) = \sqrt{x}$ im Bereich $x \in [0, 2]$ die Rechtssumme graphisch dar und berechne Sie diese für $n = 10, 50, 100$.

8.2 Gesucht sind die folgenden unbestimmten Integrale:

$$\begin{array}{lll} \text{a) } \int x^5 dx & \text{b) } \int \frac{dx}{x^2} & \text{c) } \int \sqrt[3]{z} dz \\ \text{d) } \int \frac{dx}{\sqrt[3]{x^2}} & \text{e) } \int (2x^2 - 5x + 3) dx & \text{f) } \int (1-x) \sqrt{x} dx \end{array}$$

8.3 Bestimmen Sie die folgenden bestimmten Integrale:

$$\begin{array}{lll} \text{a) } \int_0^{\frac{\pi}{2}} \sin x dx & \text{b) } \int_{\frac{1}{2}\pi}^{\frac{3}{2}\pi} (2x + \sin x - \cos x) dx & \text{c) } \int_1^a \frac{1}{x} dx \end{array}$$

8.4 Bestimmen Sie mit partieller Integration die folgenden unbestimmten Integrale:

$$\begin{array}{lll} \text{a) } \int x \cos x dx & \text{b) } \int \sin x \cos x dx & \text{c) } \int x^2 \sin x dx \\ \text{d) } \int x^2 \ln x dx & \text{e) } \int x e^x dx & \text{f) } \int x^2 e^x dx \end{array}$$

8.5 Man bestimme mittels Substitution die folgenden unbestimmten Integrale:

$$\begin{array}{lll} \text{a) } \int \frac{1}{x+2} dx & \text{b) } \int \frac{x}{x^2-1} dx & \text{c) } \int \frac{x^2}{1-2x^3} dx \\ \text{d) } \int (3s+4)^8 ds & \text{e) } \int \sin(\omega t + \varphi) dt & \text{f) } \int \cos(3t) dt \\ \text{g) } \int e^{-x} dx & \text{h) } \int \frac{\sin t}{\cos t} dt & \text{i) } \int \frac{e^x + x e^x}{x e^x} dx \\ \text{j) } \int \sin x \cos x dx & \text{k) } \int \sqrt{4+3x} dx \end{array}$$

8.6 Zeigen Sie die Gültigkeit der folgenden Gleichungen

$$\begin{array}{ll} \text{a) } \int e^{-x} (1-x) dx = x e^{-x} + C & \\ \text{b) } \int \frac{\sqrt{x^2-4}}{x} dx = \sqrt{x^2-4} - 2 \arccos\left(\frac{2}{x}\right) + C & \\ \text{c) } \int \cos(x) e^{\sin(x)} dx = e^{\sin(x)} + C & \\ \text{d) } \int \cos(3x) \cdot \sin(3x) dx = \frac{1}{6} \sin^2(3x) + C & \end{array}$$

8.7 a) Man löse das Integral $\int \frac{2-x}{1+\sqrt{x}} dx$ mit der Substitution $u = 1 + \sqrt{x}$.

b) Man löse das Integral $\int x \sqrt{1-x^2} dx$ mit der Substitution $x = \sin u$.

8.8 Bestimmen Sie mit einer geeigneten Substitution die folgenden Integrale

$$\begin{array}{lll} \text{a) } \int \frac{x^2}{\sqrt{1+x^3}} dx & \text{b) } \int (5x+12)^{\frac{1}{2}} dx & \text{c) } \int \sqrt[3]{1-t} dt \\ \text{d) } \int_0^{\pi} \cos^3 x \cdot \sin x dx & \text{e) } \int \frac{\arctan z}{1+z^2} dz & \text{f) } \int \frac{2x+6}{x^2+6x-12} dx \\ \text{g) } \int \frac{dx}{x \ln x} & \text{h) } \int x \cdot \sin(x^2) dx & \text{i) } \int \frac{3x^2-2}{2x^3-4x+2} dx \\ \text{j) } \int_{-1}^1 \frac{t}{\sqrt{1+t^2}} dt & \text{k) } \int_0^{\frac{\pi}{2}} \sin\left(3t - \frac{\pi}{4}\right) dt & \text{l) } \int_{-1}^1 \frac{5+x}{5-x} dx \\ \text{m) } \int x^2 e^{x^3-2} dx & \text{n) } \int \frac{\tan(z+5)}{\cos^2(z+5)} dz & \text{o) } \int \frac{\sqrt{4-x^2}}{x^2} dx \end{array}$$

8.9 Man löse die folgenden Integrale durch partielle Integration (mit MAPLE)

- a) $\int x \ln x \, dx$ b) $\int x \cos x \, dx$ c) $\int \ln t \, dt$
 d) $\int x \sin(3x) \, dx$ e) $\int \arctan x \, dx$ f) $\int \sin^2(\omega t) \, dt$
 g) $\int e^x \cos x \, dx$ h) $\int x^2 e^{-x} \, dx$

8.10 Lösen Sie die folgenden Integrale durch Partialbruchzerlegung:

- a) $\int \frac{1}{x^2 - a^2} \, dx$ b) $\int \frac{4x^3}{x^3 + 2x^2 - x - 2} \, dx$ c) $\int \frac{3z}{z^3 + 3z^2 - 4} \, dz$
 d) $\int \frac{4x - 2}{x^2 - 2x - 63} \, dx$ e) $\int \frac{2x + 1}{x^3 - 6x^2 + 9x} \, dx$

8.11 Man löse die folgenden Integrale mit MAPLE

- a) $\int \frac{\sqrt{\ln x}}{x} \, dx$ b) $\int \cot x \, dx$ c) $\int x \cosh x \, dx$
 d) $\int \sin x e^{\cos x} \, dx$ e) $\int \frac{x^3}{(x^2 - 1)(x + 1)} \, dx$ f) $\int \frac{x - 4}{x + 1} \, dx$
 g) $\int \frac{(\ln x)^3}{x} \, dx$ h) $\int \frac{12x^2}{2x^3 - 1} \, dx$ i) $\int x \cdot \arctan x \, dx$

8.12 Bestimmen Sie mit MAPLE $\int \ln(x + \sqrt{1 + x^2}) \, dx$, indem Sie zunächst die Substitution $u^2 = 1 + x^2$ durchführen und anschließend partiell integrieren.

8.13 Berechnen Sie mit MAPLE das unbestimmte Integral $\int \frac{1}{\sqrt{1 + \sqrt{x}}} \, dx$ durch die Substitution $u^2 = 1 + \sqrt{x}$.

8.14 Führen Sie mit MAPLE eine Partialbruchzerlegung durch und integrieren Sie anschließend

- a) $\int \frac{x^4 - x^3 - x - 1}{x^3 - x^2} \, dx$ b) $\int \frac{x^4 - x^3 + 3x}{x^2(x + 2)(x - 3)} \, dx$ c) $\int \frac{x^4 - x^3 + 3x^2 - 2x + 1}{x^3 - x^2 - x + 1} \, dx$

8.15 Man berechne für $n, m \in \mathbb{N}$ die folgenden bestimmten Integrale. (Hinweis: Man verwende die Additionstheoreme für Sinus und Kosinus.)

- a) $\int_0^{2\pi} \sin(nx) \, dx$ b) $\int_0^{2\pi} \cos(nx) \, dx$
 c) $\int_0^{2\pi} \cos(nx) \cos(mx) \, dx$ für $m = n$ und für $m \neq n$
 d) $\int_0^{2\pi} \sin(nx) \sin(mx) \, dx$ für $m = n$ und für $m \neq n$
 e) $\int_0^{2\pi} \sin(nx) \cos(mx) \, dx$

8.16 Man bestimme den Wert der uneigentlichen Integrale

- a) $\int_{-\infty}^{\infty} x^{-2} \, dx$ b) $\int_0^1 \frac{1}{\sqrt{1 - x^2}} \, dx$ c) $\int_0^{\infty} (3e^{-2x} - e^{-x}) \, dx$
 d) $\int_0^{\infty} e^{at} e^{-st} \, dt$ e) $\int_0^{\infty} \cos(at) e^{-st} \, dt$ f) $\int_0^{\infty} t^n e^{-st} \, dt$

8.17 Berechnen Sie den Flächeninhalt zwischen der Parabel $f(x) = x^2 - 2x - 1$ und der Geraden $g(x) = 3x - 1$.

Kapitel 9
Funktionenreihen

9

9	Funktionenreihen	341
9.1	Zahlenreihen	344
9.1.1	Beispiele	346
9.1.2	Majorantenkriterium.....	350
9.1.3	Quotientenkriterium	351
9.1.4	Leibnizkriterium.	353
9.2	Potenzreihen	355
9.3	Taylor-Reihen	361
9.4	Anwendungen.....	371
9.5	Komplexwertige Funktionen	376
9.5.1	Komplexe Potenzreihen.....	376
9.5.2	Die Eulersche Formel.....	378
9.5.3	Eigenschaften der komplexen Exponentialfunktion	379
9.5.4	Komplexe Hyperbelfunktionen	381
9.5.5	Differenziation und Integration.....	382
9.6	Aufgaben zu Funktionenreihen	385
 Zusätzliche Abschnitte auf der CD-Rom		
9.7	MAPLE: Zahlen-, Potenz- und Taylor-Reihen	cd

9 Funktionenreihen

Die wichtigsten, in den Anwendungen auftretenden Funktionen lassen sich als *Potenzreihen* der Form $\sum_{n=0}^{\infty} a_n (x - x_0)^n$, den sog. *Taylor-Reihen* darstellen. Diese Entwicklung liefert eine Möglichkeit, um Funktionen wie z.B. e^x , $\sin x$, $\tan x$, $\sqrt{x}$, $\ln x$ oder $\arctan x$ explizit zu berechnen, indem nur die Grundrechenoperationen $+$ $-$ $*$ $/$ angewendet werden. Darüber hinaus ist es für die Anwendungen wichtig, dass für gegebene, gegebenenfalls komplizierte Funktionen Näherungsformeln zur Verfügung stehen.

Anwendungsbeispiel 9.1 (Scheinwerferregulierung).

Bei der Einstellung von Scheinwerfern muss die Höhe des Abblendlichts laut Gesetz über eine Entfernung von 10 m um eine vorgegebene Höhe $H_{\text{opt}} = 0.1\text{ m}$ abnehmen. Aus dieser Vorgabe ergibt sich für die Hell-Dunkel-Grenze eine Zielneigung der Scheinwerfer durch $\beta_{\text{ab}} = \arctan \frac{H_{\text{opt}}}{10} \approx 0.009999 \hat{=} 0.5729^\circ$.

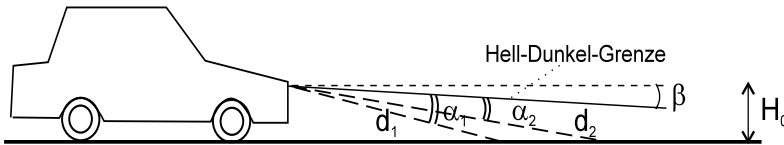


Abb. 9.1. Geometrische Anordnung der beiden Messstrahlen

Da der aktuelle Neigungswinkel β der Scheinwerfer nicht direkt ermittelbar ist, wird er optisch über die Messung zweier Distanzen d_1 und d_2 bestimmt. Die beiden gemessenen Distanzen d_1 und d_2 ergeben sich in Abhängigkeit des aktuellen Scheinwerferwinkels β durch

$$d_1 = \frac{H_0}{\sin(\alpha_1 + \beta)}, \quad d_2 = \frac{H_0}{\sin(\alpha_2 + \beta)},$$

wenn α_1 und α_2 die baubedingt, vorgegebenen Neigungswinkel sind. Geht man zum Quotienten über, ergibt sich die Formel

$$\frac{d_1}{d_2} = \frac{\sin(\alpha_2 + \beta)}{\sin(\alpha_1 + \beta)} = q(\beta),$$

Um vom Quotienten der beiden Distanzwerte auf den aktuellen Neigungswinkel β der Scheinwerfer einfach schließen zu können, wird eine Näherungsformel von $q(\beta)$ gesucht, die sich anschließend nach β auflösen lässt ($\rightarrow$ Taylor-Polynom 2. Ordnung). $\square$

9.1
9.1 Zahlenreihen

Bevor wir allgemein auf Potenz- und Taylor-Reihen zu sprechen kommen, werden zunächst Zahlenreihen und deren Konvergenzkriterien behandelt. Die Konvergenzkriterien benötigen wir dann bei der Diskussion der Konvergenz der Taylor-Reihen.

Nach Kap. 6.1 bezeichnet man eine geordnete Menge reeller Zahlen

$$(a_n)_{n \in \mathbb{N}} = a_1, a_2, a_3, \dots, a_n, \dots$$

als *reelle Zahlenfolge*. Eine Zahlenfolge heißt *konvergent*, wenn eine reelle Zahl $a \in \mathbb{R}$ existiert, so dass es zu jedem $\varepsilon > 0$ ein $n_0 \in \mathbb{N}$ gibt mit

$$|a_n - a| < \varepsilon \quad \text{für } n \geq n_0.$$

Beispiele 9.2:

Folge	allgem. Glied	Konvergenz
$(a_n)_n = 1, 2, 3, 4, \dots$	$a_n = n$	nein
$(a_n)_n = 1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots$	$a_n = \frac{1}{n}$	ja: $a_n \rightarrow 0$
$(a_n)_n = q^0, q^1, q^2, q^3, \dots$	$a_n = q^{n-1}$	$\begin{cases} \text{für } q < 1 : & a_n \rightarrow 0 \\ \text{für } q > 1 : & \text{divergent} \\ \text{für } q = 1 : & a_n \rightarrow 1 \\ \text{für } q = -1 : & \text{divergent} \end{cases}$
$(a_n)_n = 1, \frac{1}{2!}, \frac{1}{3!}, \frac{1}{4!}, \dots$	$a_n = \frac{1}{n!}$	ja: $a_n \rightarrow 0$
$(a_n)_n = -1, \frac{1}{2}, -\frac{1}{3}, \frac{1}{4}, \dots$	$a_n = (-1)^n \frac{1}{n}$	ja: $a_n \rightarrow 0$ □

Übergang zu Reihen. Wir betrachten die Zahlenfolge

$$(a_n)_n = 1, 1, \frac{1}{2!}, \frac{1}{3!}, \frac{1}{4!}, \dots, \frac{1}{(n-1)!}, \dots$$

mit dem allgemeinen Glied $a_n = \frac{1}{(n-1)!}$. Aus den Gliedern dieser Folge bilden wir sog. *Teilsummen* (= *Partialsummen*), indem wir jeweils die ersten Glieder aufsummieren:

$S_1 = 1$
 $S_2 = 1 + 1$
 $S_3 = 1 + 1 + \frac{1}{2!}$
 $S_4 = 1 + 1 + \frac{1}{2!} + \frac{1}{3!}$
 $S_5 = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!}$
 $S_6 = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!} + \frac{1}{5!}$
 $S_7 = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!} + \frac{1}{5!} + \frac{1}{6!}$
 $S_8 = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \frac{1}{4!} + \frac{1}{5!} + \frac{1}{6!} + \frac{1}{7!}$

$= 1$
 $= 2$
 $= 2,5$
 $= 2,66666$
 $= 2,70833$
 $= 2,71666$
 $= 2,71804$
 $= 2,71823$

Wir fassen die Partialsummen zu einer Folge $(S_n)_{n \in \mathbb{N}}$ zusammen. Diese Folge genügt dem Bildungsgesetz

$$S_n = 1 + 1 + \frac{1}{2!} + \frac{1}{3!} + \dots + \frac{1}{(n-1)!} = \sum_{k=1}^n \frac{1}{(k-1)!}.$$

$(S_n)_n$ bezeichnet man als *Reihe*.

Definition: (Reihen). Sei $(a_k)_{k \in \mathbb{N}}$ eine Zahlenfolge. Dann heißt

$$S_n = a_1 + a_2 + a_3 + \dots + a_n = \sum_{k=1}^n a_k$$

eine **Partialsumme** und die Folge der Partialsummen $(S_n)_{n \in \mathbb{N}}$ heißt **unendliche Reihe** (kurz: **Reihe**):

$$(S_n)_{n \in \mathbb{N}} = \left(\sum_{k=1}^n a_k \right)_{n \in \mathbb{N}} = (a_1 + a_2 + \dots + a_n)_{n \in \mathbb{N}}.$$

Bemerkungen:

- (1) Oftmals beginnt die Summation einer Reihe bei $k = 0$.
- (2) Der Name des Summationsindex kann beliebig gewählt werden:

$$\sum_{k=0}^{\infty} a_k = \sum_{i=0}^{\infty} a_i.$$

Beispiele 9.3:

allgem. Folgenglied	Partialsumme
$a_n = n$	$\sum_{k=1}^n k = 1 + 2 + 3 + \dots + n$
$a_n = \frac{1}{n}$	$\sum_{k=1}^n \frac{1}{k} = 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}$
$a_n = q^n$	$\sum_{k=0}^n q^k = 1 + q + q^2 + \dots + q^n$
$a_n = \frac{1}{n!}$	$\sum_{k=0}^n \frac{1}{k!} = 1 + 1 + \frac{1}{2!} + \dots + \frac{1}{n!}$
$a_n = (-1)^n \frac{1}{n}$	$\sum_{k=1}^n (-1)^k \frac{1}{k} = -1 + \frac{1}{2} - \frac{1}{3} \pm \dots (-1)^n \frac{1}{n} \quad \square$

Eine Reihe ist also die Folge der Partialsummen $(\sum_{k=1}^n a_k)_{n \in \mathbb{N}}$. Es stellt sich die Frage, ob diese Folgen konvergieren, d.h. ob

$$\lim_{n \rightarrow \infty} S_n = \sum_{k=1}^{\infty} a_k$$

einen endlichen Wert besitzt.

Definition:

- (1) Eine Reihe $(\sum_{k=1}^n a_k)_{n \in \mathbb{N}}$ heißt **konvergent**, wenn die Folge der Partialsummen $S_n := \sum_{k=1}^n a_k$ eine konvergente Folge ist. Liegt Konvergenz vor, so bezeichnet man den Grenzwert

$$\lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \sum_{k=1}^n a_k = \sum_{k=1}^{\infty} a_k$$

als **Summe der unendlichen Reihe**.

- (2) Eine Reihe $(\sum_{k=1}^n a_k)_{n \in \mathbb{N}}$ heißt **divergent**, wenn sie keinen Grenzwert besitzt.
- (3) Eine Reihe $(\sum_{k=1}^n a_k)_{n \in \mathbb{N}}$ heißt **absolut konvergent**, wenn die Partialsumme der Beträge $S_n := \sum_{k=1}^n |a_k|$ konvergiert.

Bemerkungen:

- (1) Eine konvergente Reihe besitzt stets einen endlichen, eindeutig bestimmten Summenwert.
- (2) Eine absolut konvergente Reihe ist stets konvergent. Die Umkehrung gilt allerdings nicht ($\rightarrow$ Beispiel 9.14)!
- (3) Eine Reihe heißt *bestimmt divergent*, wenn $\sum_{k=1}^{\infty} a_k$ entweder $+\infty$ oder $-\infty$ ist.
- (4) Die Auswertung der Partialsumme als geschlossener Ausdruck ist in manchen, seltenen Fällen möglich. Dann ist der Summenwert berechenbar. I.a. jedoch ist der Grenzwert unbekannt und man muss Konvergenzkriterien anwenden, um die Konvergenz der Reihe zu zeigen.

Wir behandeln zunächst Reihen, bei denen sich die Partialsummen auswerten lassen und lernen dann wichtige Konvergenzkriterien kennen.

► 9.1.1 Beispiele

Beispiel 9.4. Die geometrische Reihe

$$\sum_{k=0}^{\infty} q^k = 1 + q + q^2 + \dots + q^k + \dots$$

konvergiert für $|q| < 1$ und **divergiert** für $|q| \geq 1$.

Denn nach Kap. 1.2.3 gilt für die endliche geometrische Reihe:

$$S_n = \sum_{k=0}^n q^k = \frac{1 - q^{n+1}}{1 - q} \quad \text{für } q \neq 1.$$

Für $|q| < 1$ ist $\lim_{n \rightarrow \infty} q^{n+1} = 0$ und die Folge der Partialsummen hat den Grenzwert

$$S = \lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \frac{1 - q^{n+1}}{1 - q} = \frac{1}{1 - q}.$$

Folglich ist

$$\sum_{k=0}^{\infty} q^k = \frac{1}{1 - q} \quad \text{für } |q| < 1.$$

Für $|q| > 1$ divergiert q^{n+1} und damit S_n . Für $q = 1$ ist $S_n = \sum_{k=0}^n 1 = n + 1$, also divergent. Für $q = -1$ ist die Reihe ebenfalls divergent, wie das nachfolgende Beispiel zeigt. $\square$

Beispiel 9.5. Die Reihe

$$\sum_{n=0}^{\infty} (-1)^n$$

ist **divergent**. Denn die Folge der Partialsummen ist

$$\begin{aligned} S_0 &= 1, & S_1 &= 1 - 1 = 0, & S_2 &= 1 - 1 + 1 = 1, & S_3 &= 1 - 1 + 1 - 1 = 0, \\ S_4 &= 1, & S_5 &= 0, & S_6 &= 1, & S_7 &= 0, \dots \quad \text{usw.} \end{aligned}$$

Damit besitzt die Folge (S_n) keinen Grenzwert und die Reihe $\sum_{n=0}^{\infty} (-1)^n$ ist divergent. Dieses Beispiel zeigt auch, dass eine divergente Reihe nicht notwendigerweise gegen $+\infty$ oder $-\infty$ gehen muss. $\square$

Beispiel 9.6. Die arithmetische Reihe

$$\sum_{k=1}^{\infty} k = 1 + 2 + 3 + \dots + n + \dots$$

ist **divergent**. Durch vollständige Induktion wurde in Beispiel 1.3 gezeigt, dass

$$S_n = \sum_{k=1}^n k = 1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2}.$$

Folglich ist der Grenzwert

$$S = \lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \frac{1}{2} n (n+1) = \infty.$$

Die arithmetische Reihe ist damit **bestimmt divergent**. $\square$

Beispiel 9.7. Die Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k(k+1)} = \frac{1}{1 \cdot 2} + \frac{1}{2 \cdot 3} + \frac{1}{3 \cdot 4} + \dots + \frac{1}{k(k+1)} + \dots$$

ist konvergent. Wie man leicht mit vollständiger Induktion beweist, gilt für die Partialsumme

$$S_n = \sum_{k=1}^n \frac{1}{k(k+1)} = \frac{1}{1 \cdot 2} + \frac{1}{2 \cdot 3} + \dots + \frac{1}{n(n+1)} = \frac{n}{n+1}.$$

Folglich ist

$$\lim_{n \rightarrow \infty} S_n = \lim_{n \rightarrow \infty} \frac{n}{n+1} = 1 \Rightarrow \sum_{k=1}^{\infty} \frac{1}{k(k+1)} = 1. \quad \square$$

△ Beispiel 9.8. Die harmonische Reihe

$$\sum_{n=1}^{\infty} \frac{1}{n} = 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots + \frac{1}{n} + \dots$$

ist divergent: Wir vergleichen die harmonische Reihe mit einer Vergleichsreihe, deren Folgenglieder kleiner als die der harmonischen Reihe sind; die Vergleichsreihe aber schon divergiert.

Harmonische Reihe:

$$1 + \underset{\uparrow}{\frac{1}{2}} + \left(\underset{\uparrow}{\frac{1}{3}} + \underset{\uparrow}{\frac{1}{4}} \right) + \left(\underset{\uparrow}{\frac{1}{5}} + \frac{1}{6} + \frac{1}{7} + \underset{\uparrow}{\frac{1}{8}} \right) + \dots + \left(\frac{1}{2^n+1} + \dots + \frac{1}{2^{n+1}} \right) + \dots$$

Die Klammerung erfolgt dabei so, dass jeweils die Summanden

$$\frac{1}{2^n+1} + \dots + \frac{1}{2^{n+1}}$$

zusammengefasst werden. Wir ersetzen alle Terme einer Klammer durch den mit Pfeil gekennzeichneten Wert $\frac{1}{2^{n+1}}$. Dadurch verkleinern wir den Wert der Summe und erhalten die Vergleichsreihe

$$1 + \frac{1}{2} + \left(\frac{1}{4} + \frac{1}{4} \right) + \left(\frac{1}{8} + \frac{1}{8} + \frac{1}{8} + \frac{1}{8} \right) + \left(\frac{1}{16} + \dots + \frac{1}{16} \right) + \dots$$

Für diese Reihe ist

$$\sum_{i=1}^n \frac{1}{2} = n \frac{1}{2} \rightarrow \infty \quad \text{für } n \rightarrow \infty.$$

Da die Vergleichsreihe gegen ∞ divergiert, muss die harmonische Reihe, deren Glieder größer als die der Vergleichsreihe sind, ebenfalls divergieren. $\square$

Bei diesen Überlegungen geht implizit das sog. *Minorantenkriterium* ein. Es besagt, dass eine Reihe divergiert, wenn eine divergente Vergleichsreihe (Minorante) existiert, deren Reihenglieder kleiner sind als die der ursprünglichen Reihe:

Minorantenkriterium: Ist $0 < a_i \leq b_i$ ab einem $m \in \mathbb{N}$, dann gilt

$$\sum_{i=1}^{\infty} a_i \text{ divergent} \quad \Rightarrow \quad \sum_{i=1}^{\infty} b_i \text{ divergent.}$$

Folgerungen aus der Divergenz der harmonischen Reihe:

- (1) $\triangleq \lim_{n \rightarrow \infty} a_n = 0$ genügt nicht, um die Konvergenz der Reihe $\sum_{k=1}^{\infty} a_k$ sicherzustellen.
- (2) Ist a_n konvergent mit $\lim_{n \rightarrow \infty} a_n \neq 0$, dann ist die Reihe $\sum_{k=1}^{\infty} a_k$ divergent.
- (3) Ist $\sum_{k=1}^{\infty} a_k$ konvergent $\Rightarrow \lim_{n \rightarrow \infty} a_n = 0$.

$\triangleq$ **Achtung:** Eine numerische Berechnung einer Reihe reicht nicht aus, um die Konvergenz zu prüfen, bzw. im Falle der Konvergenz den Summenwert zu bestimmen!: Die harmonische Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$ ist numerisch **immer** konvergent, was im Widerspruch zu Beispiel 9.8 steht. Dieser Trugschluss rührt daher, dass numerisch nur mit einer endlichen Genauigkeit gerechnet wird. Daher ist ab einem gewissen N numerisch

$$\sum_{i=1}^N \frac{1}{i} + \frac{1}{N+1} = \sum_{i=1}^N \frac{1}{i} \quad (\text{numerisch!}),$$

da dann $\frac{1}{N+1}$ nicht mehr zum Summenwert beiträgt. Ab diesem N ändert die Reihe numerisch ihren Wert nicht mehr.

Beispiel 9.9 (Mit MAPLE-Worksheet). Um diesen Effekt zu verdeutlichen, berechnen wir die harmonische Reihe mit einer Rechengenauigkeit von 5 Stellen. Wir erhalten die folgenden Ergebnisse in Abhängigkeit von N

N	10	100	1000	10000	15000	20000	30000
<i>summe</i>	2.9290	5.1873	7.4847	9.7509	10.000	10.000	10.000

Etwa ab $N = 15000$ ändert sich der Summenwert nicht mehr, obwohl die Reihe divergiert! Ändert man die Reihenfolge der Summation, kann nahezu jeder Wert größer 10 als Summenwert erhalten werden. $\square$

Da in den wenigsten Fällen die Partialsumme als geschlossener Ausdruck vorliegt, werden Kriterien benötigt, um zu entscheiden, ob Reihen konvergieren oder nicht. Dies führt zu den sog. *Konvergenzkriterien*. Wir geben nur die drei wichtigsten an.

► 9.1.2 Majorantenkriterium

Ein sehr anschauliches Kriterium ist das *Majorantenkriterium*, welches besagt, dass eine Reihe konvergiert, wenn eine betragsmäßig größere Reihe schon konvergiert.

Majorantenkriterium:

$$\text{Ist } |a_i| \leq A_i \text{ und } \sum_{i=1}^{\infty} A_i \text{ konvergent} \Rightarrow \sum_{i=1}^{\infty} a_i \text{ konvergent.}$$

Man bezeichnet $\sum_{i=1}^{\infty} A_i$ dann als *Majorante*.

Beispiel 9.10. Für $p \geq 2$ konvergiert die Reihe $\sum_{n=1}^{\infty} \frac{1}{n^p}$:

Die konvergente Majorante ist die in Beispiel 9.7 diskutierte Reihe:

$$\sum_{k=1}^{\infty} \frac{1}{k(k+1)} = 1.$$

Denn für $p \geq 2$ gilt

$$\frac{1}{k^p} \leq \frac{1}{k^2} \leq \frac{2}{k(k+1)}.$$

Daher ist

$$\sum_{k=1}^N \frac{1}{k^p} \leq \sum_{k=1}^N \frac{2}{k(k+1)} \leq 2 \sum_{k=1}^{\infty} \frac{1}{k(k+1)} = 2$$

und $\sum_{k=1}^{\infty} \frac{1}{k^p}$ ist konvergent mit einem Summenwert ≤ 2 . $\square$

Es gilt allgemeiner der folgende Satz:

Satz: Die Reihe

$$\sum_{n=1}^{\infty} \frac{1}{n^p}$$

ist konvergent für $p > 1$ und divergent für $p \leq 1$.

Beispiele 9.11:

- ① Die Reihe $\sum_{n=1}^{\infty} \frac{1}{\sqrt{n}} = \sum_{n=1}^{\infty} \frac{1}{n^{\frac{1}{2}}}$ ist divergent, da $p = \frac{1}{2} < 1$.
- ② Die Reihe $\sum_{n=1}^{\infty} \frac{n}{\sqrt{n^5}} = \sum_{n=1}^{\infty} \frac{1}{n^{\frac{3}{2}}}$ ist konvergent, da $p = \frac{3}{2} > 1$.
- ③ Die Reihe $\sum_{n=1}^{\infty} \frac{\sin(n)}{\sqrt{n^3}}$ ist konvergent: Wegen $\left| \frac{\sin(n)}{\sqrt{n^3}} \right| \leq \frac{1}{n^{\frac{3}{2}}}$ stellt $\sum_{n=1}^{\infty} \frac{1}{n^{\frac{3}{2}}}$ eine konvergente Majorante dar. $\square$

► 9.1.3 Quotientenkriterium

Für die Anwendung bei Potenzreihen zeigt sich das folgende *Quotientenkriterium* als außerordentlich erfolgreich.

Quotientenkriterium: Die Reihe $\sum_{n=1}^{\infty} a_n$ ist konvergent, falls es ein $N \in \mathbb{N}$ und eine Zahl $q < 1$ gibt mit

$$\left| \frac{a_{n+1}}{a_n} \right| \leq q < 1$$

für alle $n \geq N$.

Begründung: Weil es auf endlich viele Glieder nicht ankommt, sei angenommen, dass $|a_{n+1}| \leq q |a_n|$ für alle n . Dann folgt induktiv

$$|a_n| \leq q |a_{n-1}| \leq q^2 |a_{n-2}| \leq q^3 |a_{n-3}| \leq \dots \leq q^n |a_0|.$$

Da $q < 1$ gilt unter Verwendung der geometrischen Reihe aus Beispiel 9.4

$$\left| \sum_{n=0}^{\infty} a_n \right| \leq \sum_{n=0}^{\infty} |a_n| \leq |a_0| \sum_{n=0}^{\infty} q^n = |a_0| \frac{1}{1-q}. \quad \square$$

Bemerkungen:

- (1) Da die geometrische Reihe für $|q| > 1$ divergiert, erhält man analog zu der Argumentation des vorherigen Beweises die Aussage:
Gibt es ein $N \in \mathbb{N}$, so dass $\left| \frac{a_{n+1}}{a_n} \right| > 1$ für alle $n \geq N$, dann divergiert die Reihe $\sum_{n=1}^{\infty} a_n$.
- (2) Für die Anwendungen bei den Potenz- und Taylor-Reihen ist es oft einfacher, die Limesform des Quotientenkriteriums zu verwenden. Hierbei berechnet man $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|$. Es gilt dann äquivalent zum Quotientenkriterium:

Limesform des Quotientenkriteriums:

$$\text{Ist } \lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| < 1 \Rightarrow \sum_{n=1}^{\infty} a_n \text{ konvergent.}$$

$$\text{Ist } \lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| > 1 \Rightarrow \sum_{n=1}^{\infty} a_n \text{ divergent.}$$

Für $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| = 1$ ist keine Aussage über die Konvergenz möglich.

Beispiele 9.12:

- ① Die Reihe

$$\sum_{n=1}^{\infty} \frac{n^2}{2^n}$$

ist konvergent. Dies folgt aus der Limesform des Quotientenkriteriums, da

$$\left| \frac{a_{n+1}}{a_n} \right| = \left| \frac{(n+1)^2}{2^{n+1}} / \frac{n^2}{2^n} \right| = \frac{2^n (n+1)^2}{2^{n+1} n^2} = \frac{1}{2} \left(1 + \frac{1}{n} \right)^2 \xrightarrow{n \rightarrow \infty} \frac{1}{2} < 1.$$

- ② Die Reihe

$$\sum_{n=0}^{\infty} \frac{1}{n!} x^n$$

ist für jedes $x \in \mathbb{R}$ konvergent. Dies folgt aus der Limesform des Quotientenkriteriums, da

$$\left| \frac{a_{n+1}}{a_n} \right| = \left| \frac{x^{n+1}}{(n+1)!} \cdot \frac{n!}{x^n} \right| = \frac{|x|}{n+1} \xrightarrow{n \rightarrow \infty} 0 < 1. \quad \square$$

Bemerkungen zum Quotientenkriterium

- (1)  **Achtung:** Im Falle $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| = 1$ ist keine Aussage möglich:
Für $a_n = \frac{1}{n}$ erhalten wir die harmonische Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$. Hier ist

$$\left| \frac{a_{n+1}}{a_n} \right| = \frac{n}{n+1} \xrightarrow{n \rightarrow \infty} 1.$$

Für $a_n = \frac{1}{n^2}$ erhalten wir die Reihe $\sum_{n=1}^{\infty} \frac{1}{n^2}$. Auch hier gilt

$$\left| \frac{a_{n+1}}{a_n} \right| = \frac{n^2}{(n+1)^2} \xrightarrow{n \rightarrow \infty} 1.$$

Für beide Fälle liefert die Limesform des Quotientenkriteriums als Wert 1. Nach Beispiel 9.8 *divergiert* die erste und nach Beispiel 9.10 *konvergiert* die zweite Reihe. Damit kann das Quotientenkriterium in solchen Fällen nicht angewendet werden!

- (2) Das Quotientenkriterium ist also nur eine hinreichende, aber keine notwendige Bedingung für die Konvergenz einer Reihe.
- (3) Man beachte, dass das Konvergenzkriterium nur Aufschluss darüber gibt, ob eine Reihe konvergiert oder nicht; es liefert **keinen** Anhaltspunkt über den Summenwert. Insbesondere stimmt $\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|$ nicht mit dem Wert der Reihe überein!



Die MAPLE-Prozedur **quotkrit** wendet auf eine gegebene Reihe mit Reihengliedern $a(n)$ das Quotientenkriterium in der Limesform an.

► 9.1.4 Leibniz-Kriterium

Für *alternierende* Reihen, Reihen deren Glieder abwechselnd positiv und negativ sind, existiert ein von Leibniz (1646 - 1716) stammendes Kriterium. Alternierende Reihen haben die Form $\sum_{n=1}^{\infty} (-1)^{n+1} a_n = a_1 - a_2 + a_3 - a_4 \pm \dots$ mit $a_n > 0$. Das Vorzeichen $(-1)^{n+1}$ wechselt dabei ständig.

Leibniz-Kriterium: Eine alternierende Reihe

$$\sum_{n=1}^{\infty} (-1)^{n+1} a_n = a_1 - a_2 + a_3 - a_4 \pm \dots$$

ist konvergent, falls $a_1 > a_2 > a_3 > a_4 > \dots > 0$ und $\lim_{n \rightarrow \infty} a_n = 0$.

Eine alternierende Reihe konvergiert also, wenn die Beträge der Glieder eine streng monoton fallende Nullfolge bilden.

Beispiel 9.13. $\sum_{n=1}^{\infty} (-1)^{n+1} \frac{1}{n!}$ ist konvergent. Die Glieder der Reihe sind alternierend und die Beträge der Glieder

$$\frac{1}{1!} > \frac{1}{2!} > \frac{1}{3!} > \dots > \frac{1}{n!} > \frac{1}{(n+1)!} > \dots > 0$$

bilden eine streng monoton fallende Nullfolge. Nach dem Leibniz-Kriterium konvergiert die Reihe. $\square$

Beispiel 9.14. Die *alternierende harmonische Reihe*

$$\sum_{n=1}^{\infty} (-1)^{n+1} \frac{1}{n} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \pm \dots$$

ist **konvergent**. Die Glieder der Reihe sind alternierend und deren Beträge

$$1 > \frac{1}{2} > \frac{1}{3} > \dots > \frac{1}{n} > \frac{1}{n+1} > \dots > 0$$

bilden eine streng monoton fallende Nullfolge. Nach dem Leibniz-Kriterium konvergiert die Reihe. $\square$

Beispiel 9.15. $\sum_{n=1}^{\infty} (-1)^{n+1}$ divergiert nach Beispiel 9.5. Das Leibniz-Kriterium ist **nicht** anwendbar, da $|a_n| = 1$ keine Nullfolge ist. $\square$

Bemerkungen:

- (1) Absolut konvergente Reihen sind auch konvergent im gewöhnlichen Sinne. Die Umkehrung gilt aber nicht!: Die alternierende harmonische Reihe ist konvergent ($\rightarrow$ Beispiel 9.14) aber nicht absolut konvergent, da die harmonische Reihe $\sum_{n=1}^{\infty} \left| (-1)^{n+1} \frac{1}{n} \right| = \sum_{n=1}^{\infty} \frac{1}{n}$ nach Beispiel 9.8 divergiert.
- (2) Bei der Anwendung des Leibniz-Kriteriums genügt es nicht, nur die Eigenschaft "alternierend" nachzuprüfen! Selbst wenn die Reihenglieder alternierendes Vorzeichen besitzen und eine Nullfolge bilden, folgt **nicht** die Konvergenz, wie die Reihe

$$\sum_{k=1}^{\infty} (-1)^k \left\{ \frac{1}{\sqrt{k+1}} + \frac{(-1)^k}{k+1} \right\}$$

zeigt. Die Reihenglieder sind alternierend, bilden aber keine betragsmäßig monoton fallende Nullfolge.

9.2 Potenzreihen

Dieser Abschnitt stellt den Übergang von den Zahlenreihen zu den Taylor-Reihen dar. Die Konvergenzkriterien der Zahlenreihen werden übertragen auf Potenzreihen, um den Definitionsbereich der Potenzreihen zu bestimmen.

Sind die Summanden in einer Reihe selbst Funktionen einer Variablen x , so stellt der Ausdruck $\sum_{n=0}^{\infty} a_n(x)$ eine Funktion dar, eine sog. **Funktionenreihe**. Ein wichtiger Spezialfall solcher Funktionenreihen sind die *Potenzreihen*.

Definition: Eine Funktion der Form

$$\sum_{n=0}^{\infty} a_n x^n = a_0 + a_1 x + \dots + a_n x^n + \dots$$

heißt **Potenzreihe**. Der Definitionsbereich einer Potenzreihe besteht aus allen reellen Zahlen x , für die $\sum_{n=0}^{\infty} a_n x^n$ konvergiert. Man nennt daher die Menge

$$K := \left\{ x \in \mathbb{R} : \sum_{n=0}^{\infty} a_n x^n \text{ konvergent} \right\}$$

den **Konvergenzbereich** der Potenzreihe.

Bemerkungen:

- (1) Man bezeichnet $a_0, a_1, a_2, \dots, a_n, \dots$ als die *Koeffizienten* der Potenzreihe.
- (2) Für jedes feste x ist eine Potenzreihe eine Zahlenreihe.
- (3) Eine etwas allgemeinere Darstellung von Potenzreihen erhält man durch Ausdrücke der Form

$$\sum_{n=0}^{\infty} a_n (x - x_0)^n = a_0 + a_1 (x - x_0) + \dots + a_n (x - x_0)^n + \dots$$

Man bezeichnet dann die Stelle x_0 als den *Entwicklungspunkt* der Reihe.

Beispiel 9.16. $\sum_{n=0}^{\infty} n x^n = 1 x + 2 x^2 + 3 x^3 + \dots + n x^n + \dots$ □

Beispiel 9.17. $\sum_{n=0}^{\infty} \frac{1}{n!} x^n = 1 + x + \frac{1}{2!} x^2 + \frac{1}{3!} x^3 + \dots + \frac{1}{n!} x^n + \dots$ □

Beispiel 9.18. f sei im Punkte $x_0 \in \mathbb{D}$ beliebig oft differenzierbar. Dann ist

$$\sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x_0) (x - x_0)^n$$

eine Potenzreihe mit Entwicklungspunkt x_0 und den Koeffizienten

$$a_n = \frac{1}{n!} f^{(n)}(x_0).$$

Eine solche Reihe bezeichnet man als *Taylor-Reihe* der Funktion f am Entwicklungspunkt x_0 ($\rightarrow$ §9.3). $\square$

Beispiel 9.19 (Geometrische Potenzreihe): Nach Beispiel 9.4 ist die Potenzreihe

$$\sum_{n=0}^{\infty} x^n = 1 + x + x^2 + \dots + x^n + \dots = \frac{1}{1-x}$$

für $|x| < 1$ konvergent und für $|x| \geq 1$ divergent. Der **Konvergenzbereich** ist daher $K = (-1, 1)$.

Beispiel 9.20. Wir berechnen den Konvergenzbereich der Potenzreihe

$$\sum_{n=1}^{\infty} \frac{1}{n} x^n = x + \frac{1}{2} x^2 + \frac{1}{3} x^3 + \dots + \frac{1}{n} x^n + \dots$$

Dazu wenden wir für ein beliebiges aber festes $x \in \mathbb{R}$ das Quotientenkriterium mit $b_n = \frac{1}{n} x^n$ an:

$$\left| \frac{b_{n+1}}{b_n} \right| = \left| \frac{x^{n+1}}{n+1} / \frac{x^n}{n} \right| = \left| \frac{x^{n+1}}{n+1} \frac{n}{x^n} \right| = \frac{n}{n+1} |x| \xrightarrow{n \rightarrow \infty} |x|.$$

Damit konvergiert die Reihe für $|x| < 1$ und divergiert für $|x| > 1$. Für $|x| = 1$ müssen getrennte Untersuchungen durchgeführt werden, indem die jeweiligen Werte in die Reihe eingesetzt werden:

Für $x = 1$ ist

$$\sum_{n=1}^{\infty} \frac{1}{n} x^n = \sum_{n=1}^{\infty} \frac{1}{n} 1^n = \sum_{n=1}^{\infty} \frac{1}{n}$$

die harmonische Reihe, also nach Beispiel 9.8 divergent.

Für $x = -1$ ist

$$\sum_{n=1}^{\infty} \frac{1}{n} x^n = \sum_{n=1}^{\infty} \frac{1}{n} (-1)^n.$$

Die alternierende harmonische Reihe ist nach Beispiel 9.14 konvergent. Damit ist der Konvergenzbereich $K = [-1, 1)$. $\square$

⊙ **Konvergenzverhalten von Potenzreihen**

Man kann für beliebige Potenzreihen $\sum_{n=0}^{\infty} a_n x^n$ das Konvergenzverhalten charakterisieren. Grundlage hierfür ist der folgende Satz.

Satz über das Konvergenzverhalten von Potenzreihen: Jede Potenzreihe

$$\sum_{n=0}^{\infty} a_n x^n = a_0 + a_1 x + a_2 x^2 + \dots + a_n x^n + \dots$$

besitzt einen eindeutig bestimmten **Konvergenzradius** ρ ($0 \leq \rho \leq \infty$) mit den Eigenschaften:

- (1) Die Reihe konvergiert für alle x mit $|x| < \rho$.
- (2) Die Reihe divergiert für alle x mit $|x| > \rho$.
- (3) Für $|x| = \rho$ ist keine allgemeine Aussage möglich.

Begründung: Zur Bestimmung von ρ wenden wir das Quotientenkriterium auf die Reihe $\sum_{n=0}^{\infty} b_n$ mit $b_n = a_n x^n$ an:

$$\left| \frac{b_{n+1}}{b_n} \right| = \left| \frac{a_{n+1} x^{n+1}}{a_n x^n} \right| = \left| \frac{a_{n+1}}{a_n} \right| |x| \xrightarrow[n \rightarrow \infty]{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| \cdot |x|.$$

Nach der Limesform des Quotientenkriteriums konvergiert die Reihe für

$$\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| \cdot |x| < 1 \quad \Leftrightarrow \quad |x| < \frac{1}{\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|} = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|$$

und sie divergiert für

$$|x| > \frac{1}{\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right|} = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|.$$

Setzen wir $\rho := \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|$, so sind die Aussagen des Satzes nachgeprüft und wir haben den Konvergenzradius berechnet. $\square$

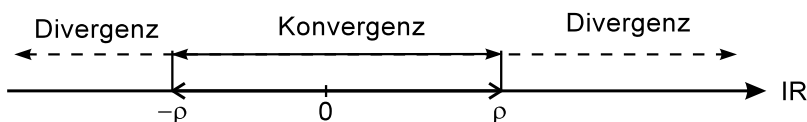
Satz: (Konvergenzradius)

Der Konvergenzradius ρ einer Potenzreihe $\sum_{n=0}^{\infty} a_n x^n$ ist gegeben durch

$$\rho = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|.$$

Bemerkungen:

- (1) Der Sonderfall $\rho = \infty$ ist möglich, denn dann ist $|x| < \rho$ immer erfüllt und der Konvergenzbereich ist $K = \mathbb{R}$.
- (2) Es gibt Potenzreihen mit $\rho = 0$. Diese Reihen konvergieren für kein $x \in \mathbb{R}$.
- (3) Für $|x| = \rho$, d.h. $x = \pm \rho$, kann keine allgemeingültige Aussage getroffen werden. An diesen Randstellen müssen getrennte Untersuchungen durchgeführt werden.
- (4) Es ist keine Aussage möglich, falls $\lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|$ nicht existiert.
- (5) **⚠ Achtung:** Beim Quotientenkriterium wird das Verhältnis $\left| \frac{a_{n+1}}{a_n} \right|$ gebildet, während für die Berechnung des Konvergenzradius das Verhältnis $\left| \frac{a_n}{a_{n+1}} \right|$ bestimmt wird!
- (6) Eine Potenzreihe $\sum_{n=0}^{\infty} a_n x^n$ konvergiert stets innerhalb des zum Nullpunkt symmetrischen Intervalls $|x| < \rho$ und divergiert außerhalb. Abb. 9.2 zeigt die graphische Darstellung des Konvergenzbereichs.

**Abb. 9.2.** Konvergenzbereich einer Potenzreihe**Beispiel 9.21.** Die Reihe

$$\sum_{n=0}^{\infty} \frac{1}{n!} x^n = 1 + x + \frac{1}{2!} x^2 + \dots + \frac{1}{n!} x^n + \dots$$

konvergiert für alle $x \in \mathbb{R}$. Denn der Konvergenzradius ist

$$\rho = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right| = \lim_{n \rightarrow \infty} \left| \frac{\frac{1}{n!}}{\frac{1}{(n+1)!}} \right| = \lim_{n \rightarrow \infty} \frac{(n+1)!}{n!} = \lim_{n \rightarrow \infty} (n+1) = \infty. \quad \square$$

Beispiel 9.22. Die Reihe

$$\sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1} = x - \frac{1}{3!} x^3 + \frac{1}{5!} x^5 - \frac{1}{7!} x^7 \pm \dots$$

konvergiert für alle $x \in \mathbb{R}$. Denn mit $a_n = \frac{(-1)^n}{(2n+1)!}$ folgt

$$\left| \frac{a_n}{a_{n+1}} \right| = \left| \frac{(-1)^n}{(2n+1)!} \frac{(2n+3)!}{(-1)^{n+1}} \right| = \frac{(2n+3)!}{(2n+1)!} = (2n+2)(2n+3)$$

$$\Rightarrow \rho = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right| = \infty \Rightarrow K = \mathbb{R}. \quad \square$$

Beispiel 9.23. Wir untersuchen die Potenzreihe

$$\sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} (x-1)^n$$

auf ihre Konvergenzeigenschaften. Dazu wenden wir das Quotientenkriterium auf die Reihe mit den Summanden $b_n = \frac{(-1)^{n+1}}{n} (x-1)^n$ an:

$$\frac{b_{n+1}}{b_n} = \frac{(-1)^{n+2}}{n+1} (x-1)^{n+1} \cdot \frac{n}{(-1)^{n+1}} \frac{1}{(x-1)^n} = \frac{n}{n+1} (-1) (x-1)$$

$$\hookrightarrow \lim_{n \rightarrow \infty} \left| \frac{b_{n+1}}{b_n} \right| = \lim_{n \rightarrow \infty} \frac{n}{n+1} |x-1| = |x-1|.$$

Somit konvergiert die Reihe für $|x-1| < 1$ und divergiert für $|x-1| > 1$. Für den Fall $|x-1| = 1$ werden getrennte Untersuchungen durchgeführt. Aus $|x-1| = 1$ folgt entweder $x-1 = 1 \hookrightarrow x = 2$ oder $x-1 = -1 \hookrightarrow x = 0$.

Für $x = 2$ ist

$$\sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} (2-1)^n = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n}$$

die alternierende harmonische Reihe und damit konvergent.

Für $x = 0$ ist

$$\sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} (0-1)^n = \sum_{n=1}^{\infty} \frac{(-1)^{n+1} (-1)^n}{n} = - \sum_{n=1}^{\infty} \frac{1}{n}$$

die harmonische Reihe und damit divergent. $\Rightarrow K = (0, 2]$. □

Bemerkung: Bei Potenzreihen der Form $\sum_{n=0}^{\infty} a_n (x-x_0)^n$ wird der Konvergenzradius ebenfalls berechnet durch die Formel

$$\rho = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|.$$

- (1) Die Reihe konvergiert für $|x-x_0| < \rho$.
- (2) Die Reihe divergiert für $|x-x_0| > \rho$.
- (3) Für $|x-x_0| = \rho$ ($\hookrightarrow x = \pm \rho + x_0$) kann keine allgemeingültige Aussage getroffen werden.

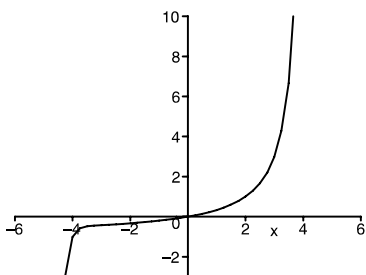


Der **Konvergenzradius** wird in MAPLE mit der Prozedur **konv_radi-us** bestimmt. Der **Konvergenzbereich** einer Potenzreihe kann aber auch graphisch visualisiert werden, indem man die Potenzreihe mit wachsender Ordnung in Form einer **Animation** darstellt.

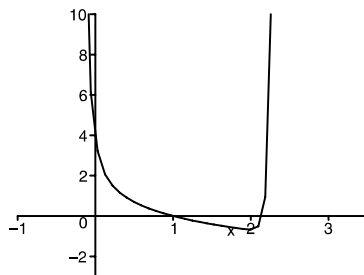
Beispiele 9.24 (Mit MAPLE-Worksheet):

- ① Der Konvergenzradius der Potenzreihe $\sum_{i=1}^{\infty} \frac{1}{4^i} x^i$ ist $\rho = 4$. Der Konvergenzbereich ist das offene Intervall $(-4, 4)$.
- ② Von der Potenzreihe $\sum_{i=1}^{\infty} (-1)^i \frac{1}{i} (x-1)^i$ ist der Konvergenzradius $\rho = 1$. Der Konvergenzbereich ist das halb offene Intervall $(0, 2]$.

Dargestellt wird in der MAPLE-Animation jeweils nur die Partialsumme bis $n = 25$ bzw. $n = 26$.



Partialsumme $\sum_{i=1}^{25} \frac{1}{4^i} x^i$



Partialsumme $\sum_{i=1}^{26} (-1)^i \frac{1}{i} (x-1)^i$



Man erkennt in der Animation, dass sich im Innern des Konvergenzbereichs die Reihen stabilisieren, außerhalb gehen sie gegen Unendlich. Bei der ersten Reihe entnimmt man den Konvergenzbereich zwischen -4 und 4 , während er bei der zweiten Reihe von 0 bis 2 geht. $\square$

Zum Abschluss dieses Abschnitts fassen wir noch einige wichtige Eigenschaften von Potenzreihen zusammen.

Wichtige Eigenschaften von Potenzreihen:

- (1) Eine Potenzreihe konvergiert innerhalb ihres Konvergenzbereichs absolut.
- (2) Eine Potenzreihe darf innerhalb ihres Konvergenzbereichs differenziert und integriert werden. Die so erhaltenen Potenzreihen besitzen den gleichen Konvergenzradius wie die ursprüngliche Reihe.
- (3) Zwei Potenzreihen dürfen im gemeinsamen Konvergenzbereich der Reihen gliedweise addiert und multipliziert werden.

9.3 Taylor-Reihen

Wir kommen nun zum zentralen Thema dieses Kapitels: die Taylor-Reihen. Die Aussage des Taylorschen Satzes ist, dass sich fast jede elementare Funktion in der Umgebung eines Punktes x_0 durch Polynome beliebig genau annähern lässt. Es zeigt sich sogar, dass diese Funktionen sich durch eine Potenzreihe der Form

$$\sum_{n=0}^{\infty} a_n (x - x_0)^n$$

darstellen lassen. Neben der Bestimmung der Koeffizienten a_n werden wir Information darüber gewinnen, welcher Fehler maximal auftritt, wenn diese Reihe nach endlich vielen Summationsgliedern abgebrochen wird. Damit erhalten wir zum einen eine Methode, die elementaren Funktionen

$$e^x, \sin x, \sqrt{x}, \ln x \quad \text{usw.}$$

mit beliebiger Genauigkeit zu berechnen, zum anderen Näherungsformeln für diese Funktionen.

Beispiel 9.25 (Einführung): Nach Beispiel 9.19 gilt für die geometrische Potenzreihe

$$1 + x + x^2 + \dots + x^n + \dots = \sum_{n=0}^{\infty} x^n = \frac{1}{1-x}$$

für $|x| < 1$. D.h. die Potenzreihe $\sum_{n=0}^{\infty} x^n$ stimmt mit der Funktion $\frac{1}{1-x}$ für alle $x \in (-1, 1)$ überein. Außerhalb dieses offenen Intervalls ist zwar $\frac{1}{1-x}$ noch definiert ($x \neq 1$), aber nicht mehr die Potenzreihe. Wir leiten eine Formel heuristisch her, die es uns erlaubt, für elementare Funktionen die zugehörige Potenzreihe aufzustellen. $\square$

Herleitung der Taylor-Polynome. Gegeben sei eine Funktion $f(x)$, siehe Abb. 9.3. Gesucht ist eine Näherung der Funktion in der Umgebung des Punktes $x_0 \in \mathbb{D}$. Die Funktion f sei in dieser Umgebung mehrmals differenzierbar.

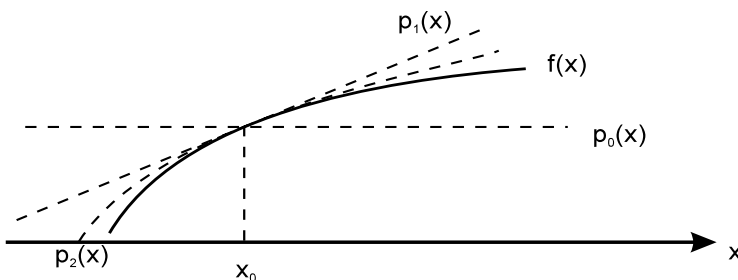


Abb. 9.3. Funktion f und Näherungen in der Umgebung von x_0

- (0.) Die "nullte" Näherung p_0 an die Funktion erhält man, wenn die konstante Funktion

$$p_0(x) = f(x_0)$$

gewählt wird. Die Funktion p_0 hat mit f nur den Funktionswert an der Stelle x_0 gemeinsam.

- (1.) Die lineare Näherung p_1 an die Funktion erhält man, wenn man die Tangente in x_0 wählt:

$$p_1(x) = f(x_0) + f'(x_0)(x - x_0).$$

Die Tangente hat mit der Funktion sowohl den Funktionswert, als auch die Ableitung an der Stelle x_0 gemeinsam.

- (2.) Gesucht ist eine quadratische Funktion p_2 , die im Punkte x_0 *zusätzlich* die gleiche Krümmung wie f aufweist:

Ansatz: $p_2(x) = f(x_0) + f'(x_0)(x - x_0) + c(x - x_0)^2.$

Bedingung: $p_2''(x_0) \stackrel{!}{=} f''(x_0).$

Wegen $p_2''(x) = 1 \cdot 2 \cdot c$, folgt $p_2''(x_0) = 1 \cdot 2 \cdot c = f''(x_0)$

$$\Rightarrow c = \frac{1}{2!} f''(x_0)$$

$$\Rightarrow p_2(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{f''(x_0)}{2!}(x - x_0)^2.$$

- (3.) Gesucht ist die kubische Funktion p_3 , die im Punkte x_0 *zusätzlich* die 3. Ableitung mit f gemeinsam hat:

Ansatz: $p_3(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2!} f''(x_0)(x - x_0)^2 + d(x - x_0)^3.$

Bedingung: $p_3'''(x_0) \stackrel{!}{=} f'''(x_0).$

Wegen $p_3'''(x_0) = 1 \cdot 2 \cdot 3 \cdot d \stackrel{!}{=} f'''(x_0)$

$$\Rightarrow d = \frac{1}{3!} f'''(x_0)$$

$$\Rightarrow p_3(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2!} f''(x_0)(x - x_0)^2 + \frac{1}{3!} f'''(x_0)(x - x_0)^3.$$

⋮

(n.) Eine bessere Approximation an die Funktion f in einer Umgebung des Punktes x_0 gewinnt man, indem jeweils Terme der Form

$$\frac{1}{n!} f^{(n)}(x_0) (x - x_0)^n$$

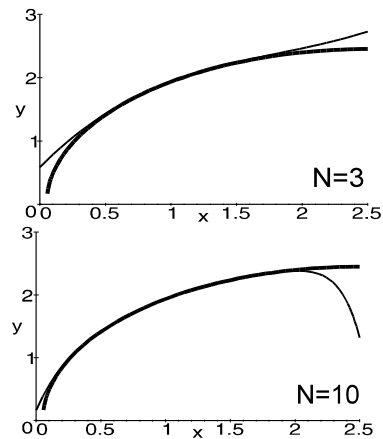
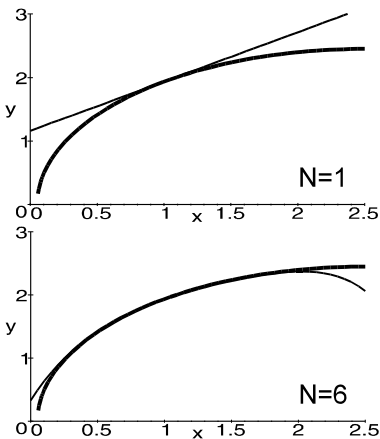
hinzugenommen werden, so dass das n -te Näherungspolynom (das **Taylor-Polynom vom Grade n**) gegeben ist durch

$$\begin{aligned} p_n(x) &= f(x_0) + f'(x_0)(x - x_0) + \dots + \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n \\ &= \sum_{i=0}^n \frac{1}{i!} f^{(i)}(x_0)(x - x_0)^i. \end{aligned}$$



Visualisierung mit MAPLE. Zur Veranschaulichung der Konvergenz der Taylor-Polynome p_n an die Funktion f wählen wir eine [Animation](#)

mit MAPLE für die Funktion $f(x) = \sqrt{6 - (x - 2.5)^2}$ am Entwicklungspunkt $x_0 = 1$. Dazu bestimmen wir die ersten 10 Taylor-Polynome.



Durch die Animation erkennt man deutlich, dass mit wachsendem Grad des Taylor-Polynoms der Bereich sich vergrößert, in dem Funktion und Taylor-Polynom graphisch übereinstimmen. Für $N = 10$ lässt sich im Bereich $0.5 \leq x \leq 1.7$ graphisch kein Unterschied zwischen der Funktion f und dem Näherungspolynom p_{10} feststellen. Es stellt sich somit die Frage, wie groß die Abweichung der Näherungsfunktion $p_n(x)$ zur Funktion f in der Umgebung von x_0 ist. Aufschluss darüber gibt der folgende Satz.

Satz von Taylor. Gegeben sei eine in $x_0 \in \mathbb{D}$ $(m+1)$ -mal stetig differenzierbare Funktion f . Dann gilt die **Taylorsche Formel**

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \dots + \frac{1}{m!} f^{(m)}(x_0)(x - x_0)^m + R_m(x)$$

mit dem Restglied

$$R_m(x) = \frac{1}{(m+1)!} f^{(m+1)}(\xi)(x - x_0)^{m+1} \quad (x \in \mathbb{D})$$

und ξ einem nicht näher bekannten Wert, der zwischen x und x_0 liegt.

Der Satz von Taylor (1685 - 1731) spezifiziert die Zwischenstelle ξ zwischen x und x_0 nicht näher. Daher kann man nicht exakt die Abweichung der Näherungsfunktion $p_n(x)$ zur Funktion f angeben. Für die konkreten Anwendungen wird diese Tatsache aber keine Rolle spielen, da wir für das Restglied $R_m(x)$ eine Obergrenze angeben. Wenn das Restglied $R_m(x) \xrightarrow{m \rightarrow \infty} 0$ erfüllt, so erhält man

Satz über Taylor-Reihen. Ist f eine in $x_0 \in \mathbb{D}$ beliebig oft differenzierbare Funktion und erfüllt das Restglied $R_m(x) \rightarrow 0$ für $m \rightarrow \infty$, so gilt

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2!} f''(x_0)(x - x_0)^2 + \dots \\ &\quad \dots + \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n + \dots \\ &= \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n. \end{aligned}$$

Diese Potenzreihe heißt die **Taylor-Reihe zur Funktion f am Entwicklungspunkt x_0** .

Bemerkungen:

- (1) Der Konvergenzradius der Taylor-Reihe ist nicht notwendigerweise > 0 .
- (2) Falls die Taylor-Reihe von f konvergiert, muss sie nicht notwendigerweise gegen $f(x)$ konvergieren.
- (3) Die Taylor-Reihe konvergiert genau dann gegen $f(x)$, wenn das Restglied $R_m(x)$ für $m \rightarrow \infty$ gegen Null geht. In diesem Fall stimmen die Funktion

und die Taylor-Reihe für alle x aus dem Konvergenzbereich der Potenzreihe überein.

- (4) Ist der Entwicklungspunkt $x_0 = 0$, so nennt man die Taylor-Reihe oftmals auch *MacLaurinsche Reihe*.
- (5) Ist f eine gerade Funktion, dann treten in der Taylor-Reihe nur Terme mit geraden Potenzen auf. Ist f eine ungerade Funktion, dann nur Terme mit ungeraden Potenzen.

Im Folgenden berechnen wir die Taylor-Reihen von wichtigen Funktionen; u.a. der Exponential- und Logarithmusfunktion bzw. den trigonometrischen Funktionen.

Beispiel 9.26 (Exponentialfunktion): Die **Taylor-Reihe von e^x** mit dem Entwicklungspunkt $x_0 = 0$:

Wegen	$f(x) = e^x$ $f'(x) = e^x$ $f''(x) = e^x$ $f'''(x) = e^x$ $\vdots$ $f^{(n)}(x) = e^x$	folgt	$f(0) = 1$ $f'(0) = 1$ $f''(0) = 1$ $f'''(0) = 1$ $\vdots$ $f^{(n)}(0) = 1$
-------	--	-------	--

Damit ist die Taylor-Reihe von e^x :

$$1 + x + \frac{1}{2!} x^2 + \frac{1}{3!} x^3 + \dots + \frac{1}{n!} x^n + \dots = \sum_{i=0}^{\infty} \frac{1}{i!} x^i.$$

Da der Konvergenzradius dieser Potenzreihe $\rho = \infty$ ($\rightarrow$ Beispiel 9.21), ist $K = \mathbb{R}$. Für das Restglied gilt

$$R_m(x) = \frac{1}{(m+1)!} f^{(m+1)}(\xi) (x - x_0)^{m+1} = \frac{x^{m+1}}{(m+1)!} e^{\xi} \quad \text{für } \xi \in [-x, x]$$

$$\Rightarrow |R_m(x)| \leq \frac{|x|^{m+1}}{(m+1)!} e^{\xi} \rightarrow 0 \quad \text{für } m \rightarrow \infty.$$

Also stimmt die Taylor-Reihe mit der Funktion überein und für alle $x \in \mathbb{R}$ gilt

$$e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n.$$

□

Beispiel 9.27 (Sinusfunktion): Die **Taylor-Reihe von** $f(x) = \sin x$ mit dem Entwicklungspunkt $x_0 = 0$:

Wegen	$f(x) = \sin x$ $f'(x) = \cos x$ $f''(x) = -\sin x$ $f'''(x) = -\cos x$ $f^{(4)}(x) = \sin x$ $f^{(5)}(x) = \cos x$ $f^{(6)}(x) = -\sin x$ $\vdots$	folgt	$f(0) = 0$ $f'(0) = 1$ $f''(0) = 0$ $f'''(0) = -1$ $f^{(4)}(0) = 0$ $f^{(5)}(0) = 1$ $f^{(6)}(0) = 0$ $\vdots$
-------	--	-------	---

Es ist also $f^{(2n)}(0) = 0$ und $f^{(2n+1)}(0) = (-1)^n$, so dass nur die ungeraden Exponenten in der Taylor-Reihe auftreten und zwar mit alternierendem Vorzeichen:

$$x - \frac{1}{3!}x^3 + \frac{1}{5!}x^5 - \frac{1}{7!}x^7 \pm \dots = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}.$$

Nach Beispiel 9.23 ist der Konvergenzradius $\rho = \infty$ und analog zum Beispiel 9.26 gilt $R_m(x) \rightarrow 0$ für $m \rightarrow \infty$. Damit stimmt die Taylor-Reihe für alle $x \in \mathbb{R}$ mit $\sin x$ überein:

$$\sin(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}.$$

□

Beispiel 9.28 (Kosinusfunktion): Die **Taylor-Reihe von** $f(x) = \cos x$ mit dem Entwicklungspunkt $x_0 = 0$ ergibt sich sofort aus obigem Beispiel: Da die Potenzreihe gliedweise innerhalb des Konvergenzbereichs differenziert werden darf, ist für alle $x \in \mathbb{R}$

$$\cos(x) = \sin'(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} (2n+1) x^{2n}$$

$$\Rightarrow \cos(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n}$$

□

Beispiel 9.29 (Logarithmusfunktion): Die **Taylor-Reihe von $\ln x$** , $x > 0$, mit dem Entwicklungspunkt $x_0 = 1$:

$f(x) = \ln x$	$f(1) = 0$
$f'(x) = x^{-1}$	$f'(1) = 1$
$f''(x) = (-1) x^{-2}$	$f''(1) = (-1)$
$f'''(x) = (-1)(-2) x^{-3}$	$f'''(1) = (-1)(-2)$
$f^{(4)}(x) = (-1)(-2)(-3) x^{-4}$	$f^{(4)}(1) = (-1)(-2)(-3)$
$f^{(5)}(x) = (-1)(-2)(-3)(-4) x^{-5}$	$f^{(5)}(1) = (-1)(-2)(-3)(-4)$
$f^{(6)}(x) = (-1)(-2)(-3)(-4)(-5) x^{-6}$	$f^{(6)}(1) = (-1)(-2)(-3)(-4)(-5)$
$\vdots$	$\vdots$
$f^{(n)}(x) = (-1)^{n+1} \frac{(n-1)!}{x^n}$	$f^{(n)}(1) = (-1)^{n+1} (n-1)!$

Damit ergeben sich die Taylor-Koeffizienten für $n \geq 1$ zu

$$\frac{f^{(n)}(1)}{n!} = \frac{(-1)^{n+1} (n-1)!}{n!} = \frac{(-1)^{n+1}}{n}.$$

Da das Restglied $R_m(x) \rightarrow 0$ für $m \rightarrow \infty$ geht, ist die Taylor-Reihe für $\ln x$ am Punkte $x_0 = 1$ gegeben durch

$$\ln x = (x-1) - \frac{1}{2} (x-1)^2 + \frac{1}{3} (x-1)^3 \pm \dots \pm \frac{(-1)^{n+1}}{n} (x-1)^n \pm \dots$$

$$\Rightarrow \ln x = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} (x-1)^n \quad \text{für } x \in (0, 2].$$

Nach Beispiel 9.23 ist der Konvergenzbereich $K = (0, 2]$. Speziell für $x = 2$ gilt

$$\ln 2 = \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n}.$$

Die Summe der alternierenden harmonischen Reihe hat den Wert $\ln 2$. □

Beispiel 9.30 (Binomische Reihe): Die **Taylor-Reihe der Binomischen Reihe** $(1+x)^\alpha$ am Entwicklungspunkt $x_0 = 0$ lautet für beliebiges $\alpha \in \mathbb{R}$:

$$(1+x)^\alpha = \sum_{k=0}^{\infty} \binom{\alpha}{k} x^k \quad \text{für } x \in (-1, 1),$$

wenn wir die verallgemeinerten Binomialkoeffizienten definieren

$$\binom{\alpha}{0} := 1 \quad \text{und} \quad \binom{\alpha}{k} := \frac{\alpha(\alpha-1)(\alpha-2) \cdots (\alpha-k+1)}{k!}.$$

Denn aus

$$\begin{array}{ll}
 f(x) = (1+x)^\alpha & f(0) = 1 \\
 f'(x) = \alpha (1+x)^{\alpha-1} & f'(0) = \alpha \\
 f''(x) = \alpha (\alpha-1) (1+x)^{\alpha-2} & f''(0) = \alpha (\alpha-1) \\
 f'''(x) = \alpha (\alpha-1) (\alpha-2) (1+x)^{\alpha-3} & f'''(0) = \alpha (\alpha-1) (\alpha-2) \\
 \vdots & \vdots \\
 f^{(n)}(x) = \alpha (\alpha-1) \cdot \dots & f^{(n)}(0) = \alpha (\alpha-1) \cdot \dots \\
 \dots (\alpha-n+1) (1+x)^{\alpha-n} & \dots (\alpha-n+1)
 \end{array}$$

folgt für die Taylor-Koeffizienten

$$\frac{f^{(k)}(x_0)}{k!} = \frac{\alpha (\alpha-1) (\alpha-2) \cdot \dots (\alpha-k+1)}{k!} = \binom{\alpha}{k}$$

und für die Taylor-Reihe

$$(1+x)^\alpha = \sum_{k=0}^{\infty} \binom{\alpha}{k} x^k.$$

Der Konvergenzbereich ergibt sich mit dem Quotientenkriterium zu $K = (-1, 1)$.

Spezialfälle:

① $\alpha = -1$ (geometrische Reihe):

$$\frac{1}{1+x} = 1 - x + x^2 - x^3 \pm \dots = \sum_{k=0}^{\infty} (-1)^k x^k.$$

② $\alpha = -2$ (Ableitung der geometrischen Reihe):

$$\frac{1}{(1+x)^2} = 1 - 2x + 3x^2 \pm \dots = \sum_{k=0}^{\infty} (-1)^k (k+1) x^k.$$

③ $\alpha = \frac{1}{2}$:

$$\sqrt{1+x} = 1 + \frac{1}{2}x - \frac{1}{2 \cdot 4}x^2 + \frac{1 \cdot 3}{2 \cdot 4 \cdot 6}x^3 - \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6 \cdot 8}x^4 \pm \dots$$

④ $\alpha = -\frac{1}{2}$:

$$\frac{1}{\sqrt{1+x}} = 1 - \frac{1}{2}x + \frac{1 \cdot 3}{2 \cdot 4}x^2 - \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6}x^3 + \frac{1 \cdot 3 \cdot 5 \cdot 7}{2 \cdot 4 \cdot 6 \cdot 8}x^4 \mp \dots \quad \square$$

Häufig wird die Berechnung der Taylor-Reihe einer Funktion durch Differenziation bzw. Integration auf bekannte Potenzreihen zurückgeführt, wie die folgenden beiden Beispiele zeigen.

Beispiel 9.31 (Arkustangensfunktion): Die **Taylor-Reihe von** $f(x) = \arctan(x)$ am Entwicklungspunkt $x_0 = 0$:

Aus $f(x) = \arctan(x)$

$$\Rightarrow f'(x) = \frac{1}{1+x^2}.$$

Nach Beispiel 9.19 ist für $|x| < 1$

$$\frac{1}{1+x^2} = \frac{1}{1-(-x^2)} = \sum_{n=0}^{\infty} (-x^2)^n = \sum_{n=0}^{\infty} (-1)^n x^{2n}.$$

Da Potenzreihen gliedweise integriert werden dürfen, folgt

$$\begin{aligned} \arctan(x) &= f(0) + \int_0^x f'(\tilde{x}) d\tilde{x} = 0 + \sum_{n=0}^{\infty} (-1)^n \int_0^x \tilde{x}^{2n} d\tilde{x} \\ &= \sum_{n=0}^{\infty} \frac{(-1)^n}{2n+1} x^{2n+1}. \end{aligned}$$

Nach dem Leibniz-Kriterium konvergiert die Potenzreihe auch für $x = \pm 1$, so dass insgesamt:

$$\arctan(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{2n+1} x^{2n+1} \quad \text{für } x \in [-1, 1]. \quad \square$$

Beispiel 9.32 (Area-Funktionen): Berechnung der **Taylor-Reihen der Area-Funktionen** am Entwicklungspunkt $x_0 = 0$ durch Zurückspielen auf die Binomische Reihe: Aus $f(x) = \operatorname{ar} \tanh(x)$ folgt

$$f'(x) = \operatorname{ar} \tanh'(x) = \frac{1}{1-x^2} = \sum_{n=0}^{\infty} x^{2n}.$$

Damit ist

$$f(x) = f(0) + \sum_{n=0}^{\infty} \frac{1}{2n+1} x^{2n+1}.$$

Da $f(0) = \operatorname{ar} \tanh(0) = 0$ ist

$$\operatorname{ar} \tanh(x) = \sum_{n=0}^{\infty} \frac{1}{2n+1} x^{2n+1} \quad \text{für } |x| < 1.$$

Auf analoge Weise werden die Taylor-Reihen von $\operatorname{ar} \sinh(x)$, $\operatorname{ar} \cosh(x)$ und $\operatorname{ar} \coth(x)$ berechnet, da $\operatorname{ar} \sinh'(x) = \frac{1}{\sqrt{1+x^2}}$ für $x \in \mathbb{R}$, $\operatorname{ar} \cosh'(x) = \frac{1}{\sqrt{x^2-1}}$ für $|x| > 1$ und $\operatorname{ar} \coth'(x) = \frac{1}{1-x^2}$ für $|x| > 1$. $\square$

Tabelle 9.1: Taylor-Reihen:

Funktion	Potenzreihenentwicklung	Konvergenz- bereich
$(1+x)^\alpha$	$\sum_{k=0}^{\infty} \binom{\alpha}{k} x^k$	$ x < 1$
$(1 \pm x)^{\frac{1}{2}}$	$1 \pm \frac{1}{2} x - \frac{1}{2 \cdot 4} x^2 \pm \frac{1 \cdot 3}{2 \cdot 4 \cdot 6} x^3 - \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6 \cdot 8} x^4 \pm \dots$	$ x \leq 1$
$(1 \pm x)^{-\frac{1}{2}}$	$1 \mp \frac{1}{2} x + \frac{1 \cdot 3}{2 \cdot 4} x^2 \mp \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6} x^3 + \frac{1 \cdot 3 \cdot 5 \cdot 7}{2 \cdot 4 \cdot 6 \cdot 8} x^4 \mp \dots$	$ x < 1$
$(1 \pm x)^{-1}$	$1 \mp x + x^2 \mp x^3 + x^4 \mp \dots$	$ x < 1$
$(1 \pm x)^{-2}$	$1 \mp 2x + 3x^2 \mp 4x^3 + 5x^4 \mp \dots$	$ x < 1$
$\sin x$	$x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \frac{x^9}{9!} - + \dots$	$ x < \infty$
$\cos x$	$1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \frac{x^6}{6!} + \frac{x^8}{8!} - + \dots$	$ x < \infty$
$\tan x$	$x + \frac{1}{3} x^3 + \frac{2}{15} x^5 + \frac{17}{315} x^7 + \frac{62}{2835} x^9 + \dots$	$ x < \frac{\pi}{2}$
e^x	$1 + \frac{x}{1!} + \frac{x^2}{2!} + \frac{x^3}{3!} + \frac{x^4}{4!} + \dots$	$ x < \infty$
$\ln x$	$(x-1) - \frac{1}{2} (x-1)^2 + \frac{1}{3} (x-1)^3 - + \dots$	$0 < x \leq 2$

Funktion	Potenzreihenentwicklung	Konvergenz- bereich
$\arcsin x$	$x + \frac{1}{2 \cdot 3} x^3 + \frac{1 \cdot 3}{2 \cdot 4 \cdot 5} x^5 + \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6 \cdot 7} x^7 + \dots$	$ x < 1$
$\arccos x$	$\frac{\pi}{2} - \left[x + \frac{1}{2 \cdot 3} x^3 + \frac{1 \cdot 3}{2 \cdot 4 \cdot 5} x^5 + \frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6 \cdot 7} x^7 + \dots \right]$	$ x < 1$
$\arctan x$	$x - \frac{1}{3} x^3 + \frac{1}{5} x^5 - \frac{1}{7} x^7 + \frac{1}{9} x^9 - + \dots$	$ x \leq 1$
$\sinh x$	$x + \frac{1}{3!} x^3 + \frac{1}{5!} x^5 + \frac{1}{7!} x^7 + \dots$	$ x < \infty$
$\cosh x$	$1 + \frac{1}{2!} x^2 + \frac{1}{4!} x^4 + \frac{1}{6!} x^6 + \dots$	$ x < \infty$
$\tanh x$	$x - \frac{1}{3} x^3 + \frac{2}{15} x^5 - \frac{17}{315} x^7 + \frac{62}{2835} x^9 - + \dots$	$ x < \frac{\pi}{2}$

 **Visualisierung der Konvergenz** Mit der Prozedur `taylor_poly` erhält man eine Animation, bei der das Taylor-Polynom mit steigendem n zusammen mit der Funktion $f(x)$ zu sehen ist. Man erkennt, dass mit wachsender Ordnung der Polynome eine gleichmäßige Anpassung an die Funktion erfolgt.

9.4 Anwendungen

⊙ Näherungspolynome einer Funktion

In vielen Anwendungen werden komplizierte Funktionen durch Taylor-Polynome $p_n(x)$ angenähert. Zum einen, damit man die Funktionen auf einfache Weise mit vorgegebener Genauigkeit auswerten kann, zum anderen, damit man z.B. bei linearer Näherung einen einfacheren physikalischen Zusammenhang erhält.

Der Fehler zwischen der Funktion $f(x)$ und dem Taylor-Polynom $p_n(x)$ ist nach dem Satz von Taylor gegeben durch das Lagrange Restglied

$$R_n(x) = \frac{1}{(n+1)!} f^{(n+1)}(\xi) (x - x_0)^{n+1},$$

wenn x_0 der Entwicklungspunkt und ξ ein nicht näher bekannter Zwischenwert zwischen x und x_0 . Für die meisten in der Praxis auftretenden Funktionen geht der Fehler gegen Null für $n \rightarrow \infty$. Bei hinreichend großem n wird also eine beliebig hohe Genauigkeit erzielt. In technischen Anwendungen werden Funktionen nahe ihrem Entwicklungspunkt oftmals nur durch das Taylor-Polynom $p_1(x)$ bzw. $p_2(x)$ ersetzt!

Beispiel 9.33 (Berechnung der Zahl e): Die Zahl e soll bis auf 6 Dezimalstellen genau berechnet werden. Dazu gehen wir von der Taylor-Entwicklung der Exponentialfunktion bei $x_0 = 0$ aus

$$e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n = 1 + x + \frac{1}{2!} x^2 + \dots + \frac{1}{n!} x^n + \dots$$

und berechnen e^1 durch das Taylor-Polynom der Ordnung n

$$e^1 \approx p_n(1) = 1 + 1 + \frac{1}{2!} 1^2 + \dots + \frac{1}{n!} 1^n.$$

Der Fehler nach dem Lagrangen Restglied ist

$$R_n(1) = \frac{1}{(n+1)!} e^\xi \leq \frac{1}{(n+1)!} e^1 < \frac{3}{(n+1)!}$$

(da $e^\xi \leq e^1 < 3$). Damit der Fehler kleiner als 6 Dezimalstellen wird, muss

$$R_n(1) < \frac{3}{(n+1)!} \stackrel{!}{<} 0.9 \cdot 10^{-6} \Rightarrow (n+1)! > \frac{3}{0.9 \cdot 10^{-6}} \approx 3\,333\,333.$$

Dies ist für $n \geq 9$ erfüllt, denn $(9 + 1)! = 3628800$. Für $n = 9$ ist e^1 bis auf 6 Dezimalstellen genau berechnet:

$$e^1 \approx \sum_{n=0}^9 \frac{1}{n!} = 2.7182815.$$

Vergleicht man diese Methode zur Berechnung der Zahl e mit der Folge $(1 + \frac{1}{n})^n$ aus Beispiel 6.3, so ist die Reihendarstellung sehr schnell konvergent. Es werden für eine Genauigkeit von 6 Dezimalstellen nur 9 Summenglieder benötigt im Vergleich zu $n > 10^5$ bei der Folgendarstellung. $\square$

Anwendungsbeispiel 9.34 (Relativistische Teilchen).

Nach A. Einstein beträgt die Gesamtenergie eines Teilchens

$$E = m c^2.$$

Dabei ist c die Lichtgeschwindigkeit und m die von der Geschwindigkeit des Teilchens v abhängige Masse:

$$m = \frac{m_0}{\sqrt{1 - (v/c)^2}},$$

m_0 ist dabei die Ruhemasse des Teilchens. Bezeichnet $E_0 = m_0 c^2$ die *Ruheenergie*, so ist die *kinetische Energie*

$$E_{kin} = E - E_0 = m c^2 - m_0 c^2 = m_0 c^2 \left(\frac{1}{\sqrt{1 - (v/c)^2}} - 1 \right).$$

Für ein nicht-relativistisches Teilchen ist $v \ll c$, d.h. $0 \approx \frac{v}{c} \ll 1$. $\frac{v}{c}$ ist also nahe dem Entwicklungspunkt $x_0 = 0$ der Funktion $\frac{1}{\sqrt{1-x^2}}$. Wir ersetzen daher nach Tabelle 9.1

$$\frac{1}{\sqrt{1-x}} \approx 1 + \frac{1}{2}x \quad \text{bzw.} \quad \frac{1}{\sqrt{1-(v/c)^2}} \approx 1 + \frac{1}{2}\left(\frac{v}{c}\right)^2.$$

Für die kinetische Energie gilt damit

$$E_{kin} = m_0 c^2 \left(\frac{1}{\sqrt{1-(v/c)^2}} - 1 \right) \approx m_0 c^2 \left(1 + \frac{1}{2}\left(\frac{v}{c}\right)^2 - 1 \right) = \frac{1}{2} m_0 v^2.$$

Der Term $\frac{1}{2} m_0 v^2$ repräsentiert die kinetische Energie eines Teilchens im Grenzfall $v \ll c$ (= *klassischer Grenzfall*). $\square$

Anwendungsbeispiel 9.35 (Scheinwerferregelung, mit MAPLE-Worksheet).

Kommen wir auf das Einführungsbeispiel der Scheinwerferregelung zurück. Um vom Quotienten der Distanzwerte d_1 und d_2 auf den aktuellen Neigungswinkel β zu schließen, müssen wir diesen Quotienten nach β auflösen. Dazu definieren wir die Funktion $q(\beta)$, die wir im Folgenden in eine Taylor-Reihe entwickeln.

$$q(\beta) := \frac{d_1}{d_2} = \frac{\sin(\alpha_2 + \beta)}{\sin(\alpha_1 + \beta)} \quad (*)$$

Gehen von den Parameterwerten $\beta_{ab} = 0.00999996$, $\alpha_1 = 0.20337$ und $\alpha_2 = 0.097913$ aus, ist der Quotient q_0 für den Winkel β_{ab} zwischen der Horizontalen und der Hell-Dunkel-Grenze beim ruhenden Fahrzeug

$$q_0 := q(\beta_{ab}) = 0.5086238522.$$

Um den Quotienten nach β von (*) aufzulösen, entwickeln wir nun die rechte Seite in eine Taylor-Reihe bis zur Ordnung 2.

$$q_2(\beta) := 0.4851497843 + 2.347500693\beta - 10.83456844(\beta - 0.00999996)^2$$

und lösen die Gleichung (*) für eine beliebige linke Seite $\frac{d_1}{d_2}$ mit der Näherung für die rechte Seite $\frac{\sin(\alpha_2 + \beta)}{\sin(\alpha_1 + \beta)} \sim q_2(\beta)$ nach β auf

$$\beta_{1/2} := 0.11833343 \pm 0.36918867 \sqrt{0.43052561 - 0.67716052 \frac{d_1}{d_2}}.$$

Von den beiden gefundenen Lösungen kommt nur diejenige in Frage, welche für die Größe $q_0 = \frac{d_1}{d_2}$ den richtigen Ablenkswinkel β_{ab} liefert. Dies ist die zweite Lösung β_2 .

Wir zeichnen die Näherungsfunktion gestrichelt und die ursprüngliche, implizit gegebene Funktion durchgezogen.

Aus der Grafik entnimmt man, dass die Näherungsformel für q zwischen 0.4 und 0.58 gut mit der impliziten Funktion übereinstimmt. Dies liefert einen Winkelbereich von -0.03 (-1.71°) bis 0.05 (2.864°), in dem die Näherung verwendet werden kann.

Um eine Näherungsformel zu erhalten, die auf die Berechnung von Wurzeln ganz verzichtet, entwickeln wir β_2 ebenfalls in eine Taylor-Reihe mit dem Ent-

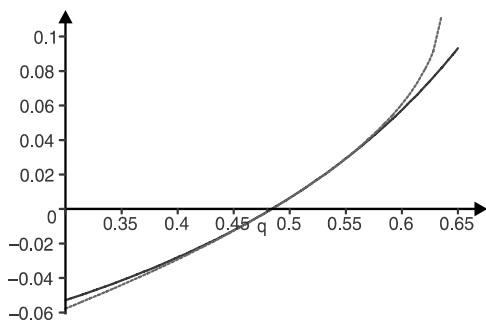


Abb. 9.4. Funktion und Näherung

wicklungspunkt $q = q_0$

$$\beta_2 = -2.373259209 + 23.42846966 q - 96.31243152 q^2 + 200.8851230 q^3 - 210.4285911 q^4 + 89.10928185 q^5$$

und zeichnen diese weitere Näherung in den obigen Graphen mit ein.

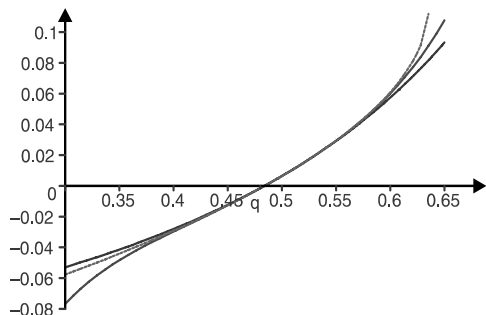


Abb. 9.5. Vergleich der Näherungen

Diese Funktion stellt im Winkelbereich zwischen -1° und 2° ebenfalls eine akzeptable Lösung dar. Der Vorteil dieser Näherungsformel besteht eben darin, dass auf die Berechnung von Wurzeln ganz verzichtet werden kann! Zur effizienten Berechnung stellen wir die Näherungsformel durch das Horner-Schema dar:

$$\beta_2 = -2.373259209 + (23.42846966 + (-96.31243152 + (200.8851230 + (-210.4285910 + 89.10928185 q) q) q) q) q. \quad \square$$

⊙ Integration durch Potenzreihenentwicklung

Potenzreihen und damit Taylor-Reihen dürfen in ihrem Konvergenzradius gliedweise differenziert bzw. integriert werden. Für

$$f(x) = \sum_{n=0}^{\infty} a_n (x - x_0)^n$$

gilt :

$$f'(x) = \frac{d}{dx} \sum_{n=0}^{\infty} a_n (x - x_0)^n = \sum_{n=0}^{\infty} a_n \frac{d}{dx} (x - x_0)^n = \sum_{n=1}^{\infty} n a_n (x - x_0)^{n-1} .$$

Man beachte, dass die Differenziation des konstanten Summanden a_0 Null ergibt und damit die abgeleitete Taylor-Reihe bei $n = 1$ beginnt.

$$\int f(x) dx = \int \sum_{n=0}^{\infty} a_n x^n dx = \sum_{n=0}^{\infty} a_n \int x^n dx = \sum_{n=0}^{\infty} \frac{a_n}{n+1} x^{n+1} + C.$$

Man beachte, dass beim bestimmten Integral die Integrationsgrenzen innerhalb des Konvergenzbereichs der Potenzreihe gelegen sein müssen.

Beispiel 9.36. Gesucht ist die Integralfunktion $F(x) = \int_0^x e^{-t^2} dt$, die nicht durch eine elementare Funktion darstellbar ist.

Mit dem Potenzreihenansatz

$$e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n$$

folgt

$$\begin{aligned} e^{-t^2} &= \sum_{n=0}^{\infty} \frac{1}{n!} (-t^2)^n = \sum_{n=0}^{\infty} \frac{1}{n!} (-1)^n t^{2n}. \\ \Rightarrow F(x) &= \int_0^x e^{-t^2} dt = \sum_{n=0}^{\infty} \frac{1}{n!} (-1)^n \int_0^x t^{2n} dt \\ &= \sum_{n=0}^{\infty} \frac{1}{n!} (-1)^n \frac{1}{2n+1} x^{2n+1} \quad (x \in \mathbb{R}). \quad \square \end{aligned}$$

⊙ Lösen von Differenzialgleichungen durch Potenzreihen

Eine in der Physik oftmals benutzte Methode zum Lösen von Differenzialgleichungen ist, die gesuchte Funktion in eine Potenzreihe zu entwickeln. Diese Potenzreihe enthält als unbekannte Größen die Koeffizienten a_n . Durch Einsetzen der Potenzreihe in die Differenzialgleichung werden über einen Koeffizientenvergleich die a_n bestimmt.

9.5 Komplexwertige Funktionen

Im Kapitel über komplexe Zahlen 5.1 benutzten wir die Eulersche Formel

$$e^{i\varphi} = \cos \varphi + i \sin \varphi \quad \varphi \in [0, 2\pi]$$

als Abkürzung. Wir zeigen in diesem Kapitel, dass diese Formel die Gleichheit der Funktion e^z und der Funktion $\cos(z) + i \sin(z)$ für beliebige komplexe Zahlen $z \in \mathbb{C}$ bedeutet.

Zunächst erklären wir e^z , $\cos z$ und $\sin z$ für $z \in \mathbb{C}$ als komplexe Funktionen $f: \mathbb{C} \rightarrow \mathbb{C}$ mit $z \mapsto f(z)$. Die Definition der Funktion muss dabei derart erfolgen, dass für $z \in \mathbb{R}$ die herkömmlichen reellen Funktionen als Spezialfall enthalten sind.

Im Komplexen stehen uns die Grundrechenoperationen $+$, $-$, $*$, $/$ zur Verfügung. Wir definieren daher komplexe Funktionen über diese Grundoperationen. Gerade aber die Exponential-, Sinus- und Kosinusfunktionen sind Standardbeispiele für die Darstellung einer Funktion durch ihre Taylor-Reihen. Da man bei der Auswertung einer Funktion über die Taylor-Reihe nur die oben genannten Grundoperationen benötigt, erklären wir die komplexen Funktionen e^z , $\sin(z)$ und $\cos(z)$ über ihre Taylor-Reihe. Zuvor geben wir jedoch die wichtigsten Ergebnisse für komplexe Potenzreihen an:

9.5.1 Komplexe Potenzreihen

Es übertragen sich alle Eigenschaften der reellen Potenzreihen sinngemäß auf den komplexen Fall. Bezüglich der Konvergenz einer komplexen Potenzreihe gilt:

Satz: Die komplexe Potenzreihe

$$\sum_{n=0}^{\infty} a_n (z - z_0)^n$$

mit $a_n \in \mathbb{C}$ und Entwicklungspunkt $z_0 \in \mathbb{C}$ hat als Majorante die **reelle** Potenzreihe $\sum_{n=0}^{\infty} |a_n| |z - z_0|^n$ und besitzt den Konvergenzradius

$$\rho = \lim_{n \rightarrow \infty} \frac{|a_n|}{|a_{n+1}|}.$$

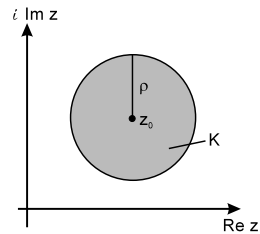
Der Konvergenzbereich ist $K = \{z \in \mathbb{C} : |z - z_0| < \rho\}$.

Begründung: Im Komplexen gelten die Rechenregeln $|z_1 + z_2| \leq |z_1| + |z_2|$ und $|a z| = |a| |z|$. Daher gilt die Abschätzung

$$\left| \sum_{n=0}^N a_n (z - z_0)^n \right| \leq \sum_{n=0}^N |a_n| |z - z_0|^n = \sum_{n=0}^N |a_n| |z - z_0|^n.$$

Somit ist $\sum_{n=0}^{\infty} |a_n| |z - z_0|^n$ eine Majorante von $\sum_{n=0}^{\infty} a_n (z - z_0)^n$. Die komplexe Potenzreihe besitzt damit den gleichen Konvergenzradius wie die reelle Majorante, nämlich $\rho = \lim_{n \rightarrow \infty} \frac{|a_n|}{|a_{n+1}|}$. $\square$

Interpretation: Erst im Komplexen erhält der Begriff *Konvergenzradius* seine volle Bedeutung, denn die Menge $K = \{z \in \mathbb{C} : |z - z_0| < \rho\}$ entspricht einem Kreis um z_0 mit Radius ρ . Innerhalb des Kreises konvergiert die Potenzreihe, außerhalb divergiert sie.



Beispiele 9.37. Aufgrund der Darstellung der Exponentialfunktion bzw. der trigonometrischen Funktionen über die Taylor-Reihe erhält man direkt die Definition der zugehörigen komplexwertigen Funktionen:

① **Komplexe Exponentialfunktion**

$$e^z := \sum_{n=0}^{\infty} \frac{1}{n!} z^n = 1 + \frac{1}{1!} z^1 + \frac{1}{2!} z^2 + \frac{1}{3!} z^3 + \dots \quad \text{für } z \in \mathbb{C}.$$

Wegen

$$\left| \sum_{n=0}^N \frac{1}{n!} z^n \right| \leq \sum_{n=0}^N \frac{1}{n!} |z|^n \leq \sum_{n=0}^{\infty} \frac{1}{n!} |z|^n = e^{|z|}$$

ist $e^{|z|}$ eine konvergente Majorante und e^z konvergiert für alle $z \in \mathbb{C}$.

② **Komplexe Sinusfunktion**

$$\sin(z) := \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} z^{2n+1} = z - \frac{1}{3!} z^3 + \frac{1}{5!} z^5 - \frac{1}{7!} z^7 \pm \dots \quad \text{für } z \in \mathbb{C}.$$

③ **Komplexe Kosinusfunktion**

$$\cos(z) := \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} z^{2n} = 1 - \frac{1}{2!} z^2 + \frac{1}{4!} z^4 - \frac{1}{6!} z^6 \pm \dots \quad \text{für } z \in \mathbb{C}.$$

Die absolute Konvergenz der Potenzreihen ② und ③ ist nach dem Majorantenkriterium für alle $z \in \mathbb{C}$ gesichert, denn die Majoranten sind die reellen Potenzreihen $\sum_{n=0}^{\infty} \frac{1}{(2n+1)!} |z|^{2n+1}$ und $\sum_{n=0}^{\infty} \frac{1}{(2n)!} |z|^{2n}$. $\square$

9.5.2 Die Eulersche Formel

Nach diesen Vorbemerkungen sind e^z , $\cos(z)$, $\sin(z)$ für jedes $z \in \mathbb{C}$ als unabhängige Funktionen definiert. Es gilt der Zusammenhang:

Satz: Für jedes $z \in \mathbb{C}$ gilt

$$e^{iz} = \cos(z) + i \sin(z) \quad (\text{Eulersche Formel})$$

Begründung: Wir stellen $\cos(z)$ und $i \sin(z)$ durch ihre Taylor-Reihen dar. Anschließend addieren wir die beiden Reihen und identifizieren die Summe als e^{iz} . Mit $i^0 = 1$, $i^1 = i$, $i^2 = -1$, $i^3 = -i$, $i^4 = 1$ und $i^5 = i$, $i^6 = -1$, $i^7 = -i$, $i^8 = 1$, usw. gilt:

$$\begin{array}{rcll}
 \cos(z) & = & 1 & -\frac{z^2}{2!} & +\frac{z^4}{4!} & -\frac{z^6}{6!} \pm \dots \\
 & = & 1 & +\frac{i^2 z^2}{2!} & +\frac{i^4 z^4}{4!} & +\frac{i^6 z^6}{6!} + \dots \\
 & = & 1 & +\frac{(iz)^2}{2!} & +\frac{(iz)^4}{4!} & +\frac{(iz)^6}{6!} + \dots \\
 \\
 i \sin(z) & = & i \frac{z}{1!} & -i \frac{z^3}{3!} & +i \frac{z^5}{5!} & \pm \dots \\
 & = & \frac{iz}{1!} & +\frac{i^3 z^3}{3!} & +\frac{i^5 z^5}{5!} & + \dots \\
 & = & \frac{(iz)}{1!} & +\frac{(iz)^3}{3!} & +\frac{(iz)^5}{5!} & + \dots \\
 \\
 \hline
 \cos(z) + i \sin(z) & = & 1 & +\frac{(iz)}{1!} & +\frac{(iz)^2}{2!} & +\frac{(iz)^3}{3!} & +\frac{(iz)^4}{4!} & +\frac{(iz)^5}{5!} & +\frac{(iz)^6}{6!} + \dots
 \end{array}$$

Folglich ist

$$\cos(z) + i \sin(z) = \sum_{n=0}^{\infty} \frac{1}{n!} (iz)^n = e^{iz} \quad \square$$

Mit dieser sehr einfachen Begründung ist die Eulersche Formel für alle $z \in \mathbb{C}$ bewiesen. Speziell für $z = \varphi \in \mathbb{R}$ gilt dann

$$e^{i\varphi} = \cos \varphi + i \sin \varphi \quad \varphi \in \mathbb{R}.$$

Wenn wir den Winkel φ durch $\varphi = \omega t$ ersetzen, so gilt die folgende Identität von Funktionen

$$e^{i\omega t} = \cos(\omega t) + i \sin(\omega t) \quad t \in \mathbb{R}.$$

Diese Gleichheit von Funktionen werden wir im folgenden Abschnitt ausnutzen, um $e^{i\omega t}$ zu differenzieren und zu integrieren.

► 9.5.3 Eigenschaften der komplexen Exponentialfunktion

E1

$$e^{z_1+z_2} = e^{z_1} \cdot e^{z_2}$$

$$z_1, z_2 \in \mathbb{C}.$$

Wie im Reellen hat auch im Komplexen das Additionstheorem für die Exponentialfunktion seine Gültigkeit. Der Beweis dieser zentralen Formel würde den Rahmen dieser Darstellung überschreiten. Festzuhalten ist die folgende Folgerung:

E2

$$e^{z+2\pi i} = e^z$$

$$z \in \mathbb{C}.$$

Denn setzt man in (E1) $z_2 = 2\pi i$, so folgt die Formel, da $e^{z_2} = e^{2\pi i} = 1$. Die komplexe Exponentialfunktion ist damit periodisch mit der komplexen Periode $2\pi i$. $\square$

E3

$$\sin(-z) = -\sin(z)$$

$$z \in \mathbb{C}$$

$$\cos(-z) = \cos(z)$$

$$z \in \mathbb{C}.$$

Wie im Reellen ist der Sinus eine ungerade Funktion, denn in der Definitionsgleichung für den Sinus treten nur ungerade Potenzen auf

$$\sin(-z) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} (-z)^{2n+1} = - \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} z^{2n+1} = -\sin(z).$$

Da per Definition $\cos(z)$ nur gerade Potenzen z^{2n} besitzt, ist $\cos(z)$ eine gerade Funktion. $\square$

E4

$$\cos(z) = \frac{1}{2} (e^{iz} + e^{-iz})$$

$$z \in \mathbb{C}$$

$$\sin(z) = \frac{1}{2i} (e^{iz} - e^{-iz})$$

$$z \in \mathbb{C}.$$

Anwendungen dieser beiden Identitäten werden wir im Kapitel über Differenzialgleichungen noch kennen lernen. Sie besagen, dass man die trigonometrischen Funktionen aus der komplexen Exponentialfunktion gewinnen

kann. Beide Identitäten sind Folgerungen aus der Eulerschen Formel, denn

$$e^{iz} = \cos(z) + i \sin(z) \quad (1)$$

$$\begin{aligned} e^{-iz} &= \cos(-z) + i \sin(-z) \\ &= \cos(z) - i \sin(z) \end{aligned} \quad (2)$$

Addiert man Gleichung (1) und (2), ist $e^{iz} + e^{-iz} = 2 \cos(z)$.

Subtrahiert man Gleichung (2) von (1), ist $e^{iz} - e^{-iz} = 2i \sin(z)$.

Durch Division der Faktoren, erhält man jeweils die Behauptung. $\square$

E5

$$\cos^2(z) + \sin^2(z) = 1 \quad z \in \mathbb{C}.$$

Man erhält (E5) aus (E4), indem man beide Gleichungen quadriert und dann addiert:

$$\begin{aligned} \cos^2(z) + \sin^2(z) &= \frac{1}{4} (e^{iz} + e^{-iz})^2 - \frac{1}{4} (e^{iz} - e^{-iz})^2 \\ &= \frac{1}{4} [(e^{iz})^2 + 2e^{iz}e^{-iz} + (e^{-iz})^2 \\ &\quad - (e^{iz})^2 + 2e^{iz}e^{-iz} - (e^{-iz})^2] \\ &= e^{iz}e^{-iz} = e^{i(z-z)} = e^0 = 1. \end{aligned} \quad \square$$

Anwendung: Additionstheoreme für Sinus und Kosinus

Für $\alpha, \beta \in \mathbb{C}$ (bzw. $\alpha, \beta \in \mathbb{R}$) gelten die **Additionstheoreme**

$$\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta \quad (A1)$$

$$\sin(\alpha + \beta) = \sin \alpha \cos \beta + \cos \alpha \sin \beta \quad (A2)$$

Begründung:

(A1): Aufgrund der Darstellung der Kosinus- und Sinusfunktion durch die komplexe Exponentialfunktion (E4) und dem Additionstheorem (E1) rechnet man:

$$\begin{aligned} \cos \alpha \cos \beta - \sin \alpha \sin \beta &= \frac{1}{2} (e^{i\alpha} + e^{-i\alpha}) \frac{1}{2} (e^{i\beta} + e^{-i\beta}) - \\ &\quad \frac{1}{2i} (e^{i\alpha} - e^{-i\alpha}) \frac{1}{2i} (e^{i\beta} - e^{-i\beta}) \\ &= \frac{1}{4} (e^{i(\alpha+\beta)} + e^{-i\alpha+i\beta} + e^{i\alpha-i\beta} + e^{-i(\alpha+\beta)}) + \\ &\quad + \frac{1}{4} (e^{i(\alpha+\beta)} - e^{-i\alpha+i\beta} - e^{i\alpha-i\beta} + e^{-i(\alpha+\beta)}) \\ &= \frac{1}{2} (e^{i(\alpha+\beta)} + e^{-i(\alpha+\beta)}) = \cos(\alpha + \beta). \end{aligned}$$

(A2): Analog zu (A1). $\square$

Folgerung: Verwandlung eines Produktes in eine Summe bzw. Differenz

Für $\alpha, \beta \in \mathbb{C}$ (bzw. $\alpha, \beta \in \mathbb{R}$) gelten die Formeln:

$$(1) \quad 2 \sin \alpha \sin \beta = \cos(\alpha - \beta) - \cos(\alpha + \beta)$$

$$(2) \quad 2 \cos \alpha \cos \beta = \cos(\alpha - \beta) + \cos(\alpha + \beta)$$

$$(3) \quad 2 \sin \alpha \cos \beta = \sin(\alpha - \beta) + \sin(\alpha + \beta)$$

Begründung: Übungsaufgabe. Man verwende die bereits bewiesenen Additionstheoreme (A1) und (A2).

9.5.4 Komplexe Hyperbelfunktionen

Definition: Für $z \in \mathbb{C}$ heißen die komplexen Funktionen

$$\cosh(z) := \frac{1}{2} (e^z + e^{-z}) \quad \text{Kosinus-Hyperbolikus}$$

$$\sinh(z) := \frac{1}{2} (e^z - e^{-z}) \quad \text{Sinus-Hyperbolikus}$$

Aufgrund der Definition der Hyperbelfunktionen und Eigenschaft (E4) gelten die folgenden Beziehungen

$$\mathbf{H1:} \quad \cos(iz) = \frac{1}{2} (e^{i(iz)} + e^{-i(iz)}) = \frac{1}{2} (e^{-z} + e^z) = \cosh(z),$$

$$\mathbf{H2:} \quad \sin(iz) = \frac{1}{2i} (e^{i(iz)} - e^{-i(iz)}) = \frac{1}{2i} (e^{-z} - e^z) = i \sinh(z).$$

Dies ist der Zusammenhang zwischen den Hyperbolikus-Funktionen und den trigonometrischen: $\cosh(z)$ und $\sinh(z)$ sind im Komplexen nichts anderes als die Kosinus- und Sinusfunktion mit dem Argument iz . Daher gelten auch die bis auf das Vorzeichen ähnlichen Formeln für beide Funktionstypen.

$$\mathbf{H3:} \quad \cosh^2(z) - \sinh^2(z) = 1 \quad z \in \mathbb{C}.$$

Gleichung (H3) erhält man durch Quadrieren von (H1) und (H2) und anschließender Addition, wenn Gleichung (E5) berücksichtigt wird. $\square$

9.5.5 Differenziation und Integration

Bei der Herleitung der komplexen Widerständen zur Berechnung von Wechselstromkreisen in Kap. 5.3.3 wurde die Funktion $e^{i\omega t}$ mit der Formel

$$(e^{i\omega t})' = i\omega e^{i\omega t}$$

nach t differenziert. Die imaginäre Einheit i wird als konstanter Faktor angesehen und die Funktion $e^{i\omega t}$ mit der Kettenregel nach t differenziert. dass diese Methode auch allgemein gilt, zeigt der folgende Satz.

Satz über die Differenziation komplexwertiger Funktionen. Seien $u, v : (a, b) \rightarrow \mathbb{R}$ reelle, differenzierbare Funktionen. Dann ist die komplexwertige Funktion $f := u + i v$ mit

$$f : (a, b) \rightarrow \mathbb{C}, \quad x \mapsto f(x) := u(x) + i v(x)$$

differenzierbar und es gilt

$$f'(x) = u'(x) + i v'(x).$$

Dieser Satz besagt, dass eine komplexwertige Funktion nach seiner reellen Variablen x differenziert wird, indem man die gewöhnliche Ableitung von Realteil und Imaginärteil bildet. **Beim Differenzieren komplexwertiger Funktionen dürfen alle Differenziationsregeln wie bei reellwertigen Funktionen benutzt werden.** Die Formel für die Ableitung folgt sofort aus der Definition der Ableitung, denn

$$\begin{aligned} f'(x) &= \lim_{h \rightarrow 0} \frac{1}{h} (f(x+h) - f(x)) \\ &= \lim_{h \rightarrow 0} \frac{1}{h} (u(x+h) + i v(x+h) - (u(x) + i v(x))) \\ &= \lim_{h \rightarrow 0} \left[\frac{1}{h} (u(x+h) - u(x)) + i \frac{1}{h} (v(x+h) - v(x)) \right] \\ &= \lim_{h \rightarrow 0} \frac{1}{h} (u(x+h) - u(x)) + i \lim_{h \rightarrow 0} \frac{1}{h} (v(x+h) - v(x)) \\ &= u'(x) + i v'(x). \quad \square \end{aligned}$$

Beispiele 9.38.

① Gesucht ist die Ableitung der Funktion $f(t) = e^{i\omega t}$. Wegen

$$e^{i\omega t} = \cos(\omega t) + i \sin(\omega t)$$

folgt

$$\begin{aligned}(e^{i\omega t})' &= \cos(\omega t)' + i \sin(\omega t)' \\ &= -\omega \sin(\omega t) + i \omega \cos(\omega t) = i\omega (\cos(\omega t) + i \sin(\omega t)) \\ &= i\omega e^{i\omega t}.\end{aligned}$$

Die komplexwertige Funktion $e^{i\omega t}$ darf wie die reellwertige Exponentialfunktion differenziert werden, wenn i als konstanter Faktor angesehen wird.

- ② Gesucht wird die Ableitung der Funktion $f(x) = x e^{ix}$. Mit der Produktregel folgt

$$f'(x) = e^{ix} + ix e^{ix} = (1 + ix) e^{ix}. \quad \square$$

Satz über die Integration komplexwertiger Funktionen. Seien $u, v : [a, b] \rightarrow \mathbb{R}$ reelle, integrierbare Funktionen. Dann ist die komplexwertige Funktion $f := u + i v$ mit

$$f : [a, b] \rightarrow \mathbb{C}, \quad x \mapsto f(x) := u(x) + i v(x)$$

integrierbar und es gilt

$$\int_a^b f(x) dx = \int_a^b u(x) dx + i \int_a^b v(x) dx.$$

Es gilt für die Integration einer komplexwertigen Funktion $f(x) = u(x) + i v(x)$, dass der Realteil und Imaginärteil integriert werden und anschließend das Integral von f sich aus beiden Teilen zusammensetzt. **Beim Integrieren komplexwertiger Funktionen dürfen alle Integrationsregeln wie bei reellwertigen Funktionen verwendet werden.** Die Formel ergibt sich analog zur Differenziationsformel.

Beispiele 9.39.

- ① Gesucht ist eine Stammfunktion von $f(x) = e^{ix}$.

$$\begin{aligned}\int f(x) dx &= \int (\cos x + i \sin x) dx = \int \cos x dx + i \int \sin x dx \\ &= \sin x + i (-\cos x) + C = -i (\cos x + i \sin x) + C \\ &= \frac{1}{i} e^{ix} + C.\end{aligned}$$

- ② Gesucht ist das unbestimmte Integral $\int x e^{ix} dx$.

Mit partieller Integration ($u = x$, $v' = e^{ix} \hookrightarrow u' = 1$, $v = -i e^{ix}$) folgt

$$\begin{aligned}\int x e^{ix} dx &= -i x e^{ix} + i \int e^{ix} dx \\ &= -i x e^{ix} + e^{ix} + C.\end{aligned}$$

Auch bei der Integration wird i wie eine Konstante behandelt. $\square$

Anwendungsbeispiel 9.40 (RC-Wechselstrom-Kreis).

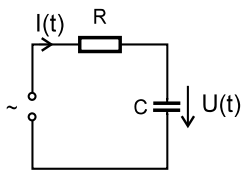


Abb. 9.6. RL-Kreis

Gegeben ist ein RC-Wechselstromkreis. Der Spannungsabfall am Kondensator ist

$$U(t) = \frac{1}{C} Q(t).$$

Da $I(t) = \frac{dQ(t)}{dt}$ ist $Q(t) = \int I(t) dt$. Der Spannungsabfall bei C lautet

$$U(t) = \frac{1}{C} \int I(t) dt.$$

Für einen komplexen Wechselstrom der Form

$$\hat{I}(t) = I_0 e^{i\omega t}$$

folgt für $\hat{U}(0) = 0$

$$\begin{aligned} \hat{U}(t) &= \frac{1}{C} \int I_0 e^{i\omega t} dt = \frac{1}{C} I_0 \frac{1}{i\omega} e^{i\omega t} \\ &= \frac{1}{i\omega C} I_0 e^{i\omega t} = \frac{1}{i\omega C} \hat{I}(t). \end{aligned}$$

Dies ist das komplexe Ohmsche Gesetz für den Kondensator, wenn als Widerstand

$$\hat{R}_C := \frac{\hat{U}(t)}{\hat{I}(t)} = \frac{1}{i\omega C}$$

gesetzt wird (vgl. Kap. 5.3.3).

□

MAPLE-Worksheets zu Kapitel 9

- Zahlenreihen mit MAPLE
- Die harmonische Reihe mit MAPLE
- Quotientenkriterium mit MAPLE
- Potenzreihen mit MAPLE
- Visualisierung der Konvergenz der Taylor-Reihen
- MAPLE-Prozedur zur Berechnung der Taylor-Polynome
- Scheinwerferregelung mit MAPLE
- Visualisierung der Eulerschen Formel
- Zusammenstellung der MAPLE-Befehle
- MAPLE-Lösungen zu den Aufgaben

9.6 Aufgaben zu Funktionenreihen

9.1 Man untersuche die folgenden Zahlenreihen auf Konvergenz

$$\begin{array}{llll} \text{a)} \sum_{n=1}^{\infty} \frac{1}{\sqrt{n}} & \text{b)} \sum_{n=1}^{\infty} n e^{-n^2} & \text{c)} \sum_{n=1}^{\infty} \frac{\sin n}{n^2} & \text{d)} \sum_{n=1}^{\infty} \frac{(-1)^n n}{2n+1} \\ \text{e)} \sum_{n=1}^{\infty} \frac{2^n}{n!} & \text{f)} \sum_{n=1}^{\infty} n \left(\frac{1}{2}\right)^{n-1} & \text{g)} \sum_{n=1}^{\infty} \frac{3^{2n}}{(2n)!} & \text{h)} \sum_{n=1}^{\infty} \frac{(-1)^{n+1} n}{5^{2n-1}} \\ \text{i)} \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{2n+1} & \text{j)} \sum_{n=1}^{\infty} \frac{1}{2^n n} & \text{k)} \sum_{n=1}^{\infty} 2^n \frac{1}{n} & \text{l)} \sum_{n=2}^{\infty} \frac{1}{(2n-1)(2n+1)} \end{array}$$

9.2 Untersuchen Sie die Konvergenz der folgenden Reihen und berechnen Sie -falls möglich- mit MAPLE ihren Wert.

$$\text{a)} \sum_{k=0}^{\infty} \frac{2^k + 3}{4^k} \quad \text{b)} \sum_{k=0}^{\infty} \frac{1}{k!} \quad \text{c)} \sum_{n=1}^{\infty} \frac{n^2}{2^n}$$

9.3 Man zeige die Divergenz der Reihe $\sum_{n=1}^{\infty} \frac{1 \cdot 3 \cdot 5 \cdot \dots \cdot (2n-1)}{6^n} (-1)^n$ und die Konvergenz der Reihe $\sum_{n=1}^{\infty} \frac{6^n}{(3^{n+1} - 2^{n+1})(3^n - 2^n)}$

9.4 Bestimmen Sie den Konvergenzradius der Reihen

$$\begin{array}{ll} \text{a)} \sum_{n=1}^{\infty} (-1)^{n+1} \frac{1}{5^n n} x^n & \text{b)} \sum_{n=1}^{\infty} \frac{\ln(n)}{n} x^n \\ \text{c)} \sum_{n=1}^{\infty} \left(1 + \frac{1}{n}\right)^{n^2} x^n & \text{d)} \sum_{n=1}^{\infty} \frac{(n!)^2}{(2n)!} x^{n+1} \end{array}$$

9.5 Berechnen Sie den Konvergenzradius von

$$\begin{array}{llll} \text{a)} \sum_{n=1}^{\infty} \frac{n x^n}{2^n} & \text{b)} \sum_{n=1}^{\infty} \frac{x^n}{n^2 + 1} & \text{c)} \sum_{n=1}^{\infty} n x^n & \text{d)} \sum_{n=1}^{\infty} (-1)^n \frac{x^n}{n} \\ \text{e)} \sum_{n=0}^{\infty} \frac{x^n}{2^n} & \text{f)} \sum_{n=1}^{\infty} \frac{n}{n+1} x^{n+1} & \text{g)} \sum_{n=1}^{\infty} \frac{n+1}{n!} x^n & \text{h)} \sum_{n=1}^{\infty} 2^n \cdot \frac{1}{n} x^n \end{array}$$

und diskutieren Sie den Konvergenzbereich K .

9.6 Bestimmen Sie den Konvergenzbereich der Potenzreihen

$$\begin{array}{ll} \text{a)} \sum_{n=1}^{\infty} n e^{-n} (x-4)^n & \text{b)} \sum_{i=0}^{\infty} \frac{i^3}{2^{i+1}} (x-1)^i \\ \text{c)} \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n} & \text{d)} \sum_{n=1}^{\infty} \frac{1}{n^n} (x-2)^n \end{array}$$

- 9.7 Zeigen Sie, dass die Taylor-Reihen von Sinus und Kosinus am Entwicklungspunkt $x_0 = 0$ gegeben sind durch

$$\sin x = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1} \quad , \quad \cos x = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n}.$$

Man bestimme den Konvergenzbereich der Potenzreihen.

- 9.8 Entwickeln Sie die Funktion

$$f(x) = \frac{1}{x^2} - \frac{2}{x} \quad , \quad x > 0,$$

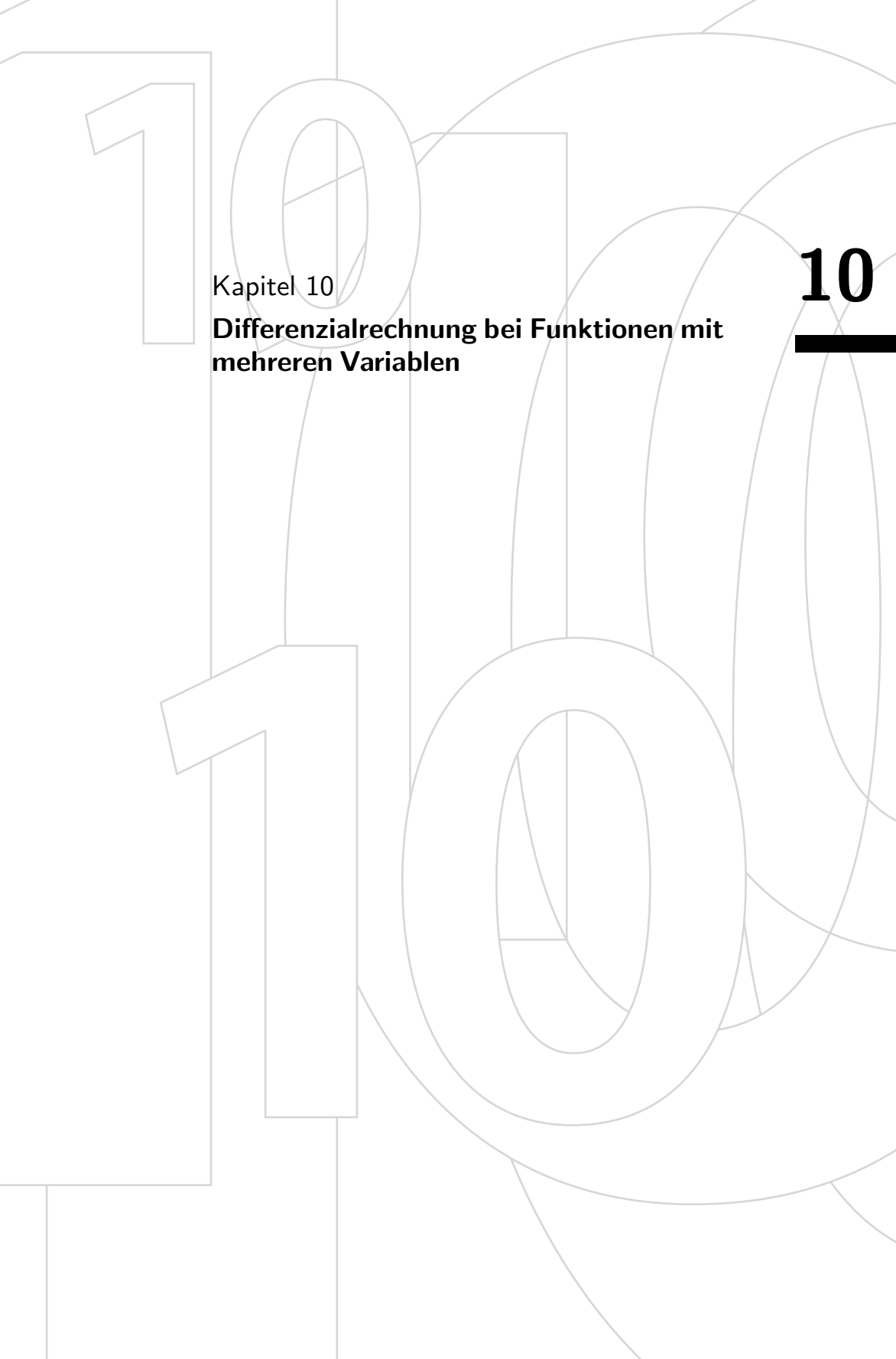
am Entwicklungspunkt $x_0 = 1$ in eine Taylor-Reihe. Geben Sie den zugehörigen Konvergenzbereich an.

- 9.9 Man berechne die Taylor-Reihe der Funktion $f(x) = \frac{1}{\sqrt{1+x}}$ an der Stelle $x_0 = 0$ und bestimme den Konvergenzbereich.
- 9.10 Berechnen Sie die Taylor-Reihen der Arkusfunktionen $\arcsin$, $\arccos$, arccot und bestimmen Sie den Konvergenzbereich.
- 9.11 Man berechne die Taylor-Reihe der Areafunktionen $\operatorname{ar sinh}$, $\operatorname{ar cosh}$, $\operatorname{ar coth}$ und bestimme den Konvergenzbereich.
- 9.12 Entwickeln Sie $f(x) = \cos x$ an der Stelle $x_0 = \frac{\pi}{3}$ in eine Taylor-Reihe und bestimmen Sie den Konvergenzbereich.
- 9.13 a) Erstellen Sie mit MAPLE eine Prozedur zur graphischen Darstellung der Taylor-Polynome einer Funktion, indem Sie den **taylor**-Befehl verwenden.
b) Bestimmen Sie damit die Taylor-Reihe von $y = |x|$ an der Stelle $x_0 = 0$ bis zur Ordnung 10.
c) Bestimmen Sie die Taylor-Reihe von $y = \sin x + \frac{1}{4} \sin 4x$ an der Stelle $x_0 = 0$. Wie groß muss die Ordnung gewählt werden, damit graphisch kein Unterschied zwischen Funktion und Taylor-Reihe im Bereich $[-\pi, \pi]$ erkennbar ist?
- 9.14 Die Funktion $f(x) = x e^{-x}$ soll in der Umgebung des Nullpunktes durch ein Polynom dritten Grades angenähert werden. Man bestimme mit der Taylorschen Reihenentwicklung diese Funktion.
- 9.15 Man berechne den Funktionswert von $f(x) = \sqrt{1-x}$ an der Stelle $x = 0.05$ auf sechs Dezimalstellen genau, wenn als Auswertepolynom ein Taylor-Reihenansatz mit Entwicklungspunkt $x_0 = 0$ gewählt wird.
- 9.16 Wie groß ist der maximale Fehler im Intervall $[0, \frac{1}{3}]$, wenn man die Funktion

$$f(x) = \frac{\sin x}{x}$$

um den Punkt $x_0 = 2$ bis zur Ordnung 2 entwickelt?

- 9.17 Berechnen Sie $\int_0^1 \frac{e^x - 1}{x} dx$ bis auf 3 Stellen genau.
- 9.18 Lösen Sie das unbestimmte Integral $F(x) = \int_0^x \frac{1}{1+t^2} dt$, indem der Integrand zunächst in eine Taylor-Reihe am Entwicklungspunkt $x_0 = 0$ entwickelt und anschließend gliedweise integriert wird.



Kapitel 10

**Differenzialrechnung bei Funktionen mit
mehreren Variablen**

10

10	Differenzialrechnung bei Funktionen mit mehreren Variablen	387
10.1	Funktionen mit mehreren Variablen	389
10.1.1	Einführung und Beispiele	389
10.1.2	Darstellung von Funktionen mit zwei Variablen	392
10.2	Stetigkeit	398
10.3	Differenzialrechnung	400
10.3.1	Partielle Ableitung	400
10.3.2	Totale Differenzierbarkeit	408
10.3.3	Gradient und Richtungsableitung	410
10.3.4	Der Taylorsche Satz	416
10.3.5	Kettenregeln	cd
10.4	Anwendungen der Differenzialrechnung	423
10.4.1	Das Differenzial als lineare Näherung	423
10.4.2	Fehlerrechnung	428
10.4.3	Lokale Extrema bei Funktionen mit mehreren Variablen	432
10.4.4	Ausgleichen von Messfehlern; Regressionsgerade	439
10.5	Aufgaben zur Differenzialrechnung	446
 Zusätzliche Abschnitte auf der CD-Rom		
10.6	MAPLE: Funktionen in mehreren Variablen	cd
10.6.1	Darstellung von Funktionen in zwei Variablen	cd
10.6.2	Differenzialrechnung	cd
10.6.3	Anwendung der Differenzialrechnung	cd

10 Differenzialrechnung bei Funktionen mit mehreren Variablen

Das Kapitel über die Funktionen von mehreren Variablen besteht aus vier Abschnitten. In 10.1 werden die Funktionen mit mehreren Variablen eingeführt und die graphische Darstellung von Funktionen mit zwei Variablen angegeben. Der Begriff der Stetigkeit wird in 10.2 verallgemeinert. Der mathematisch wichtigste Abschnitt ist 10.3, der die Differenzialrechnung zur Charakterisierung und Beschreibung dieser Funktionen behandelt. Die Konstruktion der Ableitung wird verallgemeinert und neue Begriffe wie die partielle Ableitung, die totale Differenzierbarkeit, der Gradient und die Richtungsableitung eingeführt. Der Taylorsche Satz liefert den Übergang zu den Anwendungen in 10.4, bei denen die Diskussion der Fehlerrechnung, lokale Extremwertbestimmungen und die Ausgleichsrechnung im Vordergrund stehen.

Hinweis: Auf der CD-Rom befinden sich ein zusätzliche Abschnitte über [Kettenregeln](#) und [Koordinatentransformationen](#).

10.1 Funktionen mit mehreren Variablen

10.1

In diesem Abschnitt werden Funktionen mit mehreren Variablen anhand physikalischer Problemstellungen eingeführt. Zentral dabei sind die unterschiedlichen Darstellungen von Funktionen mit zwei Variablen.

Hinweis: Auf der CD-Rom befindet sich viele [MAPLE-Worksheets](#), die Funktionen mit zwei Variablen dreidimensional visualisieren.

► 10.1.1 Einführung und Beispiele

Eine Funktion f **einer** reellen Variablen x besteht aus dem Definitions- und Zielbereich sowie der eindeutigen Funktionszuordnung:

$$f : \mathbb{D} \rightarrow \mathbb{R} \quad \text{mit} \quad x \mapsto y = f(x) .$$

Die Zuordnung erfolgt üblicherweise mit Hilfe einer Vorschrift $y = f(x)$; die Funktionen können dann in der Regel als Schaubild (= Graph der Funktion) dargestellt werden.

Beispiele 10.1 (Funktionen einer Variablen):

- ① $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = ax + b$ (Geradengleichung).
- ② $f : \mathbb{R}_{>0} \rightarrow \mathbb{R}$ mit $f(x) = \ln x \cdot \cos(x^2 - 1)$.
- ③ Potenzial in einem ebenen Plattenkondensator

$$\Phi : [0, d] \rightarrow \mathbb{R} \quad \text{mit} \quad \Phi(x) = \Delta\Phi \frac{x}{d} ,$$

wenn d der Plattenabstand und $\Delta\Phi = \Phi_2 - \Phi_1$ die Potentialdifferenz ist. □

Viele in den Naturwissenschaften auftretenden Zusammenhänge sind aber komplizierter und lassen sich nicht durch eine Funktion mit **einer** Variablen beschreiben. Die meisten physikalischen Gesetze stellen Beziehungen zwischen **mehreren** unabhängigen Größen dar.

Beispiele 10.2 (Funktionen mit zwei Variablen):

① Für ein ideales Gas gilt die **Zustandsgleichung**

$$p = R \cdot \frac{T}{V} .$$

Der Druck p hängt sowohl von der Temperatur T als auch von dem Gasvolumen V ab. R ist die universelle Gaskonstante. Jedem Wertepaar (T, V) wird durch diese Formel ein Druckwert $p(T, V)$ zugeordnet:

$$(T, V) \mapsto p(T, V) = R \cdot \frac{T}{V} .$$

Neben der Angabe der Zuordnungsvorschrift gehört noch die des Definitionsbereichs. Physikalisch sinnvoll ist $T > 0, V > 0$.

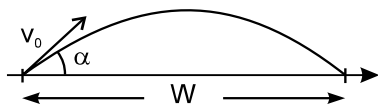


Abb. 10.1. Wurfweite W

② Die **Wurfweite** W beim schiefen Wurf bestimmt sich über die Beziehung

$$W = \frac{v_0^2}{g} \sin(2\alpha) ,$$

wenn v_0 die Anfangsgeschwindigkeit, α der Wurfwinkel und g die konstante Erdbeschleunigung ist. Die Wurfweite hängt also von v_0 und α ab; jedem Zahlenpaar (v_0, α) wird eindeutig eine Weite $W(v_0, \alpha)$ zugeordnet:

$$(v_0, \alpha) \mapsto W(v_0, \alpha) = \frac{1}{g} v_0^2 \sin(2\alpha)$$

mit $v > 0$ und $0 < \alpha \leq 90^\circ$.

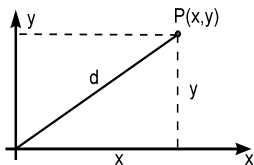


Abb. 10.2. Abstand d

③ Der **Abstand** d eines Punktes $P(x, y)$ vom Ursprung beträgt in der Ebene nach dem Satz von Pythagoras

$$d = \sqrt{x^2 + y^2} .$$

Jedem Punkt (x, y) wird genau ein Funktionswert (der Abstand) $d(x, y)$ zugeordnet

$$(x, y) \mapsto d(x, y) = \sqrt{x^2 + y^2} .$$

□

Beispiele 10.3 (Funktionen mit mehr als zwei Variablen):

① Der **Abstand** d zweier Punkte $P_1(x_1, y_1, z_1)$ und $P_2(x_2, y_2, z_2)$ beträgt im dreidimensionalen Raum

$$d = \left| \overrightarrow{P_1 P_2} \right| = \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2 + (z_1 - z_2)^2}.$$

Der Abstand d ist eine Funktion der 6 Variablen $x_1, x_2, y_1, y_2, z_1, z_2$:

$$(x_1, x_2, y_1, y_2, z_1, z_2) \mapsto d(x_1, \dots, z_2) = \sqrt{(x_1 - x_2)^2 + \dots + (z_1 - z_2)^2}.$$

② Für eine **Reihenschaltung** von n Ohmschen Widerständen $R_1, \dots, R_n$ berechnet sich der Gesamtwiderstand über

$$R = R_1 + R_2 + \dots + R_n.$$

Der Gesamtwiderstand ist eine Funktion der n Einzelwiderstände

$$(R_1, R_2, \dots, R_n) \mapsto R(R_1, \dots, R_n) = R_1 + R_2 + \dots + R_n. \quad \square$$

Definition: Eine reelle Funktion f von n reellen Variablen $x_1, \dots, x_n$ ist eine Abbildung, die jedem $(x_1, \dots, x_n) \in \mathbb{ID}$ genau einen Wert in $\mathbb{R}$ zuordnet:

$$f : \mathbb{ID} \rightarrow \mathbb{R} \quad \text{mit} \quad (x_1, \dots, x_n) \mapsto y = f(x_1, \dots, x_n).$$

Der Definitionsbereich $\mathbb{ID}$ ist dabei eine Menge von n -Tupeln $(x_1, \dots, x_n) \in \mathbb{R}^n$ reeller Zahlen, die in den Funktionsausdruck eingesetzt werden.

Die ausführliche Bezeichnung sowie die Angabe des Definitionsbereichs ist in der Praxis recht schwerfällig, so dass man in den Anwendungen stattdessen etwas nachlässig einfach von der Funktion $f(x_1, \dots, x_n)$ spricht.

Wir werden in diesem Abschnitt hauptsächlich Funktionen mit zwei Variablen behandeln. Viele Eigenschaften von Funktionen mehrerer Variablen können hier bereits verdeutlicht werden. Funktionen von zwei Variablen haben den Vorteil, dass sie sich graphisch darstellen lassen; Funktionen mit mehr als zwei Variablen nicht mehr! Im Folgenden sei daher

$$z = f(x, y), \quad (x, y) \in \mathbb{ID}$$

eine reellwertige Funktion der zwei Variablen x und y , die auf einem Gebiet $\mathbb{ID} \subset \mathbb{R}^2$ definiert ist.

Beispiel 10.4. In Abb. 10.3 sind einige Beispiele für zweidimensionale Gebiete angegeben.

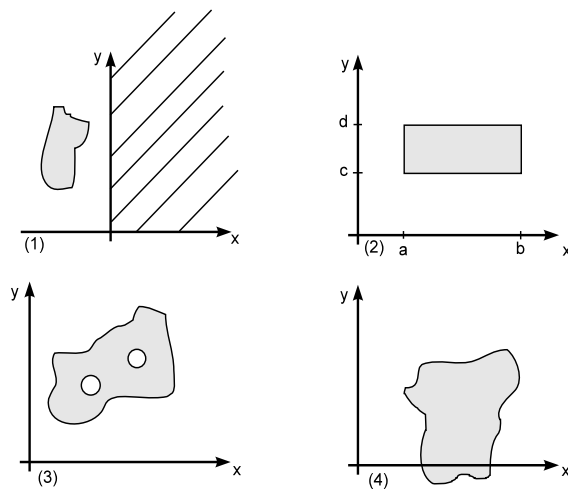


Abb. 10.3. Beispiele für zweidimensionale Gebiete

Fall (1) (*nicht zusammenhängendes Gebiet*) ist für die Anwendungen weniger wichtig; Fall (3) heißt *zusammenhängend*; Fälle (2) und (4) heißen *einfach zusammenhängend*. $\square$

➤ 10.1.2 Darstellung von Funktionen mit zwei Variablen

Zur Veranschaulichung von Funktionen mit zwei Variablen haben sich im Wesentlichen drei Darstellungsarten bewährt:

(1) Der Graph. Unter dem *Graphen* von f versteht man die Menge der Punkte $(x, y, f(x, y))$ für die (x, y) aus dem Definitionsbereich von f sind. In der Regel ist ein Graph eine gekrümmte Fläche im dreidimensionalen Raum

$$\Gamma_f := \{(x, y, z) \in \mathbb{R}^3 : z = f(x, y) \text{ und } (x, y) \in \mathbb{D}\}.$$

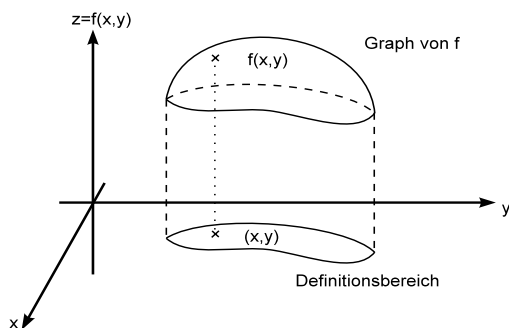


Abb. 10.4. Der Graph einer zweidimensionalen Funktion

Beispiel 10.5 (Mexikanischer Hut, mit MAPLE-Worksheet). Gegeben ist die Funktion

$$f(r) = \frac{\sin(r)}{r}.$$

Ersetzt man $r = \sqrt{x^2 + y^2}$, erhält man eine Funktion von zwei Variablen x und y :

$$f(x, y) = \frac{\sin\left(\sqrt{x^2 + y^2}\right)}{\sqrt{x^2 + y^2}}.$$

Die Darstellung der Funktion in Form eines dreidimensionalen Graphen erhält man mit MAPLE zu

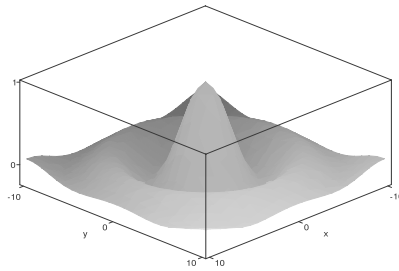


Abb. 10.5. Mexikanischer Hut

(2) Höhenlinien. Eine Funktion f kann auch durch ihre Höhenlinien (Niveaulinien) graphisch dargestellt werden. Die *Höhenlinie* von f zur Höhe c ist die Menge der Punkte $(x, y) \in \mathbb{D}$, welche die implizite Gleichung

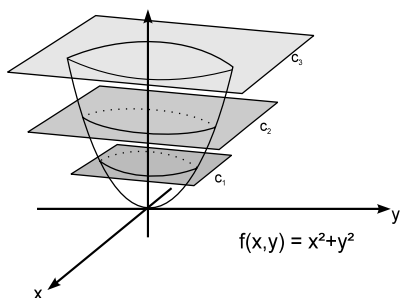
$$f(x, y) = c$$

erfüllen. Höhenlinien sind Schnitte des Graphen $(x, y, f(x, y))$ mit Ebenen parallel zur (x, y) -Ebene mit Achsenabschnitt c .

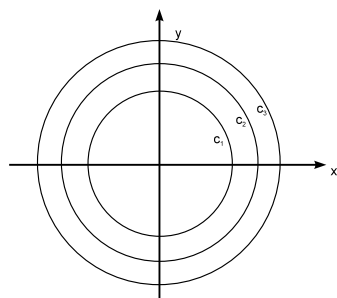
Beispiel 10.6 (Mit MAPLE-Worksheet). Gegeben ist die Funktion

$$f(x, y) = x^2 + y^2.$$

In Abb. 10.6 (links) ist der Graph der Funktion in Form einer dreidimensionalen Darstellung gezeigt. Die Höhenlinien werden als Schnitte des Graphen mit Ebenen parallel zur (x, y) -Ebene eingezeichnet. In Abb. 10.6 (rechts) sind die Höhenlinien in die (x, y) -Ebene projiziert.



Höhenlinien = Schnitte des Graphen
mit Ebenen parallel zur (x, y) -Ebene



Höhenlinien projiziert
in die (x, y) -Ebene

Abb. 10.6. Graph einer Funktion $f(x, y)$ und Höhenlinien

Zur graphischen Charakterisierung der Funktion wählt man also mehrere Höhen und zeichnet die markierten Niveaulinien in der (x, y) -Ebene. Anwendungsbeispiele sind Kurven gleichen Luftdrucks auf der Wetterkarte (*Isobare*), die Höhenlinien auf der Landkarte oder die Linien gleichen Potentials (*Äquipotenziallinien*) bei der Beschreibung elektrostatischer Felder. $\square$

Anwendungsbeispiel 10.7 (Elektrostatisches Potenzial, mit [MAPLE-Worksheet](#)).

Das **elektrostatische Potenzial** einer im Ursprung befindlichen elektrischen Ladung q ist im Abstand r bestimmt durch

$$\Phi(r) = \frac{1}{4\pi\epsilon_0} \frac{q}{r} \quad (*)$$

mit der Dielektrizitätskonstante $\epsilon_0 = 8.8 \cdot 10^{-12} \frac{F}{m}$.

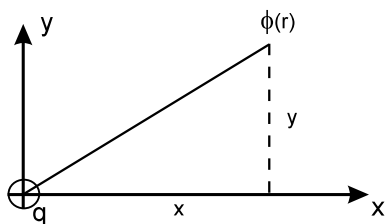


Abb. 10.7. Ladung im Ursprung

$$\Phi(x, y) = \frac{1}{4\pi\epsilon_0} \frac{q}{\sqrt{x^2 + y^2}}.$$

(1) Gesucht ist eine dreidimensionale Darstellung des Potenzialverlaufs in der (x, y) -Ebene sowie 20 Äquipotenziallinien für eine Punktladung $q = e = 1.6 \cdot 10^{-19} C$. Um eine dreidimensionale Darstellung zu erhalten, setzt man $r = \sqrt{x^2 + y^2}$ in die Potenzialformel $(*)$ ein:

Dann ist $\Phi(x, y)$ eine Funktion der zwei Variablen x und y . Für $(x, y) \rightarrow (0, 0)$ wächst das Potenzial über alle Grenzen hinweg; es wird *singulär*. Damit der funktionale Verlauf aus dem Graphen erkenntlich wird, schränkt man den dar-

zustellenden Wertebereich ein.

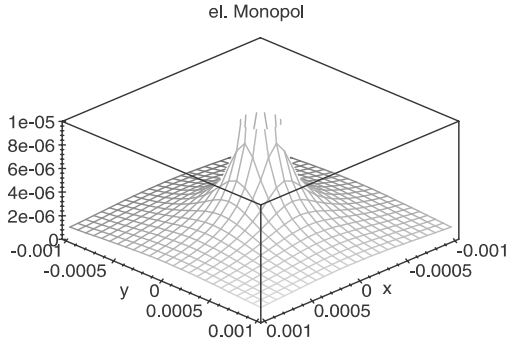


Abb. 10.8. Elektrischer Monopol

Fügt man zusätzlich Linien gleichen Potentials in den Graphen ein, erhält man die folgende graphische Darstellung:

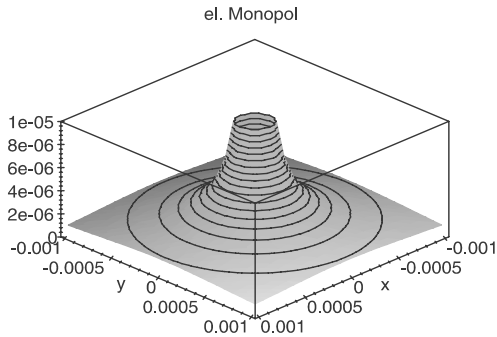


Abb. 10.9. Elektrischer Monopol mit Äquipotenziallinien

(2) Gesucht ist der Potenzialverlauf sowie die Höhenliniendarstellung eines elektrischen **Quadrupols**, wenn die Ladungen in den Ecken eines Quadrats mit Kantenlänge L angeordnet sind.

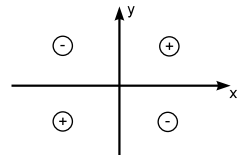
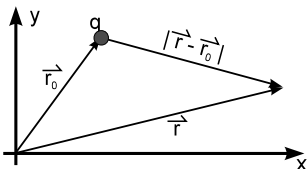


Abb. 10.10.
Quadrupol

Eine Ladung q , die bei $\vec{r}_0 = (x_0, y_0)$ lokalisiert ist, induziert am Ort $\vec{r} = (x, y)$ ein Potenzial gemäß



$$\Phi(x, y) = \frac{1}{4\pi\epsilon_0} \frac{q}{\sqrt{(x-x_0)^2 + (y-y_0)^2}}.$$

Für das Potenzial mehrerer Punktladungen $q_1, \dots, q_n$ gilt das Superpositionsprinzip

$$\Phi(x, y) = \Phi_1(x, y) + \dots + \Phi_n(x, y) .$$

Folglich ist das Potenzial des elektrischen Quadrupols am Ort (x, y) , wenn die Ladungen sich bei $(\frac{L}{2}, \frac{L}{2})$, $(\frac{L}{2}, -\frac{L}{2})$, $(-\frac{L}{2}, \frac{L}{2})$, $(-\frac{L}{2}, -\frac{L}{2})$ befinden

$$\begin{aligned} \Phi(x, y) = \frac{1}{4\pi\epsilon_0} & \left\{ \frac{q}{\sqrt{(x - \frac{L}{2})^2 + (y - \frac{L}{2})^2}} + \frac{-q}{\sqrt{(x - \frac{L}{2})^2 + (y + \frac{L}{2})^2}} \right. \\ & \left. + \frac{-q}{\sqrt{(x + \frac{L}{2})^2 + (y - \frac{L}{2})^2}} + \frac{q}{\sqrt{(x + \frac{L}{2})^2 + (y + \frac{L}{2})^2}} \right\}. \end{aligned}$$

Die graphische Darstellung des Quadrupols erfolgt entweder in Form einer dreidimensionalen Darstellung in Abb. 10.11

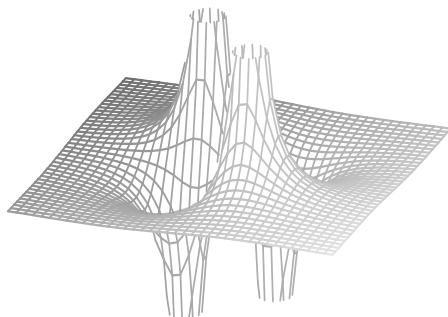


Abb. 10.11. Potenzialverlauf beim elektrischen Quadrupol

oder in Form von Äquipotenziallinien in Abb. 10.12.

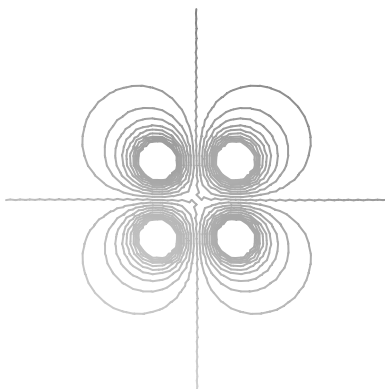


Abb. 10.12. Äquipotenziallinien beim elektrischen Quadrupol

(3) Schnittkurvendiagramme. Höhenlinien sind die Schnitte des Graphen von $z = f(x, y)$ mit Ebenen parallel zur (x, y) -Ebene. Wählt man stattdessen eine Ebene parallel zur (x, z) - oder (y, z) -Ebene, kommt man zu den sog. *Schnittkurvendiagrammen*

$$\begin{aligned} z &= f(x = c, y) && \text{Schnittebene parallel zu } (y, z), \\ z &= f(x, y = c) && \text{Schnittebene parallel zu } (x, z). \end{aligned}$$

Anwendung findet diese Darstellung in den Kennlinienbildern, wie das folgende Beispiel verdeutlichen soll.

Anwendungsbeispiel 10.8 (Zustandsgleichung idealer Gase).

Gegeben ist die Zustandsgleichung für ideale Gase

$$p = R \frac{T}{V}.$$

Hierbei ist p der Druck, T die Temperatur, V das Volumen und R die universelle Gaskonstante. Indem man für die Temperatur T konstante Werte einsetzt, $T_1 < T_2 < T_3 < T_4 < T_5$, erhält man den Druck als Funktion des Gasvolumens V . Man bezeichnet die Kurven gleicher Temperatur auch als *Isotherme*.

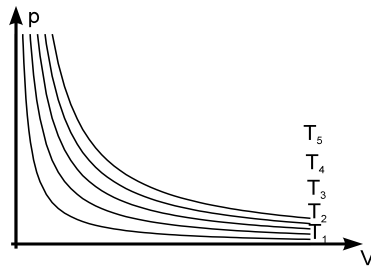


Abb. 10.13. Darstellung des Drucks über Isotherme



Auf der CD-Rom befindet sich die Prozedur [Funktion2d](#). Diese Prozedur stellt eine Funktion von zwei Variablen graphisch als Animation unter unterschiedlichen Blickwinkeln dar. Zunächst erfolgt die Darstellung als dreidimensionales, farbiges Schaubild, anschließend werden die Höhenlinien eingeblendet und nur noch Grauschattierungen der Funktion dargestellt.

10.2 Stetigkeit

Die *Stetigkeit* einer Funktion $f(x)$ bedeutet unpräzise gesprochen, dass der Graph von f keine Sprünge aufweist. In diesem Sinne bezeichnet man auch eine Funktion mit zwei Variablen als stetig.

Die Präzisierung des Stetigkeitsbegriffs im Punkte x_0 ist, dass der linksseitige und rechtsseitige Funktionsgrenzwert mit dem Funktionswert $f(x_0)$ übereinstimmt. D.h. unabhängig ob man sich von links oder von rechts an x_0 nähert, es kommt immer der selbe Funktionswert $f(x_0)$ heraus: Für jede Folge $x_n \rightarrow x_0$ konvergiert $f(x_n) \rightarrow f(x_0)$. Übertragen auf Funktionen mit zwei Variablen liefert dies die folgende Definition.

Definition: (Stetigkeit). Die Funktion f heißt im Punkte $(x_0, y_0) \in \mathbb{ID}$ **stetig**, wenn für **jede** Folge von Punkten $(x_n, y_n) \in \mathbb{ID}$ mit $x_n \rightarrow x_0$ und $y_n \rightarrow y_0$ gilt

$$f(x_n, y_n) \xrightarrow{n \rightarrow \infty} f(x_0, y_0) .$$

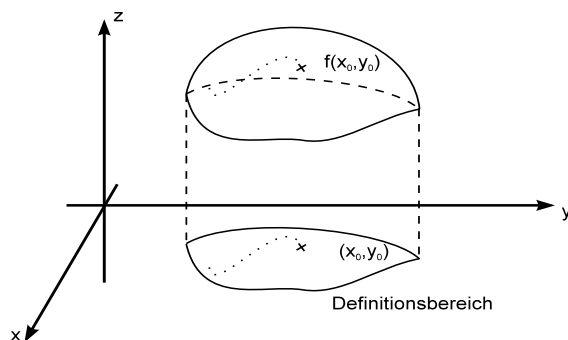


Abb. 10.14. Zur Stetigkeit einer Funktion

Bemerkungen:

- (1) Statt $f(x_n, y_n) \xrightarrow{n \rightarrow \infty} f(x_0, y_0)$ schreibt man auch $\lim_{n \rightarrow \infty} f(x_n, y_n) = f(x_0, y_0)$.
- (2) Im Folgenden schreiben wir für $x_n \rightarrow x$ und $y_n \rightarrow y$ kurz $(x_n, y_n) \rightarrow (x, y)$.
- (3)  **Achtung:** Für die Stetigkeit der Funktion f im Punkte $(x_0, y_0) \in \mathbb{ID}$ genügt es nicht, dass für **eine** spezielle Folge von Punkten $(x_n, y_n) \rightarrow (x_0, y_0)$ die Funktionsfolge $f(x_n, y_n)$ konvergiert, sondern die Betonung liegt auf **jede** Folge.
- (4) Ist f stetig für alle Punkte des Definitionsbereichs, so heißt f *stetig in ID*.

- (5) **Geometrische Interpretation:** f ist im Punkte (x_0, y_0) stetig, wenn für jede Folge aus dem Definitionsbereich, die gegen (x_0, y_0) konvergiert, die Funktionsfolgen immer gegen den Wert $f(x_0, y_0)$ streben (siehe Abb. 10.14).

Beispiele 10.9:

- ① Stetig sind z.B. alle Polynome in mehreren Variablen

$$f(x, y) = 2 - xy + 3x^2y + 5x^9 - x^2y^3$$

$$g(x, y, z) = x^3 - 6xz - 3yz + 4x^2y^3z^4.$$

- ② Folgende Funktionen sind für alle $(x, y) \in \mathbb{R}^2$ stetig

$$e^{2x-3y}, \quad \ln(1 + x^4 + x^2y^2), \quad \sin(x^2 + y^4), \quad \sqrt{x^2 + y^2}.$$

- ③ Auch rationale Funktionen, die als Quotient von Polynomen definiert sind, stellen stetige Funktionen in allen Punkten dar, in denen das Nennerpolynom nicht verschwindet:

$$F(x, y) = \frac{x^3 - 2y^3 + xy}{x^2 - y^2}, \quad G(x, y, z) = \frac{2xy^2 + 4y^2z}{x^2 + y^2 + z^2}.$$

Die Funktion F ist für alle Punkte aus $\mathbb{R}^2$ stetig, die nicht auf der Geraden $y = x$ oder $y = -x$ liegen. G ist für alle $(x, y, z) \in \mathbb{R}^3$ mit Ausnahme des Nullpunktes stetig.

- ④  **Achtung:** Die Funktion

$$f(x, y) = \frac{2xy}{x^2 + y^2}; \quad (x, y) \neq (0, 0)$$

ist außerhalb $(0, 0)$ stetig. f ist aber **nicht** in $(0, 0)$ stetig fortsetzbar, denn wählen wir als Folge im Definitionsbereich $(x_n, y_n) = (\frac{1}{n}, \frac{a}{n}) \rightarrow (0, 0)$ gilt für die Funktionsfolge

$$f(x_n, y_n) = \frac{2 \frac{1}{n} \frac{a}{n}}{\frac{1}{n^2} + \frac{a^2}{n^2}} = \frac{2a}{1 + a^2}.$$

Der Funktionsgrenzwert im Punkte $(0, 0)$ ist **abhängig** von der gewählten Folge $(x_n, y_n) \rightarrow (0, 0)$. Somit ist f dort nicht stetig fortsetzbar. $\square$

Bemerkung: Die Stetigkeit einer Funktion f von zwei Variablen in einem festen Punkt (x_0, y_0) bedeutet anschaulich gesprochen, dass der Funktionswert $f(x, y)$ beliebig nahe beim Funktionswert $f(x_0, y_0)$ liegt, wenn nur der Punkt (x, y) genügend nahe beim Punkt (x_0, y_0) liegt. Diese anschauliche Vorstellung lässt sich durch die folgende Definition präzisieren:

δ - ε -Stetigkeit einer Funktion: Die Funktion f ist im Punkte $(x_0, y_0) \in \mathbb{D}$ **stetig**, wenn es zu jeder (beliebig kleinen) Zahl $\varepsilon > 0$ eine Zahl $\delta > 0$ gibt mit der Eigenschaft: $|f(x, y) - f(x_0, y_0)| < \varepsilon$ für alle Punkte $(x, y) \in \mathbb{D}$ für die $|x - x_0| < \delta$ und $|y - y_0| < \delta$.

10.3

10.3 Differenzialrechnung

Der wichtigste Begriff bei der Differenzialrechnung von Funktionen in mehreren Variablen ist die partielle Ableitung. Rechentechnisch benötigt man die partielle Ableitung bei allen weiteren Begriffsbildungen wie z.B. der totalen Differenzierbarkeit, beim Gradient und der Richtungsableitung oder dem Taylorschen Satz.

Hinweis: Auf der CD-Rom befindet sich ein zusätzlicher Abschnitt über [Kettenregeln](#) beim Differenzieren von Funktionen mit mehreren Variablen. Viele [MAPLE-Prozeduren](#) verdeutlichen die Begriffsbildungen in diesem Abschnitt.

10.3.1 Partielle Ableitung

Bei Funktionen einer Variablen spielt der Ableitungsbegriff eine zentrale Rolle: Die Ableitung der Funktion f im Punkte x_0 ist definiert als der Grenzwert

$$f'(x_0) = \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x}.$$

Aus geometrischer Sicht ist die Ableitung der Funktion f in x_0 gleich der Steigung der Kurventangente.

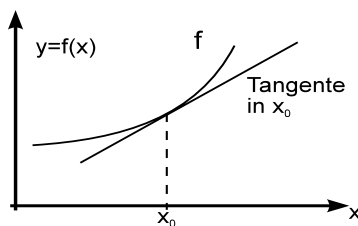


Abb. 10.15. Ableitung = Steigung der Tangente in x_0

Dieser Ableitungsbegriff wird auf Funktionen von zwei Variablen $f(x, y)$ erweitert. Wenn wir eine Variable festhalten (z.B. $y = y_0$), dann ist

$$z = f(x, y = y_0)$$

eine Funktion in der Variablen x . Ist diese Funktion im Punkte x_0 differenzierbar, so nennen wir ihre Ableitung die *partielle Ableitung nach x* . Analog wird die partielle Ableitung von f nach y definiert.

Definition: (Partielle Ableitung). Eine Funktion f heißt im Punkt $(x_0, y_0) \in \mathbb{D}$ **partiell nach x differenzierbar**, wenn der Grenzwert

$$\frac{\partial f}{\partial x}(x_0, y_0) := \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x, y_0) - f(x_0, y_0)}{\Delta x}$$

existiert. Man bezeichnet ihn als die **partielle Ableitung von f nach x im Punkte (x_0, y_0)** .

Entsprechend heißt

$$\frac{\partial f}{\partial y}(x_0, y_0) := \lim_{\Delta y \rightarrow 0} \frac{f(x_0, y_0 + \Delta y) - f(x_0, y_0)}{\Delta y}$$

die **partielle Ableitung von f nach y im Punkte (x_0, y_0)** , wenn dieser Grenzwert existiert.

Etwas lax formuliert lässt sich die Definition so zusammenfassen: **Die partiellen Ableitungen sind nichts anderes als die gewöhnlichen Ableitungen, bei denen alle Variablen bis auf eine festgehalten werden.** Die wichtige Konsequenz hiervon ist, dass sich alle Differenzierungsregeln von Funktionen einer Variablen auf die partielle Differenziation übertragen.

Bemerkungen:

- (1) Man beachte, dass die partiellen Ableitungen im Gegensatz zu den gewöhnlichen Ableitungen nicht durch Striche (oder Punkte im Falle der zeitlichen Ableitung) gekennzeichnet werden, sondern durch Indizierung mit der Differenzierungsvariablen. Allgemein übliche Bezeichnungen sind

$$f_x(x, y), \quad \frac{\partial f}{\partial x}(x, y), \quad \frac{\partial}{\partial x} f(x, y), \quad \partial_x f(x, y),$$

bzw. kurz

$$f_x, \quad \frac{\partial f}{\partial x}, \quad \frac{\partial}{\partial x} f, \quad \partial_x f.$$

Um anzudeuten, dass keine gewöhnlichen Ableitungen vorliegen, wird also auch $\frac{d}{dx}$ durch $\frac{\partial}{\partial x}$ ersetzt. Analoge Bezeichnungen gelten für die partiellen Ableitungen nach y .

- (2) In Anlehnung an die gewöhnliche Ableitung f' bezeichnet man f_x und f_y als partielle Ableitungen **1. Ordnung**.
- (3) Alternative Schreibweisen für die partiellen Differenzialquotienten sind

$$\begin{aligned} \frac{\partial f}{\partial x}(x_0, y_0) &= \lim_{x \rightarrow x_0} \frac{f(x, y_0) - f(x_0, y_0)}{x - x_0} \\ &= \lim_{h \rightarrow 0} \frac{f(x_0 + h, y_0) - f(x_0, y_0)}{h} \end{aligned}$$

$$\begin{aligned} \text{und} \quad \frac{\partial f}{\partial y}(x_0, y_0) &= \lim_{y \rightarrow y_0} \frac{f(x_0, y) - f(x_0, y_0)}{y - y_0} \\ &= \lim_{h \rightarrow 0} \frac{f(x_0, y_0 + h) - f(x_0, y_0)}{h} . \end{aligned}$$

Im Folgenden geben wir einfache Beispiele zur Berechnung der partiellen Ableitung an. Man beachte, dass hierbei insbesondere die Kettenregel zur Anwendung kommt.

Beispiele 10.10 (Partielle Ableitung):

① $f(x, y) = x^2 \cdot y^3 + x + y^2$:

$$\frac{\partial f}{\partial x}(x, y) = 2x \cdot y^3 + 1; \quad \frac{\partial f}{\partial y}(x, y) = x^2 \cdot 3y^2 + 2y .$$

② $f(x, y) = \sin(x^2 - y)$:

$$\frac{\partial f}{\partial x}(x, y) = \cos(x^2 - y) \cdot 2x; \quad \frac{\partial f}{\partial y}(x, y) = \cos(x^2 - y) (-1) .$$

③ $f(x, y) = \ln\left(2x + 4\frac{1}{y}\right)$:

$$\frac{\partial f}{\partial x}(x, y) = \frac{1}{2x + 4\frac{1}{y}} \cdot 2; \quad \frac{\partial f}{\partial y}(x, y) = \frac{1}{2x + 4\frac{1}{y}} \cdot 4(-1) y^{-2} .$$

④ $U = R \cdot I$:

$$\frac{\partial U}{\partial R} = I; \quad \frac{\partial U}{\partial I} = R .$$

⑤ $W = \frac{1}{g} v_0^2 \sin(2\alpha)$:

$$\frac{\partial W}{\partial v_0} = \frac{1}{g} 2 v_0 \sin(2\alpha); \quad \frac{\partial W}{\partial \alpha} = \frac{1}{g} v_0^2 \cos(2\alpha) \cdot 2 .$$

⑥ $p = R \cdot \frac{T}{V}$:

$$\frac{\partial p}{\partial V} = -R \frac{T}{V^2}; \quad \frac{\partial p}{\partial T} = \frac{R}{V} .$$

□

Bemerkung: Beispiel 10.10 ⑥ zeigt die physikalisch-chemische Bedeutung der partiellen Ableitung: Der Druck p eines idealen Gases ist proportional zum Quotienten $\frac{T}{V}$. Somit ist p eine Funktion der beiden Variablen T und V . $\frac{\partial p}{\partial V}$ bedeutet dann die Änderung des Druckes als Funktion des Volumens, wenn die Temperatur konstant gehalten wird. In der Chemie wird dies oftmals durch $\left(\frac{\partial p}{\partial V}\right)_{T=\text{const}}$ symbolisiert. $\frac{\partial p}{\partial T}$ bedeutet entsprechend die Änderung des Druckes bei Änderung der Temperatur aber konstantem Volumen. □

Geometrische Interpretation: Die anschauliche Bedeutung der partiellen Ableitungen erläutern wir mit Hilfe von Abb. 10.16 und Abb. 10.17.

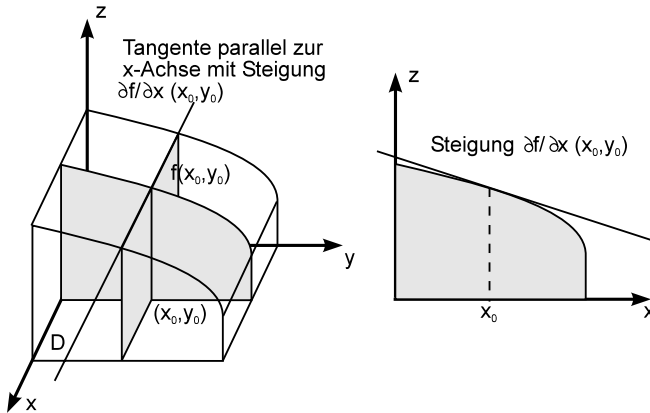


Abb. 10.16. Partielle Ableitung in x-Richtung

Die partielle Ableitung $\frac{\partial}{\partial x} f(x_0, y_0)$ von f nach x im Punkte (x_0, y_0) gibt die Steigung der Tangente im Punkte (x_0, y_0) parallel zur x -Achse an. Entsprechend gibt partielle Ableitung $\frac{\partial}{\partial y} f(x_0, y_0)$ von f nach y im Punkte (x_0, y_0) die Steigung der Tangente im Punkte (x_0, y_0) parallel zur y -Achse an.

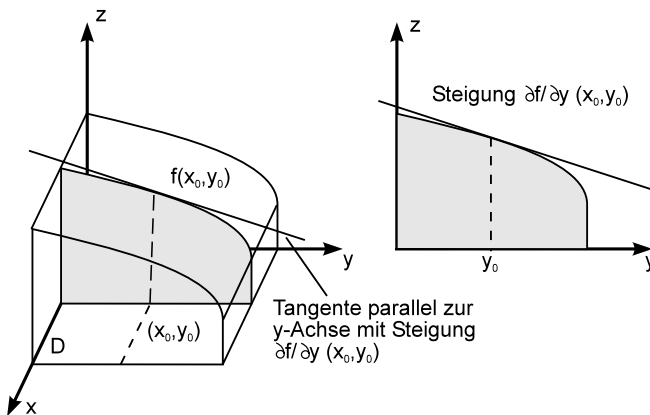


Abb. 10.17. Partielle Ableitung in y-Richtung

➤ Partielles Differenzieren von Funktionen mit mehreren Variablen

Ist f eine Funktion der Variablen $(x_1, \dots, x_n)$, so ist die partielle Ableitung von f nach der Variablen x_i in einem Punkt $(x_1^0, \dots, x_n^0)$ definiert als der Grenzwert

$$\frac{\partial f}{\partial x_i}(x_1^0, \dots, x_n^0) = \lim_{h \rightarrow 0} \frac{f(x_1^0, \dots, x_i^0 + h, \dots, x_n^0) - f(x_1^0, \dots, x_i^0, \dots, x_n^0)}{h},$$

falls dieser existiert. Man nennt ihn die **partielle Ableitung von f nach x_i** im Punkte $(x_1^0, \dots, x_n^0)$ und bezeichnet ihn mit

$$\frac{\partial f}{\partial x_i}, \quad \frac{\partial}{\partial x_i} f, \quad f_{x_i}, \quad \partial_{x_i} f.$$

Beispiel 10.11 (Mit MAPLE-Worksheet). Gesucht sind alle partiellen Ableitungen 1. Ordnung der Funktion

$$f(x, y, z) = z \cdot e^{x^2+y^2} + \sqrt{1+x^2+z^4}.$$

Mit der Kettenregel berechnet man

$$\begin{aligned} f_x(x, y, z) &= 2xz e^{x^2+y^2} + \frac{x}{\sqrt{1+x^2+z^4}} \\ f_y(x, y, z) &= 2yz e^{x^2+y^2} \\ f_z(x, y, z) &= e^{x^2+y^2} + \frac{2z^3}{\sqrt{1+x^2+z^4}}. \end{aligned}$$

Wir bestimmen noch die partiellen Ableitungen an der Stelle $(x_0, y_0, z_0) = (1, 2, 0)$ durch Einsetzen des Punktes in die obigen Formeln

$$f_x(1, 2, 0) = \frac{1}{2}\sqrt{2}; \quad f_y(1, 2, 0) = 0; \quad f_z(1, 2, 0) = e^5. \quad \square$$

⑤ Ableitungen höherer Ordnung

Sind die partiellen Ableitungen $f_x(x, y)$ und $f_y(x, y)$ ihrerseits wieder partiell differenzierbar, so bezeichnet man ihre partiellen Ableitungen $\frac{\partial}{\partial x} f_x(x, y)$, $\frac{\partial}{\partial y} f_x(x, y)$ und $\frac{\partial}{\partial x} f_y(x, y)$, $\frac{\partial}{\partial y} f_y(x, y)$ als partielle Ableitungen *zweiter Ordnung* von f . Die Schreibweise für die partiellen Ableitungen zweiter Ordnung lauten

$$f_{xx} := \frac{\partial^2 f}{\partial x^2} := \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) \quad \text{zweite partielle Ableitung nach } x.$$

$$f_{yy} := \frac{\partial^2 f}{\partial y^2} := \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial y} \right) \quad \text{zweite partielle Ableitung nach } y.$$

$$f_{xy} := \frac{\partial^2 f}{\partial y \partial x} := \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial x} \right) \quad \text{zweite partielle Ableitung nach } x \text{ und } y.$$

$$f_{yx} := \frac{\partial^2 f}{\partial x \partial y} := \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial y} \right) \quad \text{zweite partielle Ableitung nach } y \text{ und } x.$$

Deren Ableitungen wiederum, sofern sie existieren, sind die *dritten partiellen Ableitungen* von f :

$$f_{xxx}, f_{xxy}, f_{xyx}, f_{yxx}, f_{xyy}, f_{yyx}, f_{yyy} .$$

Bemerkungen:

- (1) Die Reihenfolge, in der die Differenziation durchgeführt werden muss, ist von innen nach außen (von links nach rechts): Die Ableitung f_{xy} wird gebildet, indem von der Funktion (f_x) die partielle Ableitung nach y gebildet wird:

$$f_{xy} = (f_x)_y .$$

- (2) Man bezeichnet eine Ableitung als *gemischte* Ableitung, wenn nicht nur nach einer Variablen differenziert wird.
- (3) Die *Ordnung* der partiellen Ableitung entspricht der Gesamtzahl der zu bildenden Ableitungen, d.h. der Gesamtzahl der Indizes:

$$f_{xyxx} \quad \text{ist z.B. eine Ableitung 4. Ordnung.}$$

Beispiel 10.12. Für die Funktion

$$f(x, y) = x^3 y + y$$

sind die partiellen Ableitungen bis zur Ordnung 3 gegeben durch

$$f_x = 3x^2 y, \quad f_y = x^3 + 1;$$

$$f_{xx} = 6xy, \quad f_{xy} = 3x^2 = f_{yx}, \quad f_{yy} = 0;$$

$$f_{xxx} = 6y, \quad f_{xxy} = 6x = f_{xyx} = f_{yxx}, \\ f_{xyy} = 0 = f_{yxy} = f_{yyx}, \quad f_{yyy} = 0. \quad \square$$

In diesem Beispiel kommt es nicht auf die Reihenfolge der Ableitungen an, $f_{xy} = f_{yx}$, $f_{xxy} = f_{xyx} = f_{yxx}$ usw. Diese Eigenschaft bestätigt sich für praktisch alle in den Anwendungen vorkommenden Funktionen. Es gilt die folgende wichtige Aussage

Satz von Schwarz: Vertauschbarkeit von gemischten Ableitungen. Sind für eine Funktion $f(x, y)$ in zwei Variablen die gemischten partiellen Ableitungen f_{xy} und f_{yx} **stetig**, dann kommt es nicht auf die Reihenfolge der zu bildenden Ableitungen an. D.h. es gilt dann

$$f_{xy}(x, y) = f_{yx}(x, y).$$

Verallgemeinerung: Sind für eine Funktion f alle partiellen Ableitungen bis zur Ordnung k (≥ 2) stetig, dann kommt es bei allen partiellen Ableitungen bis zur Ordnung k nicht auf die Reihenfolge der zu bildenden Ableitungen an. Analog zu den höheren partiellen Ableitungen für Funktionen von zwei Variablen bildet man sie für Funktionen mit mehr als zwei Variablen. Auch hier ist der Satz von Schwarz gültig.

Beispiel 10.13. Von der Funktion

$$f(x, y, z) = e^{x-y} \cos(5z)$$

sind alle partiellen Ableitungen bis zur Ordnung 2 gesucht:

Ableitungen erster Ordnung:

$$\begin{aligned} f_x &= e^{x-y} \cos(5z) \\ f_y &= -e^{x-y} \cos(5z) \\ f_z &= -5e^{x-y} \sin(5z). \end{aligned}$$

Ableitungen zweiter Ordnung (bezüglich einer Variablen):

$$\begin{aligned} f_{xx} &= e^{x-y} \cos(5z) \\ f_{yy} &= e^{x-y} \cos(5z) \\ f_{zz} &= -25e^{x-y} \cos(5z); \end{aligned}$$

Ableitungen zweiter Ordnung (gemischt):

$$\begin{aligned} f_{xy} &= -e^{x-y} \cos(5z) = f_{yx} \\ f_{xz} &= -5e^{x-y} \sin(5z) = f_{zx} \\ f_{yz} &= 5e^{x-y} \sin(5z) = f_{zy}. \end{aligned}$$

□

Beispiele 10.14:

$$\textcircled{1} \quad \varphi(x, y) = \sqrt{x^2 + y^2}:$$

$$\frac{\partial \varphi}{\partial x} = \frac{x}{\sqrt{x^2 + y^2}}; \quad \frac{\partial \varphi}{\partial y} = \frac{y}{\sqrt{x^2 + y^2}};$$

$$\frac{\partial^2 \varphi}{\partial x^2} = \frac{\sqrt{x^2 + y^2} - x \frac{x}{\sqrt{x^2 + y^2}}}{x^2 + y^2} = \frac{x^2 + y^2 - x^2}{(x^2 + y^2)^{\frac{3}{2}}} = \frac{y^2}{(x^2 + y^2)^{\frac{3}{2}}};$$

$$\frac{\partial^2 \varphi}{\partial y^2} = \frac{x^2}{(x^2 + y^2)^{\frac{3}{2}}}.$$

Für die Summe der zweiten Ableitungen $\varphi_{xx} + \varphi_{yy}$ gilt weiterhin:

$$\varphi_{xx} + \varphi_{yy} = \frac{x^2 + y^2}{(x^2 + y^2)^{\frac{3}{2}}} = \frac{1}{(x^2 + y^2)^{\frac{1}{2}}} = \frac{1}{\varphi(x, y)}.$$

$$\textcircled{2} \quad r(x, y, z) = \sqrt{x^2 + y^2 + z^2}:$$

$$\frac{\partial r}{\partial x} = \frac{x}{\sqrt{x^2 + y^2 + z^2}} = \frac{x}{r}; \quad \frac{\partial r}{\partial y} = \frac{y}{r}; \quad \frac{\partial r}{\partial z} = \frac{z}{r}$$

$$\frac{\partial^2 r}{\partial x^2} = \frac{r - x \frac{x}{r}}{r^2} = \frac{r^2 - x^2}{r^3}; \quad \frac{\partial^2 r}{\partial y^2} = \frac{r^2 - y^2}{r^3}; \quad \frac{\partial^2 r}{\partial z^2} = \frac{r^2 - z^2}{r^3}$$

Für die Summe der zweiten Ableitungen $r_{xx} + r_{yy} + r_{zz}$ gilt weiterhin:

$$r_{xx} + r_{yy} + r_{zz} = \frac{3r^2 - (x^2 + y^2 + z^2)}{r^3} = \frac{2r^2}{r^3} = \frac{2}{r}.$$

$$\textcircled{3} \quad f(x, y, z) = \frac{1}{r} = \frac{1}{\sqrt{x^2 + y^2 + z^2}}:$$

$$\frac{\partial f}{\partial x} = \frac{-x}{(x^2 + y^2 + z^2)^{\frac{3}{2}}} = \frac{-x}{r^3}; \quad \frac{\partial f}{\partial y} = \frac{-y}{r^3}; \quad \frac{\partial f}{\partial z} = \frac{-z}{r^3}$$

$$\frac{\partial^2 f}{\partial x^2} = -\frac{r^3 - x \frac{3}{2} 2x}{r^6} = -\frac{r^2 - 3x^2}{r^5}$$

$$\frac{\partial^2 f}{\partial y^2} = -\frac{r^2 - 3y^2}{r^5}; \quad \frac{\partial^2 f}{\partial z^2} = -\frac{r^2 - 3z^2}{r^5}$$

Für die Summe der zweiten Ableitungen $f_{xx} + f_{yy} + f_{zz}$ gilt weiterhin:

$$f_{xx} + f_{yy} + f_{zz} = -\frac{3r^2 - 3(x^2 + y^2 + z^2)}{r^5} = 0. \quad \square$$

10.3.2 Totale Differenzierbarkeit

Während differenzierbare Funktionen einer Variablen immer stetige Funktionen sind, kann man dies i.A. von partiell differenzierbaren Funktionen nicht behaupten.

Beispiel 10.15. Für die Funktion

$$f(x, y) = \frac{2xy}{x^2 + y^2}, \quad (x, y) \neq (0, 0)$$

berechnen sich die partiellen Ableitungen mit der Quotientenregel

$$f_x(x, y) = \frac{2y(y^2 - x^2)}{(x^2 + y^2)^2} \quad \text{und} \quad f_y(x, y) = \frac{2x(x^2 - y^2)}{(x^2 + y^2)^2}.$$

Im Punkte $(x, y) = (0, 0)$ existiert sowohl die partielle Ableitung von f nach x , als auch nach y : $f_x(0, 0) = f_y(0, 0) = 0$, obwohl die Funktion nach Beispiel 10.9 ④ dort nicht stetig ist! $\square$

Daher führt man den Begriff der *totalen* Differenzierbarkeit ein, der vom Begriff der *partiellen* Differenzierbarkeit zu unterscheiden ist. Man nennt eine Funktion f von zwei Variablen im Punkte (x_0, y_0) total differenzierbar, wenn sie nahe dieses Punktes durch eine Ebene angenähert werden kann:

Definition: (Totale Differenzierbarkeit). Die Funktion f heißt im Punkte $(x_0, y_0) \in \mathbb{D}$ **total differenzierbar**, wenn es Zahlen $A, B \in \mathbb{R}$ und Funktionen $\varepsilon_1(x, y), \varepsilon_2(x, y)$ gibt, so dass für alle (x, y) nahe bei (x_0, y_0) gilt

$$\begin{aligned} f(x, y) &= f(x_0, y_0) + A \cdot (x - x_0) + B \cdot (y - y_0) \\ &\quad + \varepsilon_1(x, y)(x - x_0) + \varepsilon_2(x, y)(y - y_0), \end{aligned} \tag{*}$$

wenn die Funktionen ε_1 und ε_2 gegen Null gehen für $(x, y) \rightarrow (x_0, y_0)$:

$$\begin{aligned} \varepsilon_1(x, y) &\rightarrow 0 \quad \text{für } (x, y) \rightarrow (x_0, y_0) \\ \varepsilon_2(x, y) &\rightarrow 0 \quad \text{für } (x, y) \rightarrow (x_0, y_0). \end{aligned}$$

Gleichung (*) besagt, dass in der Nähe des Punktes (x_0, y_0) die Funktionswerte $f(x, y)$ näherungsweise durch die lineare Funktion

$$z(x, y) = f(x_0, y_0) + A(x - x_0) + B(y - y_0)$$

beschrieben werden. Der Graph von z stellt eine Ebene im $\mathbb{R}^3$ dar, die durch den Punkt $(x_0, y_0, f(x_0, y_0))$ geht. Sie heißt die **Tangentialebene** von f in (x_0, y_0) , da sie sich in der Umgebung dieses Punktes an den Graphen der Funktion f anschmiegt.

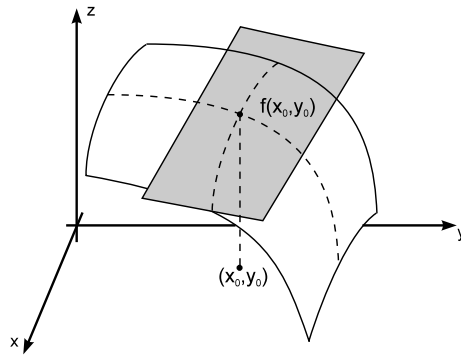


Abb. 10.18. Funktion und Tangentialebene

Aus der totalen Differenzierbarkeit folgt sowohl die partielle Differenzierbarkeit als auch die Stetigkeit von f :

Satz: Ist f in (x_0, y_0) **total differenzierbar**, dann folgt

- (1) f ist in (x_0, y_0) stetig.
- (2) Es existieren die partiellen Ableitungen

$$f_x(x_0, y_0) \quad , \quad f_y(x_0, y_0) \quad .$$

- (3) Die Zahlenwerte A und B in Gleichung (*) berechnen sich durch

$$A = f_x(x_0, y_0) \quad , \quad B = f_y(x_0, y_0) \quad .$$

- (4) Der Graph der Funktion f lässt sich in der Nähe des Punktes annähern durch die **Tangentialebene** z_t

$$f(x, y) \approx z_t = f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) \quad .$$

Die Definition für die totale Differenzierbarkeit ist anschaulich zwar einprägsam, aber im konkreten Fall schwierig nachzuprüfen. Man kann aber anhand der partiellen Ableitungen entscheiden, ob eine Funktion f total differenzierbar ist:

Satz: f ist in einer Umgebung von $(x_0, y_0) \in \mathbb{D}$ partiell nach x und y differenzierbar und die **partiellen Ableitungen f_x und f_y sind in (x_0, y_0) stetig**. Dann ist f in (x_0, y_0) **total differenzierbar**.

Beispiel 10.16 (Tangentialebene, mit MAPLE-Worksheet). Gesucht ist die Tangentialebene der Funktion

$$f(x, y) = e^{-(x^2+y^2)} \quad \text{im Punkte } (x_0, y_0) = (0.15, 0.15).$$

Die partiellen Ableitungen der Funktion

$$\begin{aligned} f_x(x, y) &= -2x e^{-(x^2+y^2)} \\ f_y(x, y) &= -2y e^{-(x^2+y^2)} \end{aligned}$$

sind stetig. Daher ist die Funktion total differenzierbar und die Tangentialebene ist gegeben durch

$$\begin{aligned} z &= f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) \\ &= e^{-\frac{9}{200}} - \frac{3}{10}e^{-\frac{9}{200}}(x - 0.15) - \frac{3}{10}e^{-\frac{9}{200}}(y - 0.15). \end{aligned}$$

Wir stellen in Abb. 10.19 sowohl die Funktion als auch die Tangentialebene graphisch dar:

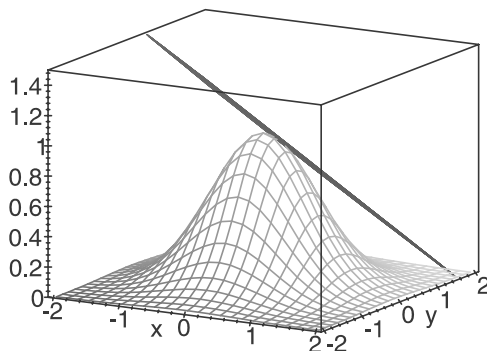


Abb. 10.19. Funktion $f(x, y) = e^{-(x^2+y^2)}$ mit Tangentialebene

Man erkennt, dass die Tangentialebene im Entwicklungspunkt (x_0, y_0) den Graphen der Funktion f berührt. $\square$

➤ 10.3.3 Gradient und Richtungsableitung

In diesem Abschnitt gehen wir von einer Funktion $f(x, y)$ mit zwei Variablen x und y aus, die stetig partiell differenzierbar nach x als auch nach y ist.

⊗ Der Gradient

Bei Funktionen einer Variablen gibt die Ableitung der Funktion im Punkte x_0 die Steigung (= Steilheit) der Funktion an. An Stellen großer Ableitung ändert sich die Funktion stark, an Stellen geringer Ableitung ändert sie sich schwach. Bei Funktionen zweier Variablen berücksichtigt man als Maß sowohl die partielle Ableitung in x -Richtung als auch in y -Richtung. Um eine Funktion f bezüglich ihrer Steigung in einem Punkt (x_0, y_0) zu charakterisieren,

führt man den Vektor $\text{grad } f(x_0, y_0) := \begin{pmatrix} \frac{\partial f}{\partial x}(x_0, y_0) \\ \frac{\partial f}{\partial y}(x_0, y_0) \end{pmatrix}$ ein, dessen Komponenten aus der partiellen Ableitung nach x und y bestehen. Er beschreibt die Neigung der Tangentialebene im Punkte (x_0, y_0) ! Da dieser Vektor den Grad der Steigung der Funktion angibt, nennt man ihn den *Gradienten*.

Definition: (Gradient)

Der Vektor

$$\text{grad } f(x_0, y_0) := \begin{pmatrix} \frac{\partial f}{\partial x}(x_0, y_0) \\ \frac{\partial f}{\partial y}(x_0, y_0) \end{pmatrix}$$

heißt der **Gradient von f** an der Stelle (x_0, y_0) .

Für den Gradienten wird oft der sog. Nabla-Operator ∇ verwendet:

$$\nabla := \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \end{pmatrix} \Rightarrow \text{grad } f = \nabla f = \begin{pmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \end{pmatrix}.$$

Beispiele 10.17 (Mit MAPLE-Worksheet):

(1) Gesucht ist der Gradient der Funktion

$$f = \sqrt{x^2 + y^2 + 1} : \quad \text{grad } f = \begin{pmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \end{pmatrix} = \begin{pmatrix} \frac{x}{\sqrt{x^2 + y^2 + 1}} \\ \frac{y}{\sqrt{x^2 + y^2 + 1}} \end{pmatrix}.$$

Stellt man dieses Ergebnis in Form einer Vektorgraphik dar, erhält man die zweidimensionale Darstellung:

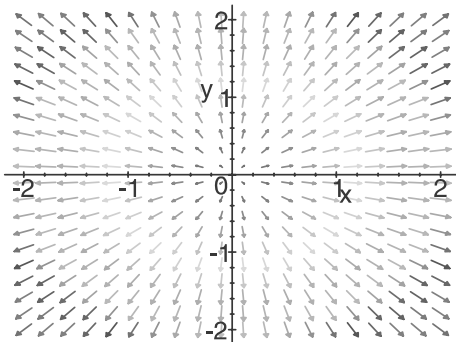


Abb. 10.20. Zweidimensionale Darstellung des Gradienten einer Funktion $f(x, y)$

(2) Gesucht ist der Gradient der Funktion

$$f = \sqrt{x^2 + y^2 + z^2 + 1} : \quad \text{grad } f = \begin{pmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \\ \frac{\partial f}{\partial z} \end{pmatrix} = \begin{pmatrix} \frac{x}{\sqrt{x^2 + y^2 + z^2 + 1}} \\ \frac{y}{\sqrt{x^2 + y^2 + z^2 + 1}} \\ \frac{z}{\sqrt{x^2 + y^2 + z^2 + 1}} \end{pmatrix}.$$

Stellt man dieses Ergebnis in Form einer Vektorgraphik dar, erhält man die dreidimensionale Darstellung:

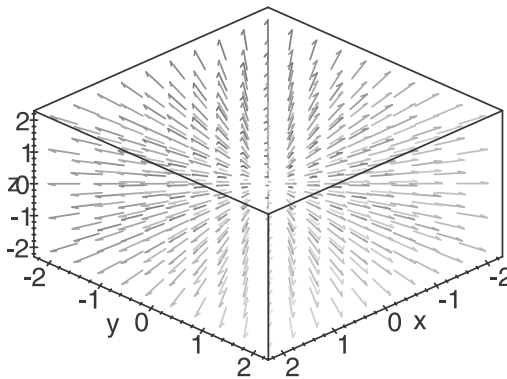


Abb. 10.21. Dreidimensionale Darstellung des Gradienten einer Funktion $f(x, y, z)$

Anwendungsbeispiel 10.18 (Elektrisches Feld einer Punktladung).

Gegeben ist eine Punktladung q an der Stelle $\vec{r}_0 = (x_0, y_0, z_0)$. Gesucht ist das Potenzial $\Phi(\vec{r})$ und das elektrische Feld $\vec{E}(\vec{r})$ in einem beliebigen Punkt des Raumes $\vec{r} = (x, y, z)$. Das durch die Punktladung induzierte Potenzial ist

$$\Phi(\vec{r}) = \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|} = \frac{1}{4\pi\epsilon_0} \frac{q}{\sqrt{(x - x_0)^2 + (y - y_0)^2 + (z - z_0)^2}}$$

und das zugehörige elektrische Feld ist definiert durch

$$\vec{E}(\vec{r}) := -\text{grad } \Phi(\vec{r}) = - \begin{pmatrix} \frac{\partial}{\partial x} \Phi(x, y, z) \\ \frac{\partial}{\partial y} \Phi(x, y, z) \\ \frac{\partial}{\partial z} \Phi(x, y, z) \end{pmatrix}.$$

Um Missverständnisse mit den partiellen Ableitungen von $\vec{E}$ auszuschließen, wird im Folgenden die x -Komponente des elektrischen Feldes statt E_x mit E_1 , die y -Komponente statt E_y mit E_2 und die z -Komponente statt E_z mit E_3

bezeichnet:

$$\begin{aligned}
 E_1(x, y, z) &= -\partial_x \Phi(x, y, z) \\
 &= -\partial_x \frac{q}{4\pi\epsilon_0} \left((x-x_0)^2 + (y-y_0)^2 + (z-z_0)^2 \right)^{-\frac{1}{2}} \\
 &= \frac{1}{4\pi\epsilon_0} \frac{q}{\sqrt{(x-x_0)^2 + (y-y_0)^2 + (z-z_0)^2}^3} (x-x_0) \\
 &= \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|^3} (x-x_0)
 \end{aligned}$$

Analog berechnen sich

$$E_2(x, y, z) = -\partial_y \Phi(x, y, z) = \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|^3} (y-y_0)$$

$$E_3(x, y, z) = -\partial_z \Phi(x, y, z) = \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|^3} (z-z_0)$$

$$\Rightarrow \vec{E}(\vec{r}) = \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|^3} \begin{pmatrix} x-x_0 \\ y-y_0 \\ z-z_0 \end{pmatrix} = \frac{1}{4\pi\epsilon_0} \frac{q}{|\vec{r} - \vec{r}_0|^3} (\vec{r} - \vec{r}_0).$$

Man nennt $\vec{E}(\vec{r})$ auch ein *Vektorfeld*. □

➤ Die Richtungsableitung

In Abb. 10.22 sind für das Zwei-Elektroden-System eines elektrolytischen Trogs Äquipotenziallinien schematisch gezeichnet.

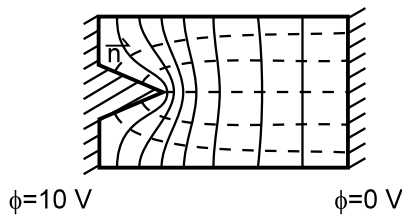
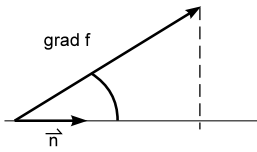


Abb. 10.22. Qualitativer Verlauf der Potenziallinien im elektrolytischen Trog

Der **Gradient** von Φ steht senkrecht zu den **Äquipotenziallinien**; der Betrag ist ein Maß für die Dichte der Potenziallinien: An Stellen mit hoher Dichte (nahe der Kante) stellt sich ein hohes elektrisches Feld ein, an Stellen mit geringer Dichte (rechte Elektrode) ein kleines Feld. Wird das elektrische Feld z.B. auf der Elektrodenoberfläche gesucht, so ist nicht die Ableitung von Φ in Richtung x oder y gesucht, sondern die Ableitung von Φ in eine vorgegebene Richtung $\vec{n}$. Dies führt auf den Begriff der *Richtungsableitung*:

Gegeben ist eine Funktion f , gesucht ist die Änderung der Funktion in Richtung $\vec{n}$, wenn $\vec{n}$ der Richtungs-Einheitsvektor ist. Zur Bestimmung der Richtungsableitung projiziert man den Gradienten in Richtung des Vektors $\vec{n}$.



Nach Kapitel 2.2.2 gilt für die Projektion eines Vektors $\vec{b}$ in Richtung $\vec{a}$ die Formel

$$\vec{b}_a = \frac{\vec{a} \cdot \vec{b}}{|\vec{a}|^2} \cdot \vec{a},$$

wenn $\vec{a} \cdot \vec{b}$ das Skalarprodukt und $|\vec{a}|$ der Betrag des Vektors $\vec{a}$ ist. Für den Fall, dass $\vec{a}$ ein Einheitsvektor ist (d.h. $|\vec{a}| = 1$), gilt für den Betrag von $\vec{b}_a$

$$|\vec{b}_a| = b_a = \vec{a} \cdot \vec{b}.$$

Übertragen auf unser Problem der Ableitung in Richtung $\vec{n}$ bedeutet dies:

Definition: (Richtungsableitung). Die Ableitung einer Funktion $f(x, y)$ mit zwei Variablen in Richtung des Einheitsvektors $\vec{n} = \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$ ist gegeben durch

$$\begin{aligned} \frac{\partial f}{\partial \vec{n}} &:= \vec{n} \cdot \text{grad } f = \begin{pmatrix} n_1 \\ n_2 \end{pmatrix} \cdot \begin{pmatrix} \partial_x f(x, y) \\ \partial_y f(x, y) \end{pmatrix} \\ &= n_1 \partial_x f(x, y) + n_2 \partial_y f(x, y). \end{aligned}$$

Sie heißt **Richtungsableitung von f** und wird auch mit

$$\partial_{\vec{n}} f, \quad \frac{\partial}{\partial \vec{n}} f, \quad D_{\vec{n}} f$$

bezeichnet.

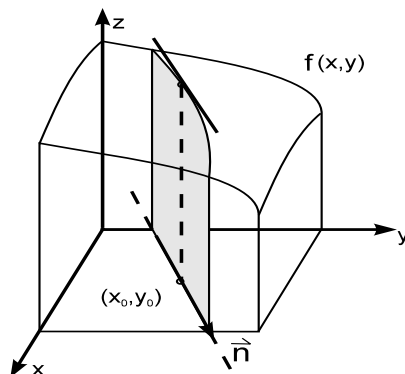


Abb. 10.23. Richtungsableitung

⚠ Achtung: Im Gegensatz zum Gradienten stellt die Richtungsableitung keinen Vektor, sondern eine skalare Größe dar. Die partiellen Ableitungen nach x und y sind Spezialfälle der Richtungsableitung.

Spezialfälle:

Für $\vec{n} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ ist $\partial_{\vec{n}} f = \frac{\partial f}{\partial x}$ die partielle Ableitung nach x .

Für $\vec{n} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ ist $\partial_{\vec{n}} f = \frac{\partial f}{\partial y}$ die partielle Ableitung nach y .

Die Richtungsableitung von f in Richtung $\vec{n}$ ist also die Projektion des Gradienten $\text{grad } f$ auf die Gerade mit Richtung $\vec{n}$. Folglich gilt:

Satz:

- (1) Die Richtungsableitung ist am größten, wenn die Richtung $\vec{n}$ parallel zum Gradienten ist.
- (2) Die Richtungsableitung ist Null, wenn $\vec{n}$ senkrecht zum Gradienten $\text{grad } f$ steht.
- (3) Der Gradientenvektor $\text{grad } f$ zeigt in Richtung des stärksten Anstiegs bzw. Abfalls der Funktion $f(x, y)$. Sein Betrag gibt die Größe der Steigung bzw. des Abfalls an.

Beispiele 10.19 (Richtungsableitung, mit MAPLE-Worksheet):

- ① Gegeben ist die Funktion $f(x, y) = \sqrt{x^2 + y^2}$. Gesucht ist die Ableitung von f in Richtung $\vec{n} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$:

Die partiellen Ableitungen von f nach x und y sind

$$f_x = \frac{x}{\sqrt{x^2 + y^2}}, \quad f_y = \frac{y}{\sqrt{x^2 + y^2}}.$$

Damit bestimmt sich die Richtungsableitung in Richtung $\vec{n}$ zu

$$\frac{\partial f}{\partial \vec{n}} = \vec{n} \cdot \text{grad } f = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \cdot \frac{1}{\sqrt{x^2 + y^2}} \begin{pmatrix} x \\ y \end{pmatrix} = \frac{1}{\sqrt{2}} \frac{x + y}{\sqrt{x^2 + y^2}}.$$

- ② $f(x, y) = x \cdot y$. Gesucht ist die Ableitung von f in Richtung $\vec{a} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$.

Die partiellen Ableitungen sind

$$f_x = y, \quad f_y = x.$$

Der zum Vektor $\vec{a} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ gehörende Normalenvektor ist

$$\begin{aligned}\vec{n} &= \frac{1}{|\vec{a}|} \vec{a} = \frac{1}{\sqrt{5}} \begin{pmatrix} 1 \\ 2 \end{pmatrix} . \\ \Rightarrow \frac{\partial f}{\partial \vec{a}} &= \frac{1}{\sqrt{5}} \begin{pmatrix} 1 \\ 2 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \frac{1}{\sqrt{5}} (2x + y) . \quad \square\end{aligned}$$

Bemerkung: Die Richtungsableitung von f in Richtung des Einheitsvektors $\vec{n} = \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$ im Punkte (x_0, y_0) wird oftmals auch in der äquivalenten Darstellung

$$\frac{\partial f}{\partial \vec{n}} = \lim_{h \rightarrow 0} \frac{f(x_0 + h n_1, y_0 + h n_2) - f(x_0, y_0)}{h}$$

definiert. Dieser Grenzwert spiegelt die Ableitung entlang der Geraden

$$\begin{pmatrix} x_0 + h n_1 \\ y_0 + h n_2 \end{pmatrix} = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix} + h \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$$

wider. Diese Gerade ist in der Punkt-Richtungs-Darstellung durch den Punkt $\begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ und die Richtung $\vec{n} = \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$ festgelegt.

➤ 10.3.4 Der Taylorsche Satz

Eine wesentliche Eigenschaft von differenzierbaren Funktionen einer Variablen besteht darin, dass sie in der Umgebung eines Punktes näherungsweise durch Polynome ersetzt werden können. Dies ist auch im mehrdimensionalen Fall möglich. Wir werden den Taylorschen Satz nicht beweisen, stattdessen geben wir eine Plausibilitätsüberlegung für die Taylorsche Formel für Funktionen mit zwei Variablen an:

Für eine $(m+1)$ -mal stetig differenzierbare Funktion $f(x)$ gilt nach der Taylorschen Formel aus Kapitel 9.3 am Entwicklungspunkt $x_0 \in \mathbb{D}$

$$\begin{aligned}f(x) &= f(x_0) + (x - x_0) f'(x_0) + \frac{1}{2} (x - x_0)^2 f''(x_0) + \dots + \\ &\quad + \frac{1}{m!} (x - x_0)^m f^{(m)}(x_0) + R_m(x) ,\end{aligned}$$

wenn die Differenz zwischen dem Polynom und der Funktion durch das Restglied

$$R_m(x) = \frac{1}{(m+1)!} f^{(m+1)}(\xi) (x - x_0)^{m+1}$$

mit einem nicht näher bekannten Wert ξ , der zwischen x und x_0 liegt, bestimmt ist. In modifizierter Schreibweise lautet die Entwicklung

$$\begin{aligned}
 f(x) &= f(x_0) + (x - x_0) \left. \frac{d}{dx} f \right|_{x_0} + \frac{1}{2!} (x - x_0)^2 \left(\frac{d}{dx} \right)^2 f \Big|_{x_0} + \dots \\
 &\quad + \frac{1}{m!} (x - x_0)^m \left(\frac{d}{dx} \right)^m f \Big|_{x_0} + R_m(x) \\
 &= \sum_{n=0}^m \frac{1}{n!} \left[(x - x_0) \frac{d}{dx} \right]^n f \Big|_{x_0} + R_m(x). \quad (*)
 \end{aligned}$$

In dieser Formel ersetzen wir $\left[(x - x_0) \frac{d}{dx} \right]^n$ durch die entsprechenden partiellen Ableitungen gemäß

$$\begin{aligned}
 (x - x_0) \frac{d}{dx} &\rightarrow (x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \\
 \left[(x - x_0) \frac{d}{dx} \right]^n &\rightarrow \left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n.
 \end{aligned}$$

Nach der Binomischen Formel (Kapitel 1.2.5) ist

$$(a + b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k,$$

so dass mit den Binomialkoeffizienten $\binom{n}{k} = \frac{n!}{(n-k)!k!}$ folgt

$$\begin{aligned}
 \left[(x - x_0) \frac{d}{dx} \right]^n &\rightarrow \sum_{k=0}^n \binom{n}{k} (x - x_0)^{n-k} \left(\frac{\partial}{\partial x} \right)^{n-k} (y - y_0)^k \left(\frac{\partial}{\partial y} \right)^k \\
 &\rightarrow \sum_{k=0}^n \binom{n}{k} \underbrace{(x - x_0)^{n-k} (y - y_0)^k}_{\text{Polynom vom Grad } n} \underbrace{\left(\frac{\partial}{\partial x} \right)^{n-k} \left(\frac{\partial}{\partial y} \right)^k}_{\text{partielle Ableitung der Ordnung } n}.
 \end{aligned}$$

In die Taylorsche Formel (*) eingesetzt, folgt

Satz von Taylor für Funktionen mit zwei Variablen: Sei $f(x, y)$ eine $(m+1)$ -mal stetig partiell differenzierbare Funktion und $(x_0, y_0) \in \mathbb{ID}$ der Entwicklungspunkt. Für $(x, y) \in \mathbb{ID}$ gilt

$$\begin{aligned} f(x, y) &= \sum_{n=0}^m \frac{1}{n!} \left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n f \Big|_{(x_0, y_0)} + R_m \\ &= \sum_{n=0}^m \frac{1}{n!} \sum_{k=0}^n \binom{n}{k} (x - x_0)^{n-k} (y - y_0)^k \cdot \\ &\quad \underbrace{f}_{n-k} \underbrace{x \dots x}_{k} \underbrace{y \dots y}_{k \text{ -mal}} (x_0, y_0) + R_m \end{aligned}$$

mit dem Restglied

$$R_m = \frac{1}{(m+1)!} \sum_{k=0}^{m+1} \binom{m+1}{k} (x - x_0)^{m+1-k} (y - y_0)^k \underbrace{f}_{m+1-k} \underbrace{x \dots x}_{k} \underbrace{y \dots y}_{k \text{ -mal}} (\xi, \eta),$$

wobei (ξ, η) ein nicht näher bekannter Punkt auf der Verbindungsgeraden von (x_0, y_0) und (x, y) , die ganz in $\mathbb{ID}$ liegen soll.

Bemerkungen:

(1) Wenn (x, y) hinreichend nahe bei (x_0, y_0) liegt, dann ist i.A. die Verbindungsgerade zwischen diesen Punkten ebenfalls in $\mathbb{ID}$ enthalten.

(2) Die Formel

$$\left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n f \Big|_{(x_0, y_0)}$$

ist folgendermaßen zu interpretieren: Man multipliziere entsprechend der Potenz n die Summe aus, wende alle Ableitungen auf f an, werte diese Ableitungen an der Stelle (x_0, y_0) aus und multipliziere mit dem Polynom $(x - x_0)^i (y - y_0)^j$, wenn i die Ordnung der Ableitung $\frac{\partial}{\partial x}$ und j die Ordnung der Ableitung $\frac{\partial}{\partial y}$ repräsentiert ($i + j = n$).

(3) **Satz von Taylor für Funktionen mit k Variablen:** Analog zu der Herleitung der Taylorschen Formel für eine Funktion mit zwei Variablen erhält man für eine Funktion f mit k Variablen $(x_1, \dots, x_k)$ die Formel

$$\begin{aligned} f(x_1, \dots, x_k) &= \sum_{n=0}^m \frac{1}{n!} \left[(x_1 - x_1^{(0)}) \frac{\partial}{\partial x_1} + \dots + (x_k - x_k^{(0)}) \frac{\partial}{\partial x_k} \right]^n \\ &\quad f \Big|_{(x_1^{(0)}, \dots, x_k^{(0)})} + R_m(x_1, \dots, x_k) \end{aligned}$$

mit dem Entwicklungspunkt $(x_1^{(0)}, \dots, x_k^{(0)})$ und dem Restglied R_m , das sich analog zum Restglied einer Funktion mit zwei Variablen berechnet.

Zur Verdeutlichung der Taylorschen Formel für eine Funktion $f(x, y)$ betrachten wir die Spezialfälle $n = 0, 1, 2$:

⊙ **n = 0: Mittelwertsatz**

$$\begin{aligned} f(x, y) &= \sum_{n=0}^0 \frac{1}{n!} \left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n f \Big|_{(x_0, y_0)} + R_0(x, y) \\ &= f(x_0, y_0) + R_0(x, y). \end{aligned}$$

Dies ist der sog. **Mittelwertsatz für Funktionen mehrerer Variablen**.

⊙ **n = 1: Linearisierung**

$$f(x, y) = \sum_{n=0}^1 \frac{1}{n!} \left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n f \Big|_{(x_0, y_0)} + R_1(x, y)$$

$$f(x, y) = f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) + R_1(x, y)$$

Diese Formel wird zur **Linearisierung von Funktionen** verwendet, indem $f(x, y)$ durch das Polynom auf der rechten Seite ersetzt wird.

⊙ **n = 2: Quadratische Näherung**

$$\begin{aligned} f(x, y) &= \sum_{n=0}^2 \frac{1}{n!} \left[(x - x_0) \frac{\partial}{\partial x} + (y - y_0) \frac{\partial}{\partial y} \right]^n f \Big|_{(x_0, y_0)} + R_2(x, y) \\ &= f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) + \\ &\quad + \frac{1}{2!} \left[(x - x_0)^2 \frac{\partial^2}{\partial x^2} + 2(x - x_0)(y - y_0) \frac{\partial}{\partial x} \frac{\partial}{\partial y} \right. \\ &\quad \left. + (y - y_0)^2 \frac{\partial^2}{\partial y^2} \right] f \Big|_{(x_0, y_0)} + R_2(x, y) \end{aligned}$$

$$\begin{aligned} f(x, y) &= f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) \\ &\quad + \frac{1}{2} \left((x - x_0)^2 f_{xx}(x_0, y_0) + 2(x - x_0)(y - y_0) f_{xy}(x_0, y_0) \right. \\ &\quad \left. + (y - y_0)^2 f_{yy}(x_0, y_0) \right) + R_2(x, y) \end{aligned}$$

Bemerkung: Führen wir den Richtungsvektor $\vec{n} = \begin{pmatrix} x - x_0 \\ y - y_0 \end{pmatrix}$ ein und definieren die *Hessesche Matrix*

$$H := \begin{pmatrix} f_{xx}(x_0, y_0) & f_{xy}(x_0, y_0) \\ f_{xy}(x_0, y_0) & f_{yy}(x_0, y_0) \end{pmatrix}$$

kann man die Formel für die quadratische Näherung schreiben in der Form

$$f(x, y) - f(x_0, y_0) = \vec{n} \cdot \operatorname{grad} f \Big|_{(x_0, y_0)} + \frac{1}{2} \vec{n}^t H \vec{n} + R_2(x, y) .$$

Beispiel 10.20 (Taylor-Polynom): Man berechne für die Funktion

$$f(x, y) = \sin(x^2 + 2y)$$

an der Stelle $(x_0, y_0) = (0, \frac{\pi}{4})$ das Taylor-Polynom bis zur Ordnung 2.

(1) Die partiellen Ableitungen bis zur Ordnung 2 lauten

$$f(x, y) = \sin(x^2 + 2y) \qquad f(0, \frac{\pi}{4}) = 1$$

$$\begin{aligned} f_x(x, y) &= 2x \cos(x^2 + 2y) & f_x(0, \frac{\pi}{4}) &= 0 \\ f_y(x, y) &= 2 \cos(x^2 + 2y) & f_y(0, \frac{\pi}{4}) &= 0 \end{aligned}$$

$$\begin{aligned} f_{xx}(x, y) &= -4x^2 \sin(x^2 + 2y) + 2 \cos(x^2 + 2y) & f_{xx}(0, \frac{\pi}{4}) &= 0 \\ f_{yy}(x, y) &= -4 \sin(x^2 + 2y) & f_{yy}(0, \frac{\pi}{4}) &= -4 \\ f_{xy}(x, y) &= -4x \sin(x^2 + 2y) & f_{xy}(0, \frac{\pi}{4}) &= 0. \end{aligned}$$

In die Formel für das Taylor-Polynom bis zur Ordnung 2 eingesetzt folgt

$$\begin{aligned} f(x, y) &= 1 + \frac{1}{2} (y - y_0)^2 \cdot (-4) + R_2(x, y) \\ &= 1 - 2 \left(y - \frac{\pi}{4}\right)^2 + R_2(x, y). \end{aligned}$$

(2) Berechnung des Restgliedes $R_2(x, y)$: Das Restglied lautet mit einem nicht näher bekannten Punkt (ξ, η) auf der Verbindungsgeraden von (x, y) und $(0, \frac{\pi}{4})$

$$\begin{aligned} R_2(x, y) &= \frac{1}{3!} \left\{ f_{xxx}(\xi, \eta) (x - x_0)^3 + 3 f_{xxy}(\xi, \eta) (x - x_0)^2 (y - y_0) \right. \\ &\quad \left. + 3 f_{xyy}(\xi, \eta) (x - x_0) (y - y_0)^2 + f_{yyy}(\xi, \eta) (y - y_0)^3 \right\} \end{aligned}$$

Zur Abschätzung von $R_2(x, y)$ bestimmen wir die partiellen Ableitungen 3. Ordnung

$$\begin{aligned} f_{xxx}(x, y) &= -12x \sin(x^2 + 2y) - 8x^3 \cos(x^2 + 2y) \\ f_{xxy}(x, y) &= -8x^2 \cos(x^2 + 2y) - 4 \sin(x^2 + 2y) \end{aligned}$$

$$f_{xyy}(x, y) = -8x \cos(x^2 + 2y)$$

$$f_{yyy}(x, y) = -8 \cos(x^2 + 2y).$$

Sei $\xi \in [0, x]$ und $\eta \in [\frac{\pi}{4}, y]$ beliebig. Wenn man von $x > 0$ und $y > \frac{\pi}{2}$ ausgeht, dann gelten die Abschätzungen

$$\begin{aligned} |f_{xxx}(\xi, \eta)| &\leq 12\xi + 8\xi^3 &\leq 12x + 8x^3 \\ |f_{xxy}(\xi, \eta)| &\leq 8\xi^2 + 4 &\leq 8x^2 + 4 \\ |f_{xyy}(\xi, \eta)| &\leq 8\xi &\leq 8x \\ |f_{yyy}(\xi, \eta)| &\leq 8. \end{aligned}$$

Daher erhalten wir eine obere Schranke für $R_2(x, y)$:

$$\begin{aligned} \Rightarrow |R_2(x, y)| &\leq \frac{1}{3!} \left\{ |f_{xxx}(\xi, \eta)| x^3 + 3 |f_{xxy}(\xi, \eta)| x^2 \left(y - \frac{\pi}{4}\right) \right. \\ &\quad \left. + 3 |f_{xyy}(\xi, \eta)| x \left(y - \frac{\pi}{4}\right)^2 + |f_{yyy}(\xi, \eta)| \left(y - \frac{\pi}{4}\right)^3 \right\} \\ &\leq \frac{1}{6} \left\{ 4x^4 (2x^2 + 3) + 12x^2 \left(y - \frac{\pi}{4}\right) (2x^2 + 1) \right. \\ &\quad \left. + 24x^2 \left(y - \frac{\pi}{4}\right)^2 + 8 \left(y - \frac{\pi}{4}\right)^3 \right\}. \end{aligned}$$

(3) Zahlenbeispiel: $(x, y) = (0.05, \frac{\pi}{4} + 0.05)$

$$\begin{aligned} f(0.05, \frac{\pi}{4} + 0.05) &= 0.99475 && \text{(Exakter Wert)} \\ f_T(0.05, \frac{\pi}{4} + 0.05) &= 0.995 && \text{(Taylor-Polynom)} \\ R_2(0.05, \frac{\pi}{4} + 0.05) &\leq 0.000455 && \text{(Abgeschätzter Fehler)} \square \end{aligned}$$



Auf der CD-Rom befindet sich die Prozedur **taylor_ani**. Diese Prozedur stellt eine Funktion von zwei Variablen zusammen mit ihrer Taylor-Entwicklung mit wachsender Ordnung graphisch in einer dreidimensionalen Animation dar.

Anwendung: Der Entwicklungspunkt (x_0, y_0) wird festgehalten und (x, y) ist ein beliebiger Punkt in der Nähe von (x_0, y_0) . Aus den Spezialfällen für $n = 1$ und $n = 2$ ergeben sich wichtige *Näherungsausdrücke*.

(1) **Linearisierung:**

$$f(x, y) \approx f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0).$$

Die rechte Seite ist eine lineare Funktion in den Variablen x und y . Das Schaubild dieser linearen Funktion stellt eine Ebene im dreidimensionalen Raum dar, die mit f den gemeinsamen Punkt $(x_0, y_0, f(x_0, y_0))$ hat und den Graphen von f dort berührt. Man nennt diese Ebene die sog. *Tangentialebene* an den Graphen von f im Punkte (x_0, y_0) ($\rightarrow$ totale Differenzierbarkeit 10.3.2).

(2) Quadratische Näherung:

$$\begin{aligned}
 f(x, y) \approx & f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) \\
 & + \frac{1}{2} (x - x_0)^2 f_{xx}(x_0, y_0) + (x - x_0)(y - y_0) f_{xy}(x_0, y_0) \\
 & + \frac{1}{2} (y - y_0)^2 f_{yy}(x_0, y_0).
 \end{aligned}$$

Die rechte Seite ist eine quadratische Funktion in den Variablen x und y . Im Allgemeinen ist die quadratische Näherung für Funktionen besser als die lineare Näherung: Der Graph von f wird nicht durch eine Ebene, sondern durch eine gekrümmte Fläche angenähert, die zusätzlich die gleiche "Krümmung" wie die Funktion im Punkte (x_0, y_0) besitzt.

Beispiel 10.21. Für die Funktion

$$f(x, y) = x \cos(x + y) + (y - 1)^2 e^{-x^2}$$

lautet die lineare Näherung am Entwicklungspunkt $(x_0, y_0) = (0, 0)$

$$f(x, y) \approx 1 + x - 2y$$

und die quadratische

$$f(x, y) \approx 1 + x - 2y - x^2 + y^2.$$

□

Anwendungsbeispiel 10.22. Die Schwingungsdauer eines Pendels der Länge l beträgt

$$T = 2\pi \sqrt{\frac{l}{g}}$$

($g = 9.81 \frac{m}{s^2}$). Gesucht ist ein linearer Ausdruck in g und l , der eine gute Näherung für T darstellt, wenn l nur wenig von 1 und g nur wenig von 9.81 abweicht:

$$\frac{\partial T}{\partial l} = \frac{\pi}{\sqrt{lg}}; \quad \frac{\partial T}{\partial g} = -\pi \sqrt{\frac{l}{g^3}}$$

$$\begin{aligned}
 T & \approx T(l_0, g_0) + \frac{\partial T}{\partial l}(l_0, g_0) \cdot (l - l_0) + \frac{\partial T}{\partial g}(l_0, g_0) \cdot (g - g_0) \\
 & \approx 2.006 + 1.0030(l - 1) - 0.1022(g - 9.81). \quad \square
 \end{aligned}$$



Auf der CD-Rom sind mehrere MAPLE-Prozeduren enthalten, um die **Tangentialebene** einer Funktion von zwei Variablen zu berechnen, die **Taylor-Reihe** einer Funktion zusammen mit der Funktion dreidimensional darzustellen, den **Gradienten** und die **Richtungsableitung** zu bestimmen zu visualisieren.

10.4 Anwendungen der Differenzialrechnung

Wir werden in diesem Abschnitt einige wichtige Anwendungen der Taylorschen Formel behandeln: Das totale Differenzial als lineare Näherung, die Fehlerrechnung, die Theorie der Maxima und Minima bei Funktionen von zwei Variablen sowie die Bestimmung von Ausgleichsfunktionen insbesondere der Regressionsgeraden.

Hinweis: Alle Themen werden durch MAPLE-Prozeduren ergänzt, die sich zusätzlich auf der CD-Rom befinden.

► 10.4.1 Das Differenzial als lineare Näherung

Wir untersuchen das Verhalten der Funktion $z = f(x, y)$ in unmittelbarer Umgebung des Punktes (x_0, y_0) , indem wir den Zuwachs der Tangentialebene dz mit dem Zuwachs der Funktion Δz vergleichen.

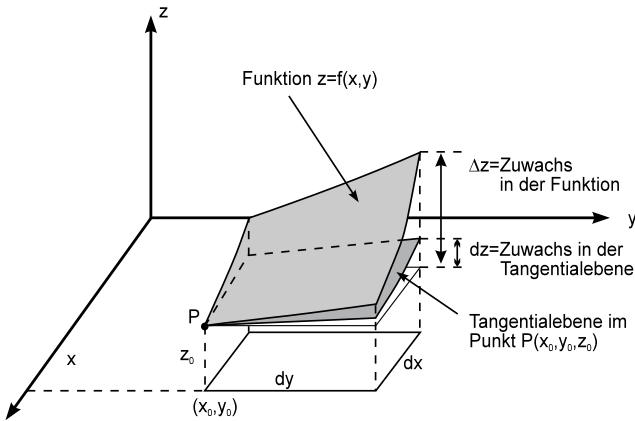


Abb. 10.24. Zum Begriff des vollständigen Differenzials

Die Tangentialebene der Funktion $z = f(x, y)$ ist im Punkte (x_0, y_0) nach Abschnitt 10.3.2 bestimmt durch

$$z_t(x, y) = f(x_0, y_0) + (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) .$$

Im Punkte $(x_0 + dx, y_0 + dy)$ hat die Tangentialebene den Wert

$$z_t(x_0 + dx, y_0 + dy) = f(x_0, y_0) + dx f_x(x_0, y_0) + dy f_y(x_0, y_0) .$$

Die Änderung der Tangentialebene dz ist daher

$$dz = z_t(x_0 + dx, y_0 + dy) - z_t(x_0, y_0) = dx f_x(x_0, y_0) + dy f_y(x_0, y_0) .$$

Wir bezeichnen

dx, dy : unabhängiges Differenzial

dz : abhängiges Differenzial (= Änderung der Tangentialebene)

und definieren

Definition: (Totales Differenzial einer Funktion mit zwei Variablen). Das totale Differenzial einer Funktion $z = f(x, y)$ im Punkte (x_0, y_0) ist

$$dz := f_x(x_0, y_0) dx + f_y(x_0, y_0) dy .$$

Es beschreibt die Änderung der Tangentialebene im Punkte (x_0, y_0) , wenn man vom Punkt (x_0, y_0) zum Punkt $(x_0 + dx, y_0 + dy)$ übergeht. Statt dz schreibt man auch df .

Beispiel 10.23. Gesucht ist das totale Differenzial der Funktion

$$f(x, y) = x \ln(x + y)$$

im Punkte $(x_0, y_0) = (\frac{1}{2}, \frac{1}{2})$.

$$\left. \begin{array}{l} \frac{\partial f}{\partial x} = \ln(x + y) + x \cdot \frac{1}{x + y} \\ \frac{\partial f}{\partial y} = x \cdot \frac{1}{x + y} \end{array} \right\} df = \left(\ln(x + y) + \frac{x}{x + y} \right) dx + \frac{x}{x + y} dy .$$

$$\Rightarrow df(x_0, y_0) = \frac{1}{2} dx + \frac{1}{2} dy .$$

□

Wir vergleichen die **Änderung der Tangentialebene** dz mit der **Änderung der Funktion** Δz , indem wir für die Funktion $z = f(x, y)$ linearisieren, d.h. die Taylorsche Formel für $n = 1$ verwenden:

$$f(x, y) - f(x_0, y_0) = (x - x_0) f_x(x_0, y_0) + (y - y_0) f_y(x_0, y_0) + R_1(x, y) .$$

Die Änderung der Funktion Δz von (x_0, y_0) nach $(x_0 + dx, y_0 + dy)$ ist

$$\begin{aligned} \Delta z &= f(x, y) - f(x_0, y_0) = f(x_0 + dx, y_0 + dy) - f(x_0, y_0) \\ &= dx f_x(x_0, y_0) + dy f_y(x_0, y_0) + R_1(x_0 + dx, y_0 + dy) \\ &= dz + R_1(x_0 + dx, y_0 + dy) . \end{aligned}$$

Sie stimmt mit der Änderung der Tangentialebene dz

$$dz = dx f_x(x_0, y_0) + dy f_y(x_0, y_0)$$

bis auf den Term $R_1(x_0 + dx, y_0 + dy)$ überein. Für $(dx, dy) \rightarrow (0, 0)$ geht $R_1(x_0 + dx, y_0 + dy) \rightarrow R_1(x_0, y_0) = 0$, so dass für kleine dx, dy gilt

$$dz \approx \Delta z .$$

□

➤ **Totales Differenzial von Funktionen mit n Variablen**

Der Begriff des totalen Differenzials überträgt sich direkt auf Funktionen mit mehr als zwei Variablen:

Definition: (Totales Differenzial einer Funktion mit n Variablen).
Unter dem totalen Differenzial einer Funktion

$$y = f(x_1, \dots, x_n)$$

versteht man den Differenzialausdruck

$$\begin{aligned} dy &:= f_{x_1} dx_1 + f_{x_2} dx_2 + \dots + f_{x_n} dx_n \\ &= \frac{\partial f}{\partial x_1} dx_1 + \frac{\partial f}{\partial x_2} dx_2 + \dots + \frac{\partial f}{\partial x_n} dx_n. \end{aligned}$$

Bemerkungen:

- (1) Statt dy schreibt man auch df .
- (2) Das totale Differenzial beschreibt näherungsweise, wie sich der Funktionswert ändert, wenn sich die Variablen geringfügig um dx_i ($i = 1, \dots, n$) ändern: $\Delta y \approx dy$.

Beispiele 10.24.

- ① Gesucht ist das totale Differenzial der Funktion

$$f(x, y, z) = x e^{x y + 4 z}$$

im Punkte $(x_0, y_0, z_0) = (1, 0, 1)$:

$$\frac{\partial f}{\partial x} = e^{x y + 4 z} + x e^{x y + 4 z} y \quad \hookrightarrow f_x(1, 0, 1) = e^4$$

$$\frac{\partial f}{\partial y} = x^2 e^{x y + 4 z} \quad \hookrightarrow f_y(1, 0, 1) = e^4$$

$$\frac{\partial f}{\partial z} = 4 x e^{x y + 4 z} \quad \hookrightarrow f_z(1, 0, 1) = 4 e^4.$$

Mit der Formel des totalen Differenzials

$$df = f_x dx + f_y dy + f_z dz$$

erhalten wir für df

$$\begin{aligned} df(x, y, z) &= (e^{x y + 4 z} + x y e^{x y + 4 z}) dx + x^2 e^{x y + 4 z} dy + 4 x e^{x y + 4 z} dz \\ df(1, 0, 1) &= e^4 dx + e^4 dy + 4 e^4 dz. \end{aligned}$$

② Ein ideales Gas genügt für 1 Mol der Zustandsgleichung

$$p(V, T) = R \cdot \frac{T}{V}.$$

Das totale Differenzial dieser Funktion lautet

$$dp = \frac{\partial p}{\partial V} dV + \frac{\partial p}{\partial T} dT = -R \frac{T}{V^2} dV + \frac{R}{V} dT.$$

Es beschreibt näherungsweise die Änderung des Gasdrucks p bei einer geringfügigen Änderung des Volumens um dV und gleichzeitig der Temperatur um dT (vgl. Beispiel 10.10 ⑥). $\square$

③ Linearisierung von Funktionen mit zwei Variablen

Für eine Funktion $z = f(x, y)$ gilt für kleine dx und dy näherungsweise

$$dz \approx \Delta z,$$

d.h. für kleine dx und dy kann die Änderung der Funktion über die Änderung der Tangentialebene (totales Differenzial) angenähert werden. Man nennt dieses Vorgehen die Linearisierung der Funktion $z = f(x, y)$ und setzt

$$\Delta z = f(x, y) - f(x_0, y_0) \approx dz = f_x(x_0, y_0) dx + f_y(x_0, y_0) dy.$$

Über die Beziehung $dx = x - x_0$ und $dy = y - y_0$ folgt insgesamt

$$f(x, y) \approx f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0)$$

(Linearisierung der Funktion $z = f(x, y)$.)

Oftmals wird die Linearisierung benutzt, um Differenzen der Form

$$z(x_0 + \Delta x, y_0 + \Delta y) - z(x_0, y_0)$$

näherungsweise zu bestimmen:

$$z(x_0 + \Delta x, y_0 + \Delta y) - z(x_0, y_0) \approx \left. \frac{\partial z}{\partial x} \right|_{(x_0, y_0)} \Delta x + \left. \frac{\partial z}{\partial y} \right|_{(x_0, y_0)} \Delta y$$

$$z(x, y) - z(x_0, y_0) \approx \left. \frac{\partial z}{\partial x} \right|_{(x_0, y_0)} (x - x_0) + \left. \frac{\partial z}{\partial y} \right|_{(x_0, y_0)} (y - y_0). \quad (2)$$

Mit einer verallgemeinerten Formel lassen sich Funktionen von n Variablen linearisieren:

Linearisierung von Funktionen mit n Variablen: In der Umgebung eines Entwicklungspunktes (Arbeitspunktes) $(x_1^0, \dots, x_n^0)$ kann eine nichtlineare Funktion

$$f(x_1, \dots, x_n)$$

näherungsweise durch die *lineare* Funktion

$$y = f(x_1^0, \dots, x_n^0) + \left(\frac{\partial f}{\partial x_1} \right)_0 (x_1 - x_1^0) + \dots + \left(\frac{\partial f}{\partial x_n} \right)_0 (x_n - x_n^0)$$

ersetzt werden. Die partiellen Ableitungen $\left(\frac{\partial f}{\partial x_i} \right)_0$ müssen am Entwicklungspunkt $(x_1^0, \dots, x_n^0)$ ausgewertet werden:

$$\left(\frac{\partial f}{\partial x_i} \right)_0 = \frac{\partial}{\partial x_i} f(x_1^0, \dots, x_n^0) .$$

Beispiele 10.25:

- ① Gesucht ist die Linearisierung der Funktion

$$f(x, y) = 5x e^{x-4y^2}$$

am Entwicklungspunkt $(x_0, y_0) = (1, \frac{1}{2})$:

Mit den partiellen Ableitungen

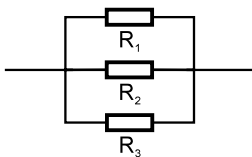
$$f_x(x, y) = 5e^{x-4y^2} + 5x e^{x-4y^2} \quad \hookrightarrow \quad f_x\left(1, \frac{1}{2}\right) = 5 + 5 = 10$$

$$f_y(x, y) = -40xy e^{x-4y^2} \quad \hookrightarrow \quad f_y\left(1, \frac{1}{2}\right) = -20$$

erhält man die lineare Approximation

$$\Rightarrow z_l(x, y) = 5 + 10(x - 1) - 20\left(y - \frac{1}{2}\right) .$$

- ② Der Gesamtwiderstand R einer Parallelschaltung aus drei Ohmschen Widerständen R_1 , R_2 und R_3 wird durch die Formel



$$\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3}$$

berechnet. Wir linearisieren die Funktion in einer Umgebung der Werte $R_1 = 100 \Omega$, $R_2 = 200 \Omega$, $R_3 = 500 \Omega$. Dazu berechnen wir die partiellen

Ableitungen von

$$R = \frac{R_1 R_2 R_3}{R_2 R_3 + R_1 R_3 + R_1 R_2}$$

mit der Quotientenregel

$$\frac{\partial R}{\partial R_1} = \frac{R_2^2 R_3^2}{(R_2 R_3 + R_1 R_3 + R_1 R_2)^2}, \quad \frac{\partial R}{\partial R_2} = \frac{R_1^2 R_3^2}{(R_2 R_3 + R_1 R_3 + R_1 R_2)^2},$$

$$\frac{\partial R}{\partial R_3} = \frac{R_1^2 R_2^2}{(R_2 R_3 + R_1 R_3 + R_1 R_2)^2},$$

und setzen den Entwicklungspunkt $(R_1^0, R_2^0, R_3^0) = (100, 200, 500)$ ein:

$$\left(\frac{\partial R}{\partial R_1} \right)_0 = 0.3460, \quad \left(\frac{\partial R}{\partial R_2} \right)_0 = 0.086, \quad \left(\frac{\partial R}{\partial R_3} \right)_0 = 0.014.$$

Die Linearisierung lautet damit

$$\begin{aligned} R_L &= R(R_1^0, R_2^0, R_3^0) + \left(\frac{\partial R}{\partial R_1} \right)_0 (R_1 - R_1^0) + \left(\frac{\partial R}{\partial R_2} \right)_0 (R_2 - R_2^0) \\ &\quad + \left(\frac{\partial R}{\partial R_3} \right)_0 (R_3 - R_3^0) \end{aligned}$$

$$R_L = 58.8235 + 0.3460 (R_1 - 100) + 0.086 (R_2 - 200) + 0.014 (R_3 - 500).$$

□

► 10.4.2 Fehlerrechnung

Eine Anwendung des totalen Differenzials tritt bei der Fehlerrechnung auf, die überall dort eine Rolle spielt, wo mit ungenauen Messwerten gearbeitet wird:

Problemstellung: Eine physikalische Größe y hänge nach einem bekannten Gesetz von n unabhängigen Größen $x_1, \dots, x_n$ ab:

$$y = f(x_1, \dots, x_n).$$

Zur Auswertung von y müssen die Werte von $x_1, \dots, x_n$ gemessen werden, was nur mit begrenzter Genauigkeit möglich ist. Die gemessenen Werte bezeichnen wir mit $x_1^0, \dots, x_n^0$, die Messfehler (Toleranzen) mit $\Delta x_1, \dots, \Delta x_n$. Es stellt sich die Frage: Um wieviel weicht der aus den Messwerten mit Toleranzen errechnete Wert $f(x_1^0 + \Delta x_1, \dots, x_n^0 + \Delta x_n)$ maximal vom Wert $f(x_1^0, \dots, x_n^0)$ ab? Die Differenz

$$\Delta f = f(x_1^0 + \Delta x_1, \dots, x_n^0 + \Delta x_n) - f(x_1^0, \dots, x_n^0)$$

heißt der **absolute Fehler** von f . Gesucht ist eine Abschätzung von Δf , wenn man weiß, wie groß die Beträge der Messfehler Δx_i höchstens sind.

Für kleine $dx_i = \Delta x_i$ stimmt das totale Differenzial df in etwa mit der Änderung der Funktion Δf überein:

$$\Delta f \approx df = \left(\frac{\partial f}{\partial x_1} \right)_0 \cdot dx_1 + \dots + \left(\frac{\partial f}{\partial x_n} \right)_0 \cdot dx_n .$$

Diese Formel verwendet man, um den Einfluss **kleiner** Fehler auf das Resultat zu berechnen. Da für die Messfehler in der Regel eine Toleranz $\pm \Delta x_i$ angegeben wird, erhält man für den Fehler in linearer Näherung eine Obergrenze durch

$$|\Delta f| \approx |df| \leq \left| \left(\frac{\partial f}{\partial x_1} \right)_0 \right| |\Delta x_1| + \dots + \left| \left(\frac{\partial f}{\partial x_n} \right)_0 \right| |\Delta x_n| = \bar{f} .$$

Man bezeichnet $\bar{f}$ als **absoluten Fehler in linearer Näherung**. Dabei sind die partiellen Ableitungen $\left(\frac{\partial f}{\partial x_i} \right)_0$ für die Messwerte $(x_1^0, \dots, x_n^0)$ auszuwerten und $|\Delta x_i|$ geben die maximalen Fehlerschranken dieser Messwerte an.

Oftmals sind auch die **relativen Fehler** von Interesse

$$\left| \frac{\bar{f}}{f_0} \right|, \quad \left| \frac{\Delta x_i}{x_i^0} \right| \quad (i = 1, \dots, n) ,$$

die in der Regel auch in Prozenten angegeben werden: $100 \left| \frac{\bar{f}}{f_0} \right| \cdot \%$.

Zusammenfassung: (Fehlerfortpflanzung nach Gauß). Die Größe y hänge von unabhängigen Größen $x_1, \dots, x_n$ gemäß dem Gesetz

$$y = f(x_1, \dots, x_n)$$

ab. Werden die Einzelgrößen $x_1, \dots, x_n$ gemessen, so liegen in der Form

$$x_i^0 \pm \Delta x_i$$

vor. Dabei ist x_i^0 der Mittelwert der Größen x_i , Δx_i die Fehlertoleranzen. Dann bestimmt man den Fehler in y aufgrund der Messungenauigkeiten von x_i näherungsweise durch

$$\bar{y} = \left| \left(\frac{\partial f}{\partial x_1} \right)_0 \right| \cdot |\Delta x_1| + \dots + \left| \left(\frac{\partial f}{\partial x_n} \right)_0 \right| \cdot |\Delta x_n| ,$$

wenn die partiellen Ableitungen $\left(\frac{\partial f}{\partial x_i} \right)_0$ an den Stellen $(x_1^0, \dots, x_n^0)$ ausgewertet werden. Die Größe y hat den Wert

$$y = f(x_1^0, \dots, x_n^0) \pm \bar{y} .$$

$\bar{y}$ heißt **absoluter maximaler Fehler** in linearer Näherung.

⚠ Achtung: Da man statt Δf den Fehler in linearer Näherung $|df|$ angibt, ist diese Vorgehensweise nur für kleine Abweichungen Δx_i sinnvoll. Für große Fehlertoleranzen ist die Änderung der Funktion Δf nicht mit der Änderung der Tangentialebene vergleichbar!

Anwendungsbeispiel 10.26.

Die Spannung U an den Enden eines elektrischen Widerstandes hängt mit der Stromstärke I eines ihn durchfließenden Gleichstromes durch das Ohmsche Gesetz zusammen



Ist U mit dU und I mit dI fehlerhaft gemessen, so hat man als Oberschranke für den Fehler in R

$$\Delta R \approx dR = \frac{\partial R}{\partial U} \cdot dU + \frac{\partial R}{\partial I} \cdot dI$$

$$\Rightarrow |\Delta R| \leq \left| \frac{1}{I} dU \right| + \left| -\frac{U}{I^2} dI \right| = \bar{R}.$$

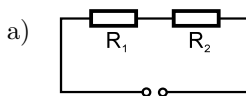
Zahlenbeispiel: $U = (110 \pm 2) \text{ V}$, $I = (20 \pm 0.5) \text{ A}$ ergibt

$$\bar{R} = \left(\frac{1}{20} \cdot 2 + \frac{110}{400} \cdot 0.5 \right) \Omega = 0.2375 \Omega \Rightarrow R = (5.5 \pm 0.24) \Omega.$$

Der relative Fehler beträgt $\frac{\bar{R}}{R} = \frac{0.24}{5.5} = 0.044 = 4.4 \%$. □

Anwendungsbeispiel 10.27.

Ein Widerstand der Größe R_0 mit 10 % Toleranz wird mit einem 5 mal größeren, zweiten Widerstand von 2 % Toleranz a) seriell b) parallel geschaltet. Man bestimme für beide Fälle den Gesamtwiderstand R_{ges} sowie den relativen Fehler.



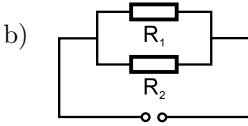
$$R_{ges} = R_1 + R_2 (= 6 R_0)$$

$$\Rightarrow \frac{\partial R_{ges}}{\partial R_1} = 1; \quad \frac{\partial R_{ges}}{\partial R_2} = 1.$$

$$\begin{aligned} \bar{R} &= \left(\frac{\partial R_{ges}}{\partial R_1} \right)_0 \cdot |\Delta R_1| + \left(\frac{\partial R_{ges}}{\partial R_2} \right)_0 \cdot |\Delta R_2| \\ &= 10 \% R_0 + 10 \% R_0 = 20 \% R_0. \end{aligned}$$

Damit ergibt sich der relative Fehler zu

$$\frac{\bar{R}}{R_{ges}} = \frac{20 \% R_0}{6 R_0} \approx 3.3 \%.$$



$$\frac{1}{R_{ges}} = \frac{1}{R_1} + \frac{1}{R_2}$$

$$\Rightarrow R_{ges} = \frac{R_1 \cdot R_2}{R_1 + R_2} (= \frac{5 R_0^2}{6 R_0} = \frac{5}{6} R_0).$$

$$\frac{\partial R_{ges}}{\partial R_1} = \frac{R_2^2}{(R_1 + R_2)^2}; \quad \frac{\partial R_{ges}}{\partial R_2} = \frac{R_1^2}{(R_1 + R_2)^2}.$$

$$\begin{aligned} \bar{R} &= \frac{R_2^2}{(R_1 + R_2)^2} |\Delta R_1| + \frac{R_1^2}{(R_1 + R_2)^2} |\Delta R_2| \\ &= \frac{25}{36} 10 \% R_0 + \frac{1}{36} 10 \% R_0 = \frac{260}{36} \% R_0 = 7.2 \% R_0. \end{aligned}$$

Der relative Gesamtfehler beträgt hierbei $\frac{\bar{R}}{R} = \frac{7.2 \% R_0}{\frac{5}{6} R_0} = 8.66 \%$. $\square$

➤ Fehlerformel für Potenzgesetze

Viele funktionale Zusammenhänge in den Naturwissenschaften haben eine einfache Beschreibung der Form

$$y = f(x_1, x_2) = \alpha x_1^{c_1} \cdot x_2^{c_2}. \quad (*)$$

Für solche Funktionen ist

$$dy = \frac{\partial f}{\partial x_1} dx_1 + \frac{\partial f}{\partial x_2} dx_2 = \alpha c_1 x_1^{c_1-1} x_2^{c_2} dx_1 + \alpha c_2 x_1^{c_1} x_2^{c_2-1} dx_2.$$

Damit folgt für den maximalen relativen Fehler

$$\frac{\bar{y}}{y} = |c_1| \left| \frac{dx_1}{x_1} \right| + |c_2| \left| \frac{dx_2}{x_2} \right|.$$

Der maximale relative Fehler der Größe $y = \alpha x_1^{c_1} \cdot x_2^{c_2}$ setzt sich additiv aus den relativen Fehlern der Einzelgrößen x_i zusammen. Als Koeffizienten treten die Beträge der Exponenten auf.

Folgerung: Beim Rechnen mit Dezimalzahlen bedeutet die Angabe der Zahl $a = 1.23$ in der Regel: Der genaue Wert von a weicht um höchstens eine halbe Einheit der letzten angegebenen Stelle ab; liegt also zwischen 1.225 und 1.235 . Man erkennt dabei, dass die Dezimalzahlen $123; 12.3; 1.23; \dots; 1.23 \cdot 10^k$ ($k \in \mathbb{Z}$) alle den gleichen relativen Fehler

$$100 \% \cdot \frac{0.005 \cdot 10^k}{1.23 \cdot 10^k} \approx 0.41 \%$$

besitzen. Der relative Fehler einer Dezimalzahl ist umso kleiner, je größer die Anzahl der geltenden Stellen ist. $\square$

10.4.3 Lokale Extrema bei Funktionen mit mehreren Variablen

Im Folgenden betrachten wir eine Funktion $z = f(x, y)$ von zwei Variablen x und y , deren partiellen Ableitungen bis zweiter Ordnung stetig sind. Wie bei Funktionen einer Variablen suchen wir zur Charakterisierung der Funktion nach solchen Punkten, in denen der Funktionswert von f am größten bzw. kleinsten wird. Bei der Diskussion schränken wir uns auf die *lokalen* Maxima und Minima ein:

Definition: (Lokales Extremum). Eine Funktion $z = f(x, y)$ besitzt an der Stelle $(x_0, y_0) \in \mathbb{D}$ ein **relatives Maximum**, wenn in einer Umgebung des Punktes (x_0, y_0) für alle $(x, y) \neq (x_0, y_0)$ stets gilt

$$f(x_0, y_0) > f(x, y) .$$

Eine Funktion $z = f(x, y)$ besitzt an der Stelle $(x_0, y_0) \in \mathbb{D}$ ein **relatives Minimum**, wenn in einer Umgebung des Punktes (x_0, y_0) für alle $(x, y) \neq (x_0, y_0)$ stets gilt

$$f(x_0, y_0) < f(x, y) .$$

Die relativen Maxima und Minima fasst man unter den Begriff "relative Extremwerte" zusammen. Relative Extremwerte werden manchmal auch als lokale Extremwerte bezeichnet, da die extreme Lage nur in unmittelbarer Umgebung von (x_0, y_0) , nicht aber global zutreffen muss.

Beispiele 10.28 (Mit MAPLE-Worksheet).

① Zeichnet man (mit MAPLE) die Funktion

$$f(x, y) = x^2 + y^2 + 4$$

erkennt man, dass sie im Punkt $(x_0, y_0) = (0, 0)$ ein lokales (sogar globales) Minimum besitzt, siehe Abb. 10.25.

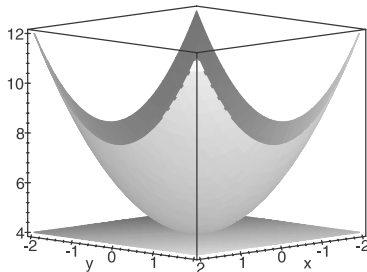


Abb. 10.25. Lokales Minimum der Funktion $f(x, y) = x^2 + y^2 + 4$ im Punkte $(0, 0)$

Neben der Funktion f ist die Tangentialebene im Punkte $(0, 0)$ gezeichnet. Man erkennt, dass sie parallel zur (x, y) -Ebene liegt.

② Die Funktion

$$f(x, y) = e^{-(x^2+y^2)}$$

besitzt im Punkte $(x_0, y_0) = (0, 0)$ ein lokales (sogar globales) Maximum, wie man aus dem Graphen in Abb. 10.26 ablesen kann:

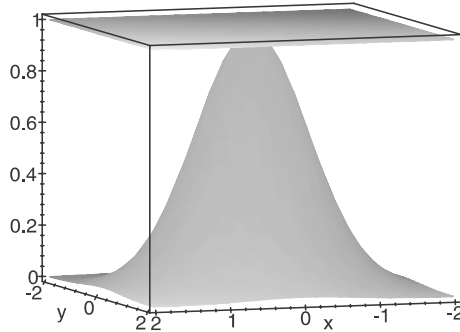


Abb. 10.26. Lokales Maximum der Funktion $f(x, y) = e^{-(x^2+y^2)}$

Wieder liegt die Tangentialebene parallel zur (x, y) -Ebene. □

Für eine stetig differenzierbare Funktion f einer Variablen x besagt die notwendige Bedingung für ein lokales Extremum $x_0 \in \mathbb{ID}$, dass die Tangente im Punkt $(x_0, f(x_0))$ parallel zur x -Achse verläuft: $f'(x_0) = 0$. Hat eine Funktion f von zwei Variablen in (x_0, y_0) ein lokales Extremum, dann ist die Tangentialebene parallel zur (x, y) -Ebene. Die Tangentialebene z_t von f lautet nach Abschnitt 10.3.2

$$z_t(x, y) = f(x_0, y_0) + \frac{\partial f}{\partial x}(x_0, y_0)(x - x_0) + \frac{\partial f}{\partial y}(x_0, y_0)(y - y_0) .$$

Sie liegt parallel zur (x, y) -Ebene, wenn sowohl die partielle Ableitung nach x als auch nach y im Punkt (x_0, y_0) verschwindet.

Satz: Notwendige Bedingung für ein relatives Extremum. In einem relativen Extremum $(x_0, y_0) \in \mathbb{ID}$ besitzt die Funktion $f(x, y)$ eine Tangentialebene parallel zur (x, y) -Ebene:

$$(x_0, y_0) \text{ relatives Extremum} \Rightarrow \frac{\partial f}{\partial x}(x_0, y_0) = 0 \ \& \ \frac{\partial f}{\partial y}(x_0, y_0) = 0 .$$

Ist (x_0, y_0) ein relatives Extremum, dann verschwindet also der Gradient von f in (x_0, y_0) :

$$\text{grad } f(x_0, y_0) = 0.$$

⚠ Achtung: Dieser Satz ist nicht umkehrbar, wie das folgende Beispiel zeigt!:

Beispiel 10.29 (Mit MAPLE-Worksheet). Die Funktion

$$f(x, y) = x \cdot y$$

hat im Punkte $(x_0, y_0) = (0, 0)$ eine waagrechte Tangentialebene:

$$\left. \begin{array}{l} \frac{\partial f}{\partial x}(x, y) = y \Rightarrow \frac{\partial f}{\partial x}(0, 0) = 0 \\ \frac{\partial f}{\partial y}(x, y) = x \Rightarrow \frac{\partial f}{\partial y}(0, 0) = 0 \end{array} \right\} \Rightarrow \text{grad } f(0, 0) = 0.$$

Aber in einer Umgebung des Punktes $(0, 0)$ existieren positive **und** negative Funktionswerte, wie man dem Funktionsgraphen in Abb. 10.27 entnimmt:

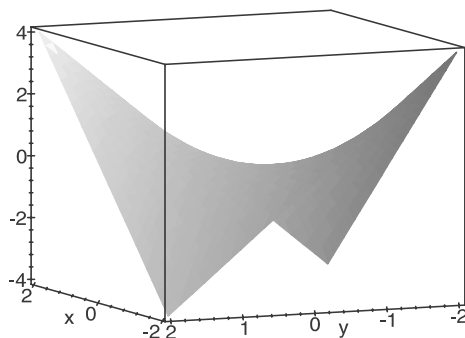


Abb. 10.27. Funktion mit Sattelpunkt

Folglich hat f bei $(0, 0)$ **kein** lokales Extremum. Man nennt diesen Punkt einen *Sattelpunkt*. □

Definition: (Stationärer Punkt). Der Punkt (x_0, y_0) heißt **stationärer Punkt** von f , wenn

$$\text{grad } f(x_0, y_0) = 0.$$

Relative Extrema sind demnach stationäre Punkte, aber nicht alle stationären Punkte sind nach Beispiel 10.29 relative Extrema. Ein stationärer Punkt P ist dadurch charakterisiert, dass die Tangentialebene von f in P parallel zur (x, y) -Ebene verläuft.

Um zu entscheiden, ob in einem stationären Punkt (x_0, y_0) ein lokales Extremum vorliegt, entwickeln wir f nach dem Satz von Taylor in der Umgebung von (x_0, y_0) bis zur Ordnung 2

$$\begin{aligned} f(x, y) = & f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) \\ & + \frac{1}{2!} \left(f_{xx}(x_0, y_0)(x - x_0)^2 + 2f_{xy}(x_0, y_0)(x - x_0)(y - y_0) \right. \\ & \left. + f_{yy}(x_0, y_0)(y - y_0)^2 \right) + R_2(x, y) \end{aligned}$$

Da $f_x(x_0, y_0) = f_y(x_0, y_0) = 0$, entfallen die Ableitungen erster Ordnung; für die Ableitungen zweiter Ordnung schreiben wir

$$\begin{aligned} f(x, y) = & f(x_0, y_0) + \frac{1}{2} f_{xx}(x_0, y_0) \left\{ (x - x_0) + \frac{f_{xy}(x_0, y_0)}{f_{xx}(x_0, y_0)} (y - y_0) \right\}^2 \\ & + \frac{1}{2} \frac{1}{f_{xx}(x_0, y_0)} \{ f_{xx}(x_0, y_0) f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) \} (y - y_0)^2 + R_2(x, y). \end{aligned}$$

Aus dieser Darstellung der Funktion kann man ablesen:

- i) Ist $f_{xx}(x_0, y_0) > 0$ und $f_{xx}(x_0, y_0) f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) > 0$

$$\Rightarrow f(x, y) > f(x_0, y_0) + R_2(x, y) .$$

D.h. in einer Umgebung von (x_0, y_0) sind die Funktionswerte **größer** als $f(x_0, y_0)$. Es liegt also in (x_0, y_0) ein lokales Minimum vor.

- ii) Ist $f_{xx}(x_0, y_0) < 0$ und $f_{xx}(x_0, y_0) f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) > 0$

$$\Rightarrow f(x, y) < f(x_0, y_0) + R_2(x, y) .$$

D.h. in einer Umgebung von (x_0, y_0) sind die Funktionswerte **kleiner** als $f(x_0, y_0)$. Es liegt also in (x_0, y_0) ein lokales Maximum vor.

Folglich entscheidet der Ausdruck

$$f_{xx}(x_0, y_0) f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) > 0,$$

ob in einem stationären Punkt ein Extremum vorliegt:

Satz: Hinreichende Bedingung für ein lokales Extremum.

$f(x, y)$ sei zweimal stetig partiell differenzierbar in $(x_0, y_0) \in \mathbb{D}$. Im Punkt (x_0, y_0) liegt ein lokales Extremum vor, falls

$$(1) \quad f_x(x_0, y_0) = 0, \quad f_y(x_0, y_0) = 0.$$

$$(2) \quad \Delta := f_{xx}(x_0, y_0) \cdot f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) > 0.$$

Für $f_{xx}(x_0, y_0) < 0$ liegt ein **relatives Maximum** vor.

Für $f_{xx}(x_0, y_0) > 0$ liegt ein **relatives Minimum** vor.

Ist hingegen in einem stationären Punkt (x_0, y_0)

$$f_{xx}(x_0, y_0) f_{yy}(x_0, y_0) - f_{xy}^2(x_0, y_0) < 0,$$

so liegt kein Extremum, sondern ein **Sattelpunkt** vor.

Bemerkungen:

- (1) Wie bei Funktionen einer Variablen wird die zweite Ableitung herangezogen, um zu entscheiden, ob ein lokales Extremum vorliegt oder nicht.
- (2) Definiert man die sog. *Hessesche Matrix*

$$\text{Hess } f := \begin{pmatrix} f_{xx} & f_{xy} \\ f_{yx} & f_{yy} \end{pmatrix},$$

entscheidet die Determinante der Hesseschen Matrix im Punkte (x_0, y_0) , ob ein Extremum vorliegt. Es gilt:

$$\begin{array}{ll} \det(\text{Hess } f) < 0 & \text{kein Extremwert, sondern Sattelpunkt.} \\ \det(\text{Hess } f) = 0 & \text{keine Entscheidung möglich, ob an} \\ & \text{der Stelle } (x_0, y_0) \text{ ein Extremum vorliegt.} \\ \det(\text{Hess } f) > 0 & \text{ein lokales Extremum liegt vor.} \end{array}$$

- (3) Auf der CD-Rom befindet sich ein Abschnitt, in dem die **lokalen Extrema** von Funktionen mit mehr als zwei Variablen bestimmt werden. Man bestimmt dann die Extremwerte, indem man den Gradienten gleich Null setzt und die stationären Punkte in die entsprechende Hessesche Matrix einsetzt.

Beispiel 10.30. Die Funktion

$$f(x, y) = x^2 + y^2 + c$$

hat im Punkte $(x_0, y_0) = (0, 0)$ ein lokales Minimum:

(i) Aus $\text{grad } f = 0$ folgen die stationären Punkte:

$$\left. \begin{aligned} f_x(x, y) = 2x &\stackrel{!}{=} 0 \Rightarrow x = 0 \\ f_y(x, y) = 2y &\stackrel{!}{=} 0 \Rightarrow y = 0 \end{aligned} \right\} \Rightarrow (x, y) = (0, 0) \text{ ist stationärer Punkt.}$$

(ii) Einsetzen des stationären Punktes in die Δ -Formel:

$$\left. \begin{aligned} f_{xx}(0, 0) &= 2 \\ f_{yy}(0, 0) &= 2 \\ f_{xy}(0, 0) &= 0 \end{aligned} \right\} \Rightarrow \Delta = f_{xx}(0, 0) \cdot f_{yy}(0, 0) - f_{xy}^2(0, 0) = 4 > 0.$$

$\Rightarrow (0, 0)$ ist relatives Extremum.

Wegen $f_{xx}(0, 0) > 0$ liegt ein relatives Minimum vor. $\square$

Beispiel 10.31. Gesucht sind die lokalen Extrema der Funktion

$$f(x, y) = (x^2 + y^2)^2 - 2(x^2 - y^2).$$

(i) Aus $\text{grad } f = 0$ folgen die stationären Punkte:

$$\begin{aligned} f_x(x, y) &= 4x(x^2 + y^2) - 4x \stackrel{!}{=} 0 \\ f_y(x, y) &= 4y(x^2 + y^2) + 4y \stackrel{!}{=} 0. \end{aligned}$$

Dies ist ein gekoppeltes System nichtlinearer Gleichungen für die Unbekannten x und y . Durch Ausklammern folgt

$$4x(x^2 + y^2 - 1) = 0 \tag{1}$$

$$4y(x^2 + y^2 + 1) = 0. \tag{2}$$

Aus Gleichung (1) folgt $x = 0$ oder $x^2 + y^2 - 1 = 0$.

Setzen wir $x = 0$ in Gleichung (2) ein, gilt: $4y(y^2 + 1) = 0$ bzw. $y = 0$.
 $\Rightarrow (x, y) = (0, 0)$ ist stationärer Punkt.

Für $x^2 + y^2 - 1 = 0$ folgt $x^2 + y^2 = 1$. In Gleichung (2) eingesetzt, gilt:
 $4y \cdot 2 = 0$ bzw. $y = 0$. Eingesetzt wiederum in $x^2 + y^2 - 1 = 0$ folgt
dann $x^2 - 1 = 0$ bzw. $x = \pm 1$.

$\Rightarrow (x, y) = (1, 0)$ und $(x, y) = (-1, 0)$ sind ebenfalls stationäre Punkte.

Damit besitzt die Funktion drei stationäre Punkte: $(0, 0)$; $(1, 0)$ und $(-1, 0)$.
Höchstens an diesen Stellen kann f ein Extremum annehmen.

(ii) Einsetzen der stationären Punkte in die Δ -Formel: Dazu berechnen wir die zweiten partiellen Ableitungen

$$f_{xx}(x, y) = 12x^2 + 4y^2 - 4$$

$$f_{yy}(x, y) = 12y^2 + 4x^2 + 4$$

$$f_{xy}(x, y) = 8xy$$

und setzen die stationären Punkte ein:

$$\Delta(0, 0) = f_{xx}(0, 0) \cdot f_{yy}(0, 0) - f_{xy}^2(0, 0) = -4 \cdot 4 - 0 = -12 < 0.$$

$\Rightarrow$ In $(0, 0)$ liegt **kein** Extremum, sondern ein Sattelpunkt vor.

$$\Delta(1, 0) = f_{xx}(1, 0) \cdot f_{yy}(1, 0) - f_{xy}^2(1, 0) = 8 \cdot 8 - 0 = 64 > 0.$$

$\Rightarrow$ In $(1, 0)$ liegt ein lokales Minimum vor, da $f_{xx}(1, 0) = 8 > 0$.

Wegen der Symmetrie $f(-x, y) = f(x, y)$ liegt auch im Punkt $(-1, 0)$ ein lokales Minimum vor. $\square$

Beispiel 10.32. Die Funktion

$$f(x, y) = c + x^2 - y^2$$

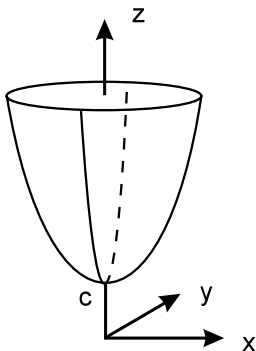
hat in $(0, 0)$ einen Sattelpunkt:

Aus $f_x(x, y) = 2x \stackrel{!}{=} 0$ und $f_y(x, y) = 2y \stackrel{!}{=} 0 \Rightarrow (x, y) = (0, 0)$ ist einziger stationärer Punkt.

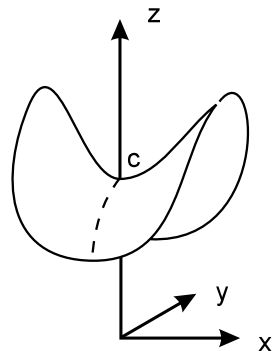
Da

$$f_{xx}(x, y) \cdot f_{yy}(x, y) - f_{xy}^2(x, y) = -4 < 0$$

für alle (x, y) ist $(0, 0)$ **kein** lokales Extremum, sondern ein Sattelpunkt. $\square$



Funktion $f(x, y) = c + x^2 + y^2$



Funktion $f(x, y) = c + x^2 - y^2$

➤ 10.4.4 Ausgleichen von Messfehlern; Regressionsgerade

Eine wichtige Anwendung der Theorie der Extremwerte ist das Ausgleichen von Messfehlern: Durch Messungen, welche die Abhängigkeit einer Größe y von einer anderen Größe x ermittelt, seien n Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$ erfasst worden. Die Aufgabe der Ausgleichsrechnung besteht darin, eine Funktion f zu finden, die sich einerseits den vorliegenden Messpunkten möglichst gut anschmiegt und die andererseits einen möglichst glatten Verlauf besitzt.

Problemstellung: Gemessen wird die Temperaturabhängigkeit eines Ohmschen Widerstandes. Gesucht ist eine Gerade, welche die Messpunkte "geeignet" repräsentiert.

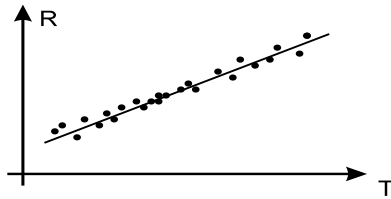


Abb. 10.28. Temperaturabhängigkeit eines Ohmschen Widerstandes

➤ Prinzip der kleinsten Quadrate

Die gesuchte Funktion f soll die Eigenschaft besitzen, dass die Abstandsquadrate der Messpunkte zur Ausgleichsfunktion minimal werden:

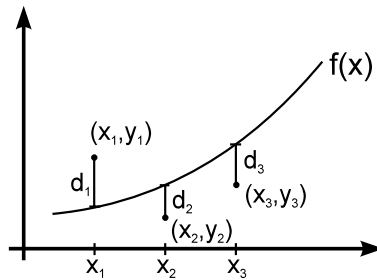


Abb. 10.29. Abstände d_i zur Ausgleichsfunktion

Der Abstand des Messpunkts y_i zur Ausgleichsfunktion f im Punkte x_i ist

$$d_i = y_i - f(x_i) .$$

Die Summe über alle Abstandsquadrate ist daher

$$\sum_{i=1}^n d_i^2 = \sum_{i=1}^n (y_i - f(x_i))^2 .$$

Je nach vermutetem funktionalen Zusammenhang wählt man sich einen Funktionsansatz.

Tabelle 10.1: Ausgleichsfunktionen

Ausgleichsfunktion		Parameter
lineare Funktion	$f(x) = ax + b$	a, b
quadratische Funktion	$f(x) = ax^2 + bx + c$	a, b, c
Polynomfunktion	$f(x) = a_n x^n + \dots + a_1 x + a_0$	$a_n, a_{n-1}, \dots, a_0$
Potenzfunktion	$f(x) = ax^b$	a, b
Exponentialfunktion	$f(x) = ae^{bx}$	a, b

In jeder Ansatzfunktion $f(x)$ sind Parameter enthalten, die so bestimmt werden müssen, dass die Summe der Abstandsquadrate minimal wird (*Prinzip der kleinsten Quadrate*):

$$F(a, b, c, \dots) := \sum_{i=1}^n (y_i - f(x_i))^2 \rightarrow \text{minimal.}$$

Diese Summe F ist wiederum eine Funktion der Parameter $a, b, c, \dots$. Gesucht ist ein Minimum dieser Funktion. Aus der Theorie der Extrema in 10.4.3 kennen wir eine notwendige Bedingung für ein lokales Extremum: Alle partiellen Ableitungen müssen im stationären Punkt verschwinden

$$\Rightarrow \frac{\partial F}{\partial a} = 0, \quad \frac{\partial F}{\partial b} = 0, \quad \frac{\partial F}{\partial c} = 0, \dots$$

Dies liefert genau so viele (nichtlineare) Gleichungen wie Parameter im Ansatz enthalten sind.

Methode der kleinsten Quadrate: Zu n vorgegebenen Messpunkten $(x_i, y_i)_{i=1, \dots, n}$ lässt sich eine Ausgleichskurve wie folgt bestimmen:

- (1) Auswahl einer Ansatzfunktion (Gerade, Parabel, ...). Dieser Ansatz enthält zunächst unbestimmte Parameter $a, b, c, \dots$.
- (2) Man bildet die Summe der Abstandsquadrate

$$F(a, b, c, \dots) = \sum_{i=1}^n (y_i - f(x_i))^2.$$

Diese Summe ist eine Funktion der freien Parameter.

- (3) Die Parameter werden dann so bestimmt, dass die Funktion $F(a, b, c, \dots)$ minimal wird, d.h.

$$\frac{\partial F}{\partial a} = 0, \quad \frac{\partial F}{\partial b} = 0, \quad \frac{\partial F}{\partial c} = 0, \dots$$

⊙ Regressionsgerade

Wir werden nur den für die Anwendungen wichtigsten Spezialfall der Regressionsgeraden diskutieren. Wir nehmen an, dass zu n verschiedenen x -Werten $x_1, \dots, x_n$ die zugehörigen y -Werte $y_1, \dots, y_n$ vorliegen, so dass n Messpunkte

$$(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$$

gegeben sind. Gesucht sind die Parameter a und b in der Regressionsgeraden

$$f(x) = ax + b,$$

so dass die Summe der Abstandsquadrate minimal wird. Der Abstand d_i der Messpunkten zur Ausgleichsgeraden lautet

$$d_i = y_i - ax_i - b.$$

Die Summe der Abstandsquadrate

$$F(a, b) = \sum_{i=1}^n d_i^2 = \sum_{i=1}^n (y_i - ax_i - b)^2$$

ist daher eine Funktion der beiden Parameter a und b . Von dieser Funktion $F(a, b)$ suchen wir das Minimum. Um die stationären Punkte zu bestimmen, setzen wir die partiellen Ableitungen von F nach a und b gleich Null:

$$\begin{aligned} \frac{\partial F}{\partial a} &= 2 \sum_{i=1}^n (y_i - ax_i - b) (-x_i) \\ &= -2 \sum_{i=1}^n y_i x_i + 2a \sum_{i=1}^n x_i^2 + 2b \sum_{i=1}^n x_i \stackrel{!}{=} 0. \\ \frac{\partial F}{\partial b} &= 2 \sum_{i=1}^n (y_i - ax_i - b) (-1) \\ &= -2 \sum_{i=1}^n y_i + 2a \sum_{i=1}^n x_i + 2 \sum_{i=1}^n b \stackrel{!}{=} 0. \end{aligned}$$

Diese notwendigen Bedingungen bilden ein lineares Gleichungssystem für a, b :

$$\begin{aligned} \left(\sum_{i=1}^n x_i^2 \right) a + \left(\sum_{i=1}^n x_i \right) b &= \sum_{i=1}^n x_i y_i \\ \left(\sum_{i=1}^n x_i \right) a + n b &= \sum_{i=1}^n y_i. \end{aligned}$$

Da für die Koeffizientenmatrix gilt

$$D = \begin{vmatrix} \sum x_i^2 & \sum x_i \\ \sum x_i & n \end{vmatrix} = n \sum_{i=1}^n x_i^2 - \left(\sum_{i=1}^n x_i \right)^2 = \frac{1}{2} \sum_{j=1}^n \sum_{i=1}^n (x_i - x_j)^2 > 0,$$

erhalten wir z.B. mit der Cramerschen Regel die eindeutige Lösung des linearen Gleichungssystems

$$a = \hat{a} := \frac{1}{D} \left(n \cdot \sum_{i=1}^n x_i y_i - \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right) \right)$$

$$b = \hat{b} := \frac{1}{D} \left(\left(\sum_{i=1}^n x_i^2 \right) \left(\sum_{i=1}^n y_i \right) - \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n x_i y_i \right) \right).$$

Um zu überprüfen, dass F im Punkte $(\hat{a}, \hat{b})$ ein lokales Minimum annimmt, berechnen wir die zweiten partiellen Ableitungen

$$F_{aa} = 2 \sum_{i=1}^n x_i^2; \quad F_{bb} = 2n; \quad F_{ab} = 2 \sum_{i=1}^n x_i$$

und setzen sie in die Δ -Formel $\Delta = F_{aa}(\hat{a}, \hat{b}) \cdot F_{bb}(\hat{a}, \hat{b}) - F_{ab}^2(\hat{a}, \hat{b})$ ein:

$$\Delta = F_{aa} \cdot F_{bb} - F_{ab}^2 = 4n \sum_{i=1}^n x_i^2 - 4 \left(\sum_{i=1}^n x_i \right)^2 = 4D > 0.$$

Da $\Delta > 0$ und $F_{aa} > 0$ ist dies ein lokales Minimum.

Zusammenfassung: Zu n vorgegebenen Messpunkten $(x_i, y_i)_{i=1, \dots, n}$ bestimmt man durch die Methode der kleinsten Quadrate die Gerade

$$y = \hat{a}x + \hat{b}.$$

Diese Gerade heißt **Regressionsgerade** oder auch **Ausgleichsgerade**. Die Koeffizienten berechnen sich über die Formeln

$$\hat{a} := \frac{1}{D} \left(n \cdot \sum_{i=1}^n x_i y_i - \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right) \right)$$

$$\hat{b} := \frac{1}{D} \left(\left(\sum_{i=1}^n x_i^2 \right) \left(\sum_{i=1}^n y_i \right) - \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n x_i y_i \right) \right),$$

wenn

$$D = n \sum_{i=1}^n x_i^2 - \left(\sum_{i=1}^n x_i \right)^2.$$

Beispiel 10.33 (Mit MAPLE-Worksheet). Gesucht ist eine Ausgleichsgerade zu den folgenden Messwerten:

x	0	1	2	3	4	5
y	3	5	7	8	10	10

Dann ist

$$D = 6 \cdot (1 + 2^2 + 3^2 + 4^2 + 5^2) - (1 + 2 + 3 + 4 + 5)^2 = 105$$

$$\hat{a} = \frac{1}{105} [6 \cdot (1 \cdot 5 + 2 \cdot 7 + 3 \cdot 8 + 4 \cdot 10 + 5 \cdot 10) - (1 + 2 + 3 + 4 + 5)(3 + 5 + 7 + 8 + 10 + 10)] = \frac{51}{35}.$$

Analog berechnet sich $\hat{b} = \frac{74}{21}$. Die Regressionsgerade hat somit die Form

$$y = \frac{51}{35}x + \frac{74}{21}.$$

In Abb. 10.30 ist die Ausgleichsgerade zusammen mit den Messwerten graphisch dargestellt:

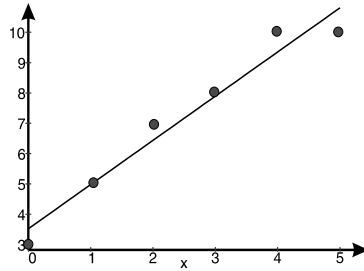


Abb. 10.30. Messwerte mit Regressionsgerade

Für eine größere Anzahl von Messpunkten muss die Bestimmung der Regressionsgeraden auf einem Rechner durchgeführt werden. $\square$

Wir haben den Fall der Ausgleichsgeraden ausführlich behandelt, da er durch Modifikation der Messwerte auch die logarithmische, die exponentielle sowie die Potenzanpassung beinhaltet:

Bemerkungen:

- (1) Ist eine **logarithmische Anpassung**

$$f(x) = a \ln x + b$$

an die Messwerte gesucht, führt man die Hilfsvariable $z = \ln x$ ein. Aus den x -Werten der Messung wird der Logarithmus gebildet und von den Mess-

werten (x_i, y_i) , $i = 1, \dots, n$ zu den Wertepaaren $(\ln(x_i), y_i)$, $i = 1, \dots, n$ übergegangen. Mit diesen Daten bestimmt man die Ausgleichsgerade.

- (2) Ist eine **exponentielle Anpassung**

$$f(x) = a e^{bx}$$

an die Messwerte gesucht, bildet man von der Gleichung

$$y = a e^{bx}$$

den Logarithmus

$$\ln y = \ln a + b \cdot x = Ax + B.$$

Aus den (x_i, y_i) -Messwerten geht man zu den Wertepaaren $(x_i, \ln(y_i))$ über und bestimmt die Parameter A und B der Ausgleichsgeraden. Dann ist $b = A$ und $a = e^B$.

- (3) Ist eine **Potenzanpassung**

$$f(x) = a x^b$$

an die Messwerte gesucht, bildet man von der Gleichung

$$y = a x^b$$

den Logarithmus

$$\ln y = \ln a + b \ln x.$$

Aus den (x_i, y_i) -Messwerten geht man zu den Wertepaaren $(\ln(x_i), \ln(y_i))$ über und bestimmt die Ausgleichsgerade

$$y = Ax + B.$$

Dann ist $b = A$ und $a = e^B$.

- (4) Die logarithmische und exponentielle Anpassung bzw. die Potenzanpassung entsprechen der Darstellung der Messwerte in einer logarithmischen bzw. doppellogarithmischen Auftragung (siehe Abschnitt 4.1.1).
- (5) In Verallgemeinerung der Vorgehensweise bei der Berechnung der Parameter der Ausgleichsgeraden bestimmt die MAPLE-Prozedur **ausgleich** die freien Parameter einer vorgegebenen Polynomfunktion nach der Methode der kleinsten Quadrate.

Beispiel 10.34 (Mit MAPLE-Worksheet). Zu den Messwerten aus Beispiel 10.33 ($[0,3]$, $[1,5]$, $[2,7]$, $[3,8]$, $[4,10]$, $[5,10]$) soll eine Ausgleichsparabel bestimmt werden. Mit der MAPLE-Prozedur **ausgleich** erhält man

Die Ausgleichsfunktion lautet, $-\frac{5}{28}x^2 + \frac{47}{20}x + \frac{41}{14}$

Das Abstandsquadrat ist, .4857142857

mit der graphischen Darstellung der Werte zusammen mit der Parabel. $\square$

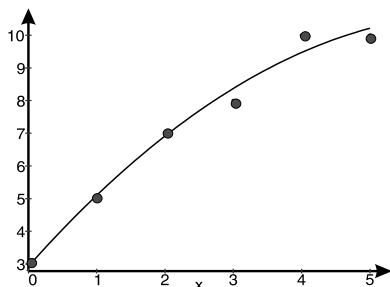


Abb. 10.31. Messwerte mit Ausgleichsparabel



Auf der CD-Rom sind mehrere MAPLE-Prozeduren zu diesem Kapitel enthalten: **differenzial** bestimmt das totale Differenzial einer Funktion, **fehler** berechnet den absoluten maximalen Fehler in linearer Näherung mit der Gaußschen Fehlerfortpflanzung, **stationaer** bestimmt die stationären Punkte einer Funktion in n Variablen. **extremum_2d** und **extremum_nd** berechnen die Extremwerte von Funktionen mit zwei bzw. n Variablen. Die Prozedur **Regressionsgerade** bestimmt die Ausgleichsgerade und stellt diese zusammen mit den Messwerten graphisch dar, während **ausgleich** das Ausgleichspolynom n -ter Ordnung angibt.

MAPLE-Worksheets zu Kapitel 10



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 10 mit MAPLE zur Verfügung.

- Darstellung von Funktionen mit zwei Variablen mit MAPLE
- Partielle Ableitung mit MAPLE
- Totale Differenzierbarkeit, Tangentialebene mit MAPLE
- Gradient und Richtungsableitung mit MAPLE
- Taylorsche Formel mit MAPLE
- Totales Differenzial mit MAPLE
- Fehlerrechnung mit MAPLE
- Bestimmung der Extrema mit MAPLE
- Ausgleichsrechnung mit MAPLE

10.5 Aufgaben zur Differenzialrechnung

- 10.1 Stellen Sie mit dem **plot3d**-Befehl die folgenden Funktionen in MAPLE graphisch dar

$$\begin{array}{lll} \text{a) } z = x \cdot y & \text{b) } z = x + y & \text{c) } z = x^2 - y^2 \\ \text{d) } z = x^2 + y^2 & \text{e) } z = (x - y)^2 & \text{f) } z = e^{-(x^2 + y^2)} \end{array}$$

- 10.2 Fügen Sie durch die Option **style = contour** 20 Höhenlinien in die Schaubilder ein und variieren Sie interaktiv den Blickwinkel.

- 10.3 Berechnen Sie für die folgenden Funktionen alle partiellen Ableitungen 1. Ordnung

$$\begin{array}{ll} \text{a) } f(x, y) = x^3 + x \cdot y - y^{-2} & \text{b) } f(a, t) = 3 \cdot a \cdot x + y \cdot \ln(t^2) \\ \text{c) } f(u, v) = \frac{u+v}{u+v} & \text{d) } f(x, y, z) = \operatorname{arcsinh}(x^2 + z^2) \\ \text{e) } f(x_1, x_2, x_3) = x_2 & \text{f) } f(a, b) = (a x + b x^2)^{-1} + y \cdot e^{a b} \end{array}$$

- 10.4 Man berechne die partiellen Ableitungen 1. und 2. Ordnung der folgenden Funktionen

$$\begin{array}{ll} \text{a) } f(x, y) = 3x^2 + 4xy - 2y^2 & \text{b) } f(x, y) = 2 \cos(3xy) \\ \text{c) } f(x, y) = (3x - 5y)^4 & \text{d) } f(x, y) = \frac{x^2 - y^2}{x + y} \\ \text{e) } f(x, y) = 3x \cdot e^{xy} & \text{f) } f(x, y) = \sqrt{x^2 - 2xy} \end{array}$$

- 10.5 Gegeben ist die Funktion $f(x, y) = \sin(x^2 + 2y)$.
Man bestätige den Satz von Schwarz, dass $f_{xy} = f_{yx}$.

- 10.6 Berechnen Sie die partiellen Ableitungen 2. Ordnung für die Funktion

$$f(x_1, x_2, x_3) = x_1 \cdot \ln(x_2^2 + x_3^2).$$

- 10.7 Gesucht sind alle partiellen Ableitungen zweiter Ordnung der Funktion

$$\text{a) } f(x, y) = (3x - 5y)^4 \quad \text{b) } f(x, y, z) = e^{x-y} \cdot \cos(5z)$$

- 10.8 Zeigen Sie, dass die Funktion

$$f(x, y, z) = \frac{a}{\sqrt{x^2 + y^2 + z^2}}$$

Lösung der Laplace-Gleichung $f_{xx} + f_{yy} + f_{zz} = 0$ ist.

- 10.9 Zeigen Sie durch Einsetzen, dass $z = x \cdot e^{y/x}$
der partiellen Differenzialgleichung $x \frac{\partial z}{\partial x} + y \frac{\partial z}{\partial y} = z$ genügt.

- 10.10 Zeigen Sie, dass die Funktion

$$f(x, y) = \frac{1}{2} \cdot \ln(x^2 + y^2)$$

die partielle Differenzialgleichung $f_{xx} + f_{yy} = 0$ erfüllt.

- 10.11 Bestimmen Sie im Punkte $P(1, 0)$ die Gleichung der Tangentialebene an die Fläche $z = (3x + x \cdot y)^2$.

- 10.12 Berechnen Sie an der Stelle $P(1, 2, 0)$ das totale Differenzial von

$$f(x, y, z) = y \cdot \cos(z) + \frac{\ln(1 + x^2)}{y}.$$

- 10.13 Berechnen Sie den Gradienten und die Richtungsableitung in Richtung $\vec{a} = \begin{pmatrix} 2 \\ 4 \end{pmatrix}$ für die Funktion

$$f(x, y) = (3x + x \cdot y)^2.$$

- 10.14 Berechnen Sie Gradient und die Richtungsableitung in Richtung $\vec{a} = \begin{pmatrix} 3 \\ -1 \\ 2 \end{pmatrix}$ für die Funktion

$$f(x, y, z) = y \cdot \cos(z) + \frac{\ln(1+x^2)}{y}.$$

- 10.15 Man bestimme mit MAPLE das totale Differenzial der Funktionen

a) $z(x, y) = 4x^3y - 3x \cdot e^y$

b) $z(x, y) = \frac{x^2 + y^2}{x - y}$

c) $f(x, y, z) = \ln \sqrt{x^2 + y^2 + z^2}.$

- 10.16 Betrachten Sie die differenzierbaren Funktionen $f_1, f_2 : \mathbb{R} \rightarrow \mathbb{R}$ und $g : \mathbb{R}^2 \rightarrow \mathbb{R}$ und bilden die Verkettung $h(x_1, x_2) = g(f_1(x_1), f_2(x_2))$. Berechnen Sie h und die ersten partiellen Ableitungen von h in folgenden Fällen

a) $f_1(x_1) = a_0 + a_1 x_1; \quad f_2(x_2) = b_0 + b_1 x_2; \quad g(u_1, u_2) = c_0 + c_1 u_1 + c_2 u_2.$

b) $f_1(x_1) = \sin x_1; \quad f_2(x_2) = \cos x_2; \quad g(u_1, u_2) = u_1^2 + u_1 u_2.$

- 10.17 Man berechne die Taylor-Reihe der Funktion f an der Stelle (x_0, y_0) bis zur Ordnung 2 für

a) $f(x, y) = \frac{(x-y)}{(x+y)}, \quad (x_0, y_0) = (1, 1)$

b) $f(x, y) = e^{x^2+y^2}, \quad (x_0, y_0) = (1, 0)$

- 10.18 Man berechne das totale Differenzial von

a) $f(x, y) = \sin(x^2 + 2y)$

b) $f(x, y) = 3x^2 + 4xy - 2y^2$

c) $f(x, y) = y \cdot \cos(x - 2y)$

d) $f(x, y, z) = x^2 z - y z^3 + x^4$

- 10.19 Für den Durchmesser eines geraden Kreiszylinders hat man $(6.0 \pm 0.003) m$ gemessen, für die Höhe $(4.0 \pm 0.02) m$. Wie groß ist der größte, absolute und relative Fehler des Zylindervolumens?

- 10.20 Zur Berechnung eines elektrischen Widerstandes $R = \frac{U}{I}$ werden die Stromstärke $I = (15 \pm 0.3) A$ und die Spannung $U = (110 \pm 2) V$ gemessen. Gesucht ist der relative Maximalfehler von R .

- 10.21 Der Elastizitätsmodul eines zylindrischen Drahtes (r : Radius des Drahtquerschnitts, l : Länge des Drahtes) wird bestimmt, indem die Längenzunahme z des Drahtes unter dem Einfluss der Kraft k gemessen wird. Es gilt

$$E = \frac{l \cdot k}{\pi r^2 \cdot z} \quad (\text{E-Modul}).$$

Wie groß und mit welcher Genauigkeit ist E bestimmt, wenn die Messwerte $l = (2000 \pm 3) mm$, $r = (0.2 \pm 0.002) mm$, $k = (200 \pm 0.05) N$ und $z = (15 \pm 0.1) mm$ betragen?

- 10.22 Zu bestimmen ist die Dichte ρ eines Messingstücks nach der Auftriebsmethode: Sei m das Gewicht in Luft, $\bar{m}$ das Gewicht in Wasser, dann gilt

$$\rho = \frac{m}{m - \bar{m}} = \frac{\text{Gewicht in Luft}}{\text{Volumen}}.$$

Wie groß ist der relative Fehler von ρ , wenn $m = (100 \pm 5 \cdot 10^{-3}) \text{ g}$ und $\bar{m} = (88 \pm 8 \cdot 10^{-3}) \text{ g}$?

- 10.23 Linearisieren Sie die Funktion

$$f(x, y, z) = y \cdot \cos(z) + \frac{\ln(1 + x^2)}{y}$$

an der Stelle $(x_0, y_0, z_0) = (1, 2, 0)$.

- 10.24 Bestimmen Sie für die folgenden Funktionen zunächst die kritischen Stellen und entscheiden Sie, ob (und wenn ja um welche) es sich um lokale Extremstellen handelt

a) $f(x, y) = x^2 + \cos(y)$

b) $f(x, y) = 3y^2 + 3xy - 18y^2$

c) $f(x, y) = (x - y)^3 + 12xy$

- 10.25 Welcher Punkt der Fläche $z = \sqrt{1 + (x - 2y)^2}$ hat den kleinsten Abstand vom Punkt $(1, -2, 0)$?

- 10.26 Zeigen Sie, dass die Funktion

$$f(x, y) = c - x^2 - y^2$$

im Punkte $(0, 0)$ ein lokales Maximum besitzt.

- 10.27 Bestimmen Sie die relativen Extrema der Funktion

$$f(x, y) = x^3 + y^3 - 3x - 12y + 20.$$

- 10.28 Bestimmen Sie mit MAPLE die relativen Extrema der Funktionen

a) $f(x, y) = 3xy - x^3 - y^3$

b) $f(x, y) = x^2 + y^2 + x - y$

c) $f(x, y) = 1 - x + y - 2xy + x^2 - y^2$

d) $f(x, y) = e^{x^2+y^2} - 2x^4 - 2y^2$

- 10.29 Bestimmen Sie alle stationären Punkte der Funktion

$$f(x, y, z) = e^{-(x^2+y^2+z^2)} \cdot (x^2 - z^2)$$

mit der Prozedur **stationaer** und überprüfen Sie mit **extremum_nd**, ob lokale Extrema vorliegen.

- 10.30 Bestimmen Sie mit MAPLE zu den folgenden Messreihen jeweils die Ausgleichsgerade

a)	$\begin{array}{c ccccccc} x_i & 0 & 1 & 2 & 3 & 4 & 5 & 6 \\ \hline y_i & 2.1 & 0.81 & -0.5 & -2.1 & -3.4 & -4.3 & -5.8 \end{array}$
b)	$\begin{array}{c ccccccc} x_i & 1.5 & 1.8 & 2.4 & 3.0 & 3.5 & 4.0 & 4.5 & 6.0 \\ \hline y_i & 1.9 & 2.1 & 2.8 & 3.4 & 4.0 & 4.1 & 5.1 & 6.1 \end{array}$

Tragen Sie die Punkte zusammen mit der Ausgleichsgeraden in ein Schaubild ein!

Kapitel 11

Integralrechnung bei Funktionen mit mehreren Variablen

11

11	Integralrechnung bei Funktionen mit mehreren Variablen	449
11.1	Doppelintegrale (Gebietsintegrale)	451
11.1.1	Definition	451
11.1.2	Berechnung von Doppelintegralen	454
11.1.3	Reduktion von Doppelintegralen auf einfache Integrationen	457
11.1.4	Anwendungen	459
11.2	Dreifachintegrale	464
11.2.1	Definition und Berechnung von Dreifachintegralen	464
11.2.2	Anwendungen	466
11.3	Aufgaben zur Integralrechnung	471
Zusätzliche Abschnitte auf der CD-Rom		
11.4	MAPLE: Doppel- und Dreifachintegrale	cd
11.4.1	Doppelintegrale mit MAPLE	cd
11.4.2	Dreifachintegrale mit MAPLE	cd
11.4.3	Berechnung der Eigenschaften starrer Körper mit MAPLE	cd
11.5	Substitutionsregeln, Koordinatentransformationen	cd
11.6	Linien- oder Kurvenintegrale	cd
11.6.1	Vektordarstellung einer Kurve	cd
11.6.2	Differentiation eines Vektors nach einem Parameter	cd
11.6.3	Vektor- oder Kraftfelder	cd
11.6.4	Linien- oder Kurvenintegrale	cd
11.6.5	Anwendungsbeispiele	cd
11.7	Oberflächenintegrale	cd
11.7.1	Oberflächenintegral eines Vektorfeldes	cd

11 Integralrechnung bei Funktionen mit mehreren Variablen

In diesem Kapitel wird der Begriff des bestimmten Integrals auf Doppel-, Dreifach- und Kurvenintegrale sowie auf Oberflächenintegrale erweitert. Bei jedem dieser Begriffe wird die Berechnung des Integralwertes auf die eines bestimmten Integrals zurück gespielt. Zunächst führen wir in 11.1 Doppelintegrale z.B. zur Beschreibung von Volumina, Schwerpunkten von ebenen Flächen und Flächenmomenten ein. Anschließend übertragen wir in 11.2 die Vorgehensweise auf Dreifachintegrale, um Schwerpunkte und Massenträgheitsmomente von Körpern zu berechnen.

Hinweis: Eine weitere Notwendigkeit, den Integralbegriff auf Funktionen mit mehreren Variablen zu erweitern, besteht in der Integration entlang einer Linie. Dies führt auf den Begriff der [Linienintegrale](#) in 11.4, die in der Elektrodynamik und Thermodynamik zur Berechnung der Energie herangezogen werden. In 11.5 werden [Oberflächenintegrale](#) diskutiert. Substitutionsregeln und Koordinatentransformationen werden in 11.3 behandelt. Diese beiden Abschnitte befinden sich zusätzlich auf der CD-Rom.

11.1 Doppelintegrale (Gebietsintegrale)

11.1

Wir beschäftigen uns in diesem Abschnitt mit zweidimensionalen Integralen. Die Übertragung vom eindimensionalen auf den mehrdimensionalen Fall bereitet im Prinzip keine Schwierigkeiten, wird von seiner Konstruktion her aber aufwändiger, da nicht ein Intervall sondern nun ein zweidimensionales Gebiet aufgeteilt werden muss.

Hinweis: Die Anwendungsbeispiele werden durch [MAPLE-Prozeduren](#) ergänzt, die sich zusätzlich auf der CD-Rom befinden.

➤ 11.1.1 Definition

Das bestimmte Integral einer positiven Funktion $f(x)$ im Intervall $[a, b]$ repräsentiert die Fläche, welche die Kurve $f(x)$ mit der x -Achse im Bereich $[a, b]$ einschließt. Definiert wird das bestimmte Integral

$$\int_a^b f(x) dx$$

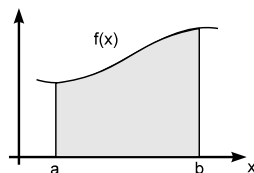


Abb. 11.1.

als Grenzwert über die Summe aller Rechteckflächen $(x_k - x_{k-1}) f(\xi_k)$, wenn die Intervallbreite $\Delta x_k = (x_k - x_{k-1})$ der Unterteilung des Intervalls

$$a = x_0 < x_1 < \dots < x_{n-1} < x_n = b$$

für $n \rightarrow \infty$ gegen Null strebt (siehe Abb. [11.2](#)).

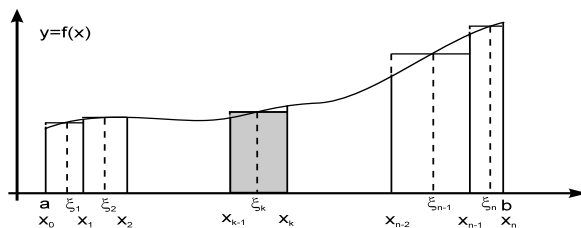
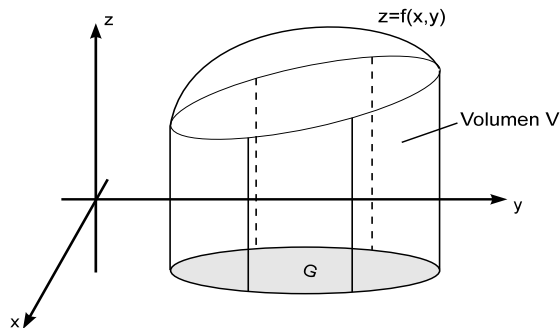


Abb. 11.2. Summe der Rechteckflächen

Um diesen Integralbegriff auf zweidimensionale Gebiete zu übertragen, reicht es für unsere Zwecke vollkommen aus, nur einen beschränkten, einfach zusammenhängenden Bereich $G \subset \mathbb{R}^2$ in der (x, y) -Ebene zu betrachten, der einen "glatten" Rand besitzt. Diesen Bereich nennen wir im Folgenden ein *Gebiet*. Es sei $z = f(x, y)$ eine auf dem Gebiet G stetige, positive Funktion. Gesucht ist das Volumen zwischen dem Funktionsgraphen und G :

Abb. 11.3. Funktion $f(x, y)$ über einem zweidimensionalen Gebiet G

Zur Bestimmung des Volumens V zerlegen wir das Gebiet G in Rechteckflächen der Länge Δx_i und Breite Δy_j . Die Anzahl der Unterteilungen in x -Richtung sei n , die Anzahl der Unterteilungen in y -Richtung m . D.h. i variiert zwischen 1 und n , j variiert zwischen 1 und m . Für jedes Rechteck mit Index (i, j) wählen wir einen beliebigen Punkt $P(\xi_i, \eta_j)$ und bestimmen den zugehörigen Funktionswert $f(\xi_i, \eta_j)$. Das Zylindervolumen über dem Rechteck beträgt Grundfläche mal Höhe:

$$V_{ij} = f(\xi_i, \eta_j) \Delta x_i \Delta y_j .$$

Anschließend bildet man die Summe aller Zylindervolumen über dem Gebiet G . (Für Rechtecke außerhalb von G setze man das Volumen auf Null.) Die Zwischensumme aller Zylindervolumen bildet eine Näherung für das zu be-

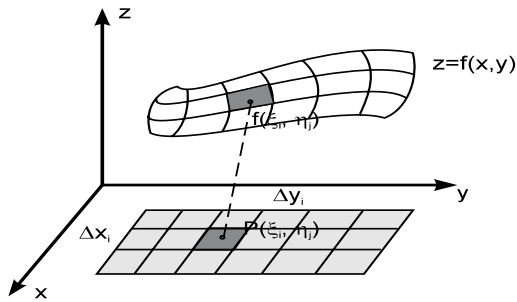


Abb. 11.4. Zur Konstruktion des Doppelintegrals

rechnende Volumen

$$V \approx \sum_{j=1}^m \sum_{i=1}^n f(\xi_i, \eta_j) \Delta x_i \Delta y_j .$$

Je feiner die Unterteilung des Gebietes G ausfällt, umso genauer ist diese Näherung. Wir lassen daher die Anzahl der Unterteilungen in x - und y -Richtung anwachsen. Strebt beim Grenzübergang $n \rightarrow \infty$, $m \rightarrow \infty$ die Zwischensumme gegen einen Grenzwert, so bezeichnet man ihn als *Doppelintegral* bzw. als *Gebietsintegral*.

Definition: (Doppelintegral, Gebietsintegral). Der Grenzwert

$$\iint_{(G)} f(x, y) dG := \lim_{m \rightarrow \infty} \lim_{n \rightarrow \infty} \sum_{j=1}^m \sum_{i=1}^n f(\xi_i, \eta_j) \Delta x_i \Delta y_j$$

wird (falls er existiert) als *Doppelintegral* bzw. *Gebietsintegral* von f über G bezeichnet. Man nennt dann f über G integrierbar.

Bemerkungen:

- (1) Oftmals wird die symbolische Schreibweise $\int_{(G)} f(x, y) dG$ nur mit einem Integralzeichen benutzt.
- (2) Man nennt

(x, y)	die Integrationsvariablen
$f(x, y)$	den Integrand
dG	das Flächenelement
(G)	den Integrationsbereich.

- (3) Für stetige Funktionen ist das Doppelintegral immer definiert.

- (4) Existiert der Grenzwert, dann ist er unabhängig von der speziellen Gebietszerlegung.
- (5) Für die algebraische Definition des Doppelintegrals ist $f(x, y) \geq 0$ nicht erforderlich, sondern nur für die geometrische Interpretation als Volumen zwischen dem Graphen von f und G .

► 11.1.2 Berechnung von Doppelintegralen

Um Doppelintegrale zu berechnen, zerlegen wir $\iint_{(G)} f(x, y) dG$ in zwei nacheinander auszuführende einfache Integrale. Dazu überdecken wir das Gebiet G mit einem achsenparallelen Rechteckgitter mit Maschenweiten Δx und Δy und bilden die Summe über alle Zylinder, deren Grundfläche in G liegt:

$$V \approx \sum_{i=1}^n \left(\sum_{j=1}^m f(\xi_i, \eta_j) \Delta y \right) \Delta x .$$

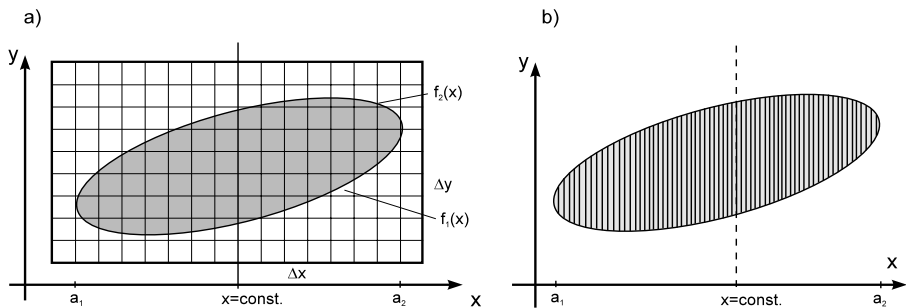


Abb. 11.5. Berechnung eines Doppelintegrals über eine Zerlegung der Grundfläche

Stellen wir den unteren und oberen Rand durch Funktionen $f_1(x)$ und $f_2(x)$ dar, dann konvergiert die innere Summe für $\Delta y \rightarrow 0$ ($m \rightarrow \infty$) für festes $\xi_i \in [x_{i-1}, x_i]$ gegen das Integral

$$I_y(\xi_i) = \int_{f_1(\xi_i)}^{f_2(\xi_i)} f(\xi_i, y) dy .$$

Dies ist ein gewöhnliches Integral mit der Integrationsvariablen y . Für festes ξ_i wird entlang der y -Achse von $f_1(\xi_i)$ bis $f_2(\xi_i)$ integriert.

$$\begin{aligned} \Rightarrow V &\approx \lim_{m \rightarrow \infty} \sum_{i=1}^n \left(\sum_{j=1}^m f(\xi_i, \eta_j) \Delta y \right) \Delta x \\ &= \sum_{i=1}^n I_y(\xi_i) \Delta x . \end{aligned}$$

Lassen wir nun auch noch $\Delta x \rightarrow \infty$ gehen ($n \rightarrow \infty$), dann konvergiert die verbleibende Summe gegen das Integral

$$\lim_{n \rightarrow \infty} \sum_{i=1}^n I_y(\xi_i) \Delta x = \int_{a_1}^{a_2} I_y(x) dx.$$

Dieses Ergebnis ist der doppelte Grenzwert $n \rightarrow \infty$, $m \rightarrow \infty$, also genau das Doppelintegral von f über dem Gebiet G .

$$\Rightarrow \iint_{(G)} f(x, y) dx dy = \int_{x=a_1}^{a_2} \left(\int_{y=f_1(x)}^{f_2(x)} f(x, y) dy \right) dx.$$

Berechnung von Doppelintegralen:

Die Berechnung des Doppelintegrals $\iint_{(G)} f(x, y) dG$ erfolgt durch zwei nacheinander auszuführenden, gewöhnlichen Integrationen:

$$\iint_{(G)} f(x, y) dG = \int_{x=a_1}^{a_2} \underbrace{\left(\int_{y=f_1(x)}^{f_2(x)} f(x, y) dy \right)}_{x=\text{const}, y \text{ variiert zwischen } f_1(x) \text{ und } f_2(x)} dx. \quad (D1)$$

- (1) Das innere Integral wird für festes x nach y integriert. Die Grenzen sind $y_1 = f_1(x)$ und $y_2 = f_2(x)$. Das Ergebnis der ersten Integration ist eine Funktion, in der nur noch x als Variable vorkommt.
- (2) Das äußere Integral wird durch Integration über x bestimmt.

Bemerkungen:

- (1) Die Integrationsformel (D1) entspricht der Zerlegung des Gebietes in Streifen parallel zur y -Achse und anschließender Aufsummierung aller Streifen in x -Richtung (siehe Abb. 11.5 b). Die einzelnen, zur y -Achse parallelen Streifen besitzen Anfangs- und Endwerte, y_1 und y_2 , die wiederum von der Lage des Streifens, d.h. von der x -Koordinate abhängen: $y_1 = f_1(x)$ und $y_2 = f_2(x)$. Das Aufsummieren erfolgt über alle Streifen zwischen $x_{\min} = a_1$ und $x_{\max} = a_2$.
- (2) Es wird bei dieser Darstellung davon ausgegangen, dass der Rand sich durch zwei Funktionen $f_1(x)$ und $f_2(x)$ beschreiben lässt. Für komplizierte Gebiete muss G in entsprechende Teilbereiche unterteilt werden (siehe Abb. 11.6, links).

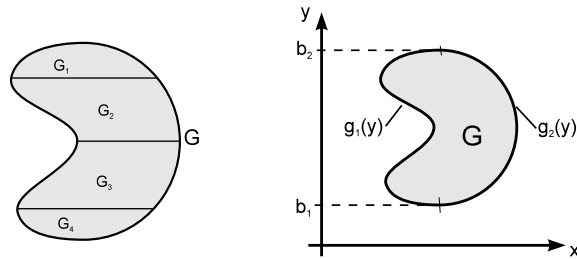


Abb. 11.6. Zerlegung in Teilgebiete

- (3) Vertauscht man die Rolle von x und y und stellt den linken Rand durch die Funktion $g_1(y)$ und den rechten Rand durch $g_2(y)$ dar, erhält man für das Doppelintegral

$$\iint_{(G)} f(x, y) dG = \int_{y=b_1}^{b_2} \underbrace{\left(\int_{x=g_1(y)}^{g_2(y)} f(x, y) dx \right)}_{y=\text{const, integriert wird über } x} dy. \quad (\text{D2})$$

Das innere Integral wird für festes y von $g_1(y)$ bis $g_2(y)$ nach x integriert. Das Ergebnis enthält nur noch die Variable y . Das äußere Integral wird durch Integration über y bestimmt (siehe Abb. 11.6, rechts).

- (4) Graphisch entspricht die Integrationsformel (D2) einer Zerlegung des Gebietes in Streifen parallel zur x -Achse und anschließender Aufsummierung aller Streifen in y -Richtung. Anfangs- und Endpunkte der Streifen hängen hier von y ab: $g_1(y)$ und $g_2(y)$. Durch Aufsummieren werden alle Streifen von $y_{\min} = b_1$ bis $y_{\max} = b_2$ berücksichtigt.
- (5) Die Reihenfolge der Integration ist durch die Reihenfolge der Differenziale dx und dy von Innen nach Außen festgelegt.
- (6) Das eigentliche Problem bei der Berechnung von Doppelintegralen besteht im Auffinden der Funktionen $f_1(x)$ und $f_2(x)$ bzw. $g_1(y)$ und $g_2(y)$. Sind diese, den Rand beschreibenden Funktionen gefunden, reduziert sich die Berechnung auf die Bestimmung gewöhnlicher Integrale.

Beispiele 11.1 (Die Integrationsgrenzen sind vorgegeben):

- ① Die Bestimmung des Doppelintegrals

$$I = \int_{y=0}^1 \int_{x=-2}^y x y dx dy$$

erfolgt, indem zunächst das innere Integral (y fest, x variiert von $x = -2$ bis y)

$$\int_{x=-2}^y x y \, dx = y \int_{x=-2}^y x \, dx = y \cdot \left[\frac{x^2}{2} \right]_{-2}^y = y \left[\frac{y^2}{2} - 2 \right] = \frac{1}{2} y^3 - 2y$$

nach der Variablen x integriert wird. Anschließend wird die äußere Integration über y durchgeführt

$$I = \int_{y=0}^1 \left(\frac{1}{2} y^3 - 2y \right) dy = \left[\frac{1}{8} y^4 - y^2 \right]_0^1 = -\frac{7}{8}.$$

② Zur Bestimmung des Doppelintegrals

$$I = \int_{x=0}^3 \int_{y=0}^{\pi} x^2 \sin(y) \, dy \, dx$$

berechnen wir zunächst das innere Integral (x fest, y variiert von 0 bis π)

$$\int_{y=0}^{\pi} x^2 \sin(y) \, dy = x^2 \int_{y=0}^{\pi} \sin(y) \, dy = x^2 [-\cos(y)]_0^{\pi} = 2x^2,$$

dann das äußere

$$I = \int_{x=0}^3 2x^2 \, dx = \frac{2}{3} x^3 \Big|_0^3 = 18. \quad \square$$

► 11.1.3 Reduktion von Doppelintegralen auf einfache Integrationen

Beispiel 11.2 (Mit Gebietszerlegung).

Gegeben ist eine Funktion $z = f(x, y)$, die auf dem Gebiet G (siehe Abb. 11.7) definiert ist. Gesucht ist $\iint_{(G)} f(x, y) \, dG$.

Um das Doppelintegral auf zwei einfache Integrationen zu reduzieren, zerlegen wir das Gebiet G in Streifen parallel zur x -Achse. Die Streifen beginnen bei $x = y$ und enden bei $x = y + 4$. Anschließend müssen alle Streifen von $y = -1$ bis $y = 1$ berücksichtigt werden. Dies bedeutet, dass wir Integralformel (D2) wählen und als innere Integrationsvariable x setzen. D.h. wir halten $y = \text{const}$, dann variiert x zwischen den Werten $g_1(y) = y$ und $g_2(y) = y + 4$ (siehe gestrichelte Linie). Anschließend variiert im äußeren Integral y von -1 bis 1.

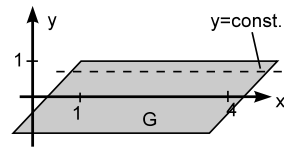


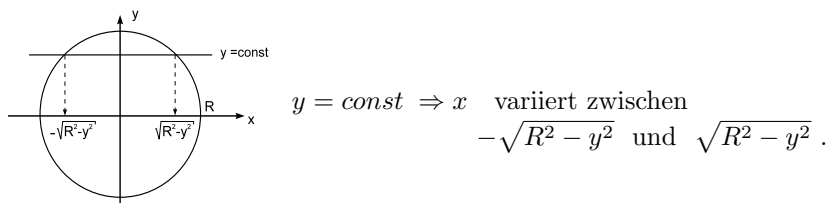
Abb. 11.7. Gebiet G

$$\Rightarrow \iint_{(G)} f(x, y) \, dG = \int_{y=-1}^1 \underbrace{\left(\int_{x=y}^{y+4} f(x, y) \, dx \right)}_{y=\text{const}} dy. \quad \square$$

Beispiel 11.3 (Mit Gebietszerlegung).

Gesucht ist das Doppelintegral $\iint_{(G)} f(x, y) dG$, wenn das Gebiet G einen Kreis in der (x, y) -Ebene mit Radius R darstellt.

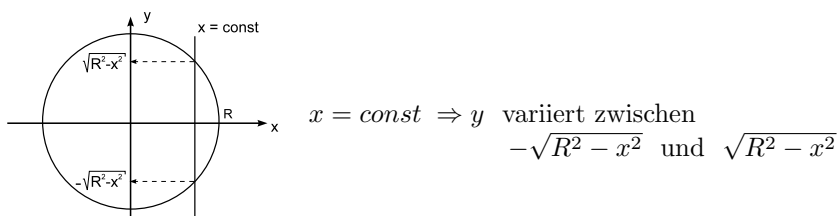
- (i) Zerlegung des Gebietes in Streifen parallel zur x -Achse: Führen wir die Integration mit Formel (D2) aus, erfolgt die innere Integration für konstantes y über die Variable x .



Für festes y variieren also die zugehörigen x -Werte zwischen $g_1(y) = -\sqrt{R^2 - y^2}$ und $g_2(y) = \sqrt{R^2 - y^2}$. Zur Bestimmung des äußeren Integrals muss y dann zwischen $b_1 = -R$ und $b_2 = R$ variieren.

$$\Rightarrow \iint_{(G)} f(x, y) dG = \int_{y=-R}^R \left(\int_{x=-\sqrt{R^2-y^2}}^{\sqrt{R^2-y^2}} f(x, y) dx \right) dy .$$

- (ii) Zerlegung des Gebietes in Streifen parallel zur y -Achse: Führen wir die Integration gemäß Formel (D1) aus, erfolgt die innere Integration für konstantes x über die Variable y .



Für festes x variieren nun die y -Werte zwischen $f_1(x) = -\sqrt{R^2 - x^2}$ und $f_2(x) = \sqrt{R^2 - x^2}$. Zur Bestimmung des äußeren Integrals muss anschließend x zwischen $-R$ und R variieren

$$\Rightarrow \iint_{(G)} f(x, y) dG = \int_{x=-R}^R \left(\int_{y=-\sqrt{R^2-x^2}}^{\sqrt{R^2-x^2}} f(x, y) dy \right) dx . \quad \square$$

11.1.4 Anwendungen

Doppelintegrale kommen häufig bei der Berechnung von ebenen Flächen, bei der Schwerpunktsberechnung ebener Flächen, bei der Bestimmung von Flächenmomenten und Volumenberechnungen vor. Auf jede dieser Problemstellungen werden wir im Folgenden eingehen.

Flächenberechnungen

Setzt man die Funktion $z = f(x, y) = 1$, entspricht das Doppelintegral über G dem Volumen des Körpers mit der Grundfläche G und der konstanten Höhe 1. Dies ist zahlenmäßig gerade der Flächeninhalt von G :

Der **Flächeninhalt** A eines ebenen Gebietes $G \subset \mathbb{R}^2$ ist

$$A = \iint_{(G)} dG .$$

Beispiel 11.4. Gesucht ist die Fläche, die durch die Gerade $f_1(x) = x + 2$ und die Parabel $f_2(x) = 4 - x^2$ begrenzt wird.

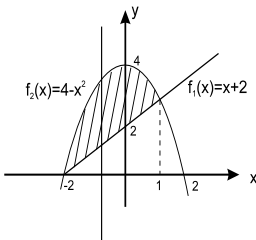


Abb. 11.8. Fläche zwischen zwei Funktionen

Zur Berechnung des Doppelintegrals zerlegen wir das Gebiet in Streifen parallel zur y -Achse; wir wählen also Formel (D1): Für festes x variiert y von $f_1(x) = x + 2$ bis $f_2(x) = 4 - x^2$; das verbleibende äußere Integral über dx wird gebildet mit den Grenzen $a_1 = -2$ und $a_2 = +1$:

$$\begin{aligned} A &= \iint_{(G)} dG = \int_{x=-2}^1 \underbrace{\left(\int_{y=x+2}^{4-x^2} dy \right)}_{x \text{ fest}} dx = \int_{x=-2}^1 [y]_{x+2}^{4-x^2} dx \\ &= \int_{x=-2}^1 (-x^2 - x + 2) dx = \left[-\frac{1}{3}x^3 - \frac{1}{2}x^2 + 2x \right]_{-2}^1 = 4.5. \quad \square \end{aligned}$$

Beispiel 11.5. Gesucht ist der Flächeninhalt des Gebietes G , welches in Abb. 11.9 definiert ist.

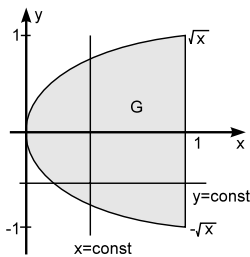


Abb. 11.9.

(i) Zerlegung des Gebietes in Streifen parallel zur y -Achse: Für $x = \text{const}$ variiert y von $-\sqrt{x}$ bis $\sqrt{x}$. Mit Formel (D1) gilt damit

$$\begin{aligned} A &= \iint_{(G)} dG = \int_{x=0}^1 \underbrace{\left(\int_{y=-\sqrt{x}}^{\sqrt{x}} dy \right)}_{x=\text{const}} dx \\ &= \int_0^1 2\sqrt{x} dx = \int_0^1 2x^{\frac{1}{2}} dx = \frac{4}{3} \left[x^{\frac{3}{2}} \right]_0^1 = \frac{4}{3}. \end{aligned}$$

(ii) Zerlegung des Gebietes in Streifen parallel zur x -Achse: Mit Formel (D2) folgt ebenfalls ($y = \text{const} \Rightarrow x$ variiert von y^2 bis 1)

$$\begin{aligned} A &= \iint_{(G)} dG = \int_{y=-1}^1 \underbrace{\left(\int_{x=y^2}^1 dx \right)}_{x=y^2} dy \\ &= \int_{y=-1}^1 (1 - y^2) dy = \left[y - \frac{1}{3} y^3 \right]_{-1}^1 = \frac{4}{3}. \end{aligned}$$

□

⊗ Schwerpunktsberechnung ebener Flächen

In §8.6.6 sind Formeln die Schwerpunktskoordinaten einer Fläche unter einem Graphen f und der x -Achse hergeleitet. Für ebene Gebiete gilt allgemein:

Schwerpunkt einer homogenen ebenen Fläche: Die Koordinaten des Schwerpunktes $S = (x_s, y_s)$ eines ebenen Gebietes G bestimmen sich über

$x_s = \frac{1}{A} \iint_{(G)} x dG,$	$y_s = \frac{1}{A} \iint_{(G)} y dG,$
---------------------------------------	---------------------------------------

wenn A der Flächeninhalt von G ist.

Beispiel 11.6. Gesucht ist der Schwerpunkt für die Fläche aus Beispiel 11.5: Für die x -Koordinate von S gilt

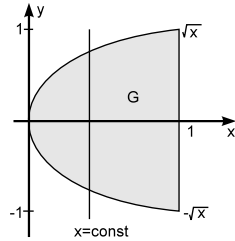
$$x_s = \frac{1}{A} \iint_{(G)} x \, dG = \frac{1}{A} \int_{x=0}^1 \underbrace{\left(\int_{y=-\sqrt{x}}^{\sqrt{x}} x \, dy \right)}_{x=\text{const}} dx$$

Das innere Integral über dy ist für festes x

$$\int_{y=-\sqrt{x}}^{\sqrt{x}} x \, dy = x \int_{y=-\sqrt{x}}^{\sqrt{x}} dy = 2x x^{\frac{1}{2}} = 2x^{\frac{3}{2}}$$

und das äußere Integral

$$\int_{x=0}^1 2x^{\frac{3}{2}} dx = \frac{4}{5} x^{\frac{5}{2}} \Big|_0^1 = \frac{4}{5} \curvearrowright x_s = \frac{1}{A} \cdot \frac{4}{5} = \frac{3}{5}.$$



Aufgrund der Symmetrie ist $y_s = 0$. $\Rightarrow S = (\frac{3}{5}, 0)$. $\square$

Beispiel 11.7 (Mit MAPLE-Worksheet). Gesucht sind die Schwerpunktskoordinaten für das Gebiet aus Beispiel 11.4: Mit Beispiel 11.4 gilt für die x -Koordinate des Schwerpunktes

$$\begin{aligned} x_s &= \frac{1}{A} \iint_{(G)} x \, dG \\ &= \frac{1}{4.5} \int_{x=-2}^1 \left(\int_{y=x+2}^{4-x^2} x \, dy \right) dx \end{aligned}$$

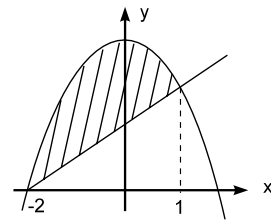


Abb. 11.10.

$$\begin{aligned} \text{bzw. } x_s &= \frac{2}{9} \int_{x=-2}^1 x(4 - x^2 - x - 2) dx = \frac{2}{9} \int_{x=-2}^1 (-x^3 - x^2 + 2x) dx \\ &= \frac{2}{9} \left[-\frac{1}{4}x^4 - \frac{1}{3}x^3 + x \right]_{-2}^1 = -\frac{1}{2}. \end{aligned}$$

Entsprechend gilt für die y -Komponente $y_s = 2.4$. $\square$

Beispiel 11.8 (Mit MAPLE-Worksheet). Gesucht sind die Schwerpunktskoordinaten des Viertelkreises.

Wir wählen zur Berechnung des Doppelintegrals Formel (D2): Dabei erfolgt die innere Integration bei konstantem y über die Variable x von $g_1(y) = 0$ bis $g_2(y) = \sqrt{R^2 - y^2}$. Zur Bestimmung des äußeren In-

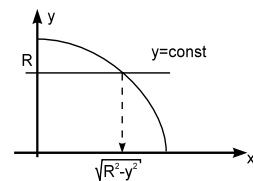


Abb. 11.11.

tegrals über y variiert dann die Variable y von $0 \leq y \leq R$:

$$x_s = \frac{1}{A} \iint_{(G)} x \, dG = \frac{1}{A} \int_{y=0}^R \left(\int_{x=0}^{\sqrt{R^2-y^2}} x \, dx \right) dy;$$

$$y_s = \frac{1}{A} \iint_{(G)} y \, dG = \frac{1}{A} \int_{y=0}^R \left(\int_{x=0}^{\sqrt{R^2-y^2}} y \, dx \right) dy.$$

Mit MAPLE folgt für die Koordinaten mit $A = \frac{\pi R^2}{4}$:

$$x_s = \frac{4}{3} \frac{R}{\pi} \text{ und } y_s = \frac{4}{3} \frac{R}{\pi}$$

□

⊗ Flächenmomente

In der Festigkeitslehre benötigt man zur Beschreibung von Biegungen **Flächenmomente** von Querschnittsflächen. Sie sind jeweils bezogen auf bestimmte Achsen. Es wird unterschieden zwischen sog. *axialen* Momenten, bei denen die Bezugsachse in der Flächenebene liegt und *polaren* Momenten, bei denen die Achse senkrecht zur Flächenebene orientiert ist. Es gelten die Formeln

Axiales Flächenmoment bzg. der y -Achse	$I_y = \iint_{(G)} x^2 \, dG$
Axiales Flächenmoment bzg. der x -Achse	$I_x = \iint_{(G)} y^2 \, dG$
Polares Flächenmoment	$I_p = \iint_{(G)} (x^2 + y^2) \, dG$

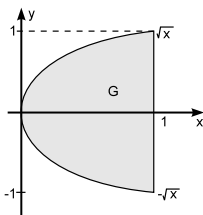


Abb. 11.12.

Beispiel 11.9. Gesucht sind die axialen Flächenmomente des Gebietes G aus Beispiel 11.5. Gemäß der Zerlegung aus Beispiel 11.5 berechnen wir

$$\begin{aligned} I_x &= \iint_{(G)} y^2 \, dG = \int_0^1 \left(\int_{-\sqrt{x}}^{\sqrt{x}} y^2 \, dy \right) dx \\ &= \int_0^1 \left[\frac{1}{3} y^3 \right]_{-\sqrt{x}}^{\sqrt{x}} dx = \frac{2}{3} \int_0^1 x^{\frac{3}{2}} dx = \frac{2}{3} \frac{2}{5} \left[x^{\frac{5}{2}} \right]_0^1 = \frac{4}{15}. \end{aligned}$$

$$\begin{aligned}
 I_y &= \iint_{(G)} x^2 dG = \int_0^1 \left(\int_{-\sqrt{x}}^{\sqrt{x}} x^2 dy \right) dx = \int_0^1 x^2 \left(\int_{-\sqrt{x}}^{\sqrt{x}} dy \right) dx \\
 &= 2 \int_0^1 x^2 x^{\frac{1}{2}} dx = 2 \int_0^1 x^{\frac{5}{2}} dx = 2 \frac{2}{7} \left[x^{\frac{7}{2}} \right]_0^1 = \frac{4}{7}. \quad \square
 \end{aligned}$$

➤ Volumenberechnung

Aufgrund seiner Definition dient das Doppelintegral zur Berechnung von *Volumeninhalten*, die ein Funktionsgraph einer positiven Funktion $z = f(x, y) > 0$ mit der (x, y) -Ebene über einem Gebiet G einschließt

$$V = \iint_{(G)} f(x, y) dG.$$

Beispiel 11.10 (Mit MAPLE-Worksheet). Gesucht ist das Volumen V , das durch den Graphen von

$$z = f(x, y) = 1 - x^2 - y^2$$

und der (x, y) -Ebene eingeschlossen wird. Das zugehörige Gebiet G soll der Einheitskreis in der (x, y) -Ebene darstellen.

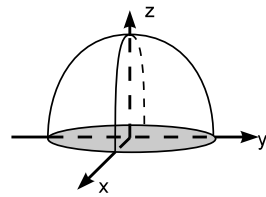


Abb. 11.13.

Für konstantes x wird y begrenzt durch $f_1(x) = -\sqrt{1-x^2}$ und $f_2(x) = \sqrt{1-x^2}$ mit $a_1 = -1$ und $a_2 = 1$. Nach der Integrationsformel (D1) für Doppelintegrale ist

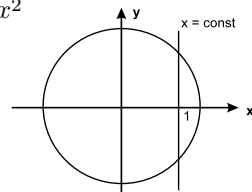


Abb. 11.14.

$$\begin{aligned}
 V &= \iint_{(G)} f(x, y) dG \\
 &= \int_{x=-1}^1 \left(\int_{y=-\sqrt{1-x^2}}^{\sqrt{1-x^2}} (1 - x^2 - y^2) dy \right) dx \\
 &= \int_{x=-1}^1 \left[(1 - x^2) y - \frac{1}{3} y^3 \right]_{y=-\sqrt{1-x^2}}^{\sqrt{1-x^2}} dx = \int_{x=-1}^1 \frac{4}{3} (1 - x^2)^{3/2} dx \\
 &= \left[\frac{1}{3} (1 - x^2)^{3/2} + \frac{1}{2} (1 - x^2)^{1/2} + \frac{1}{2} \arcsin(x) \right]_{x=-1}^1 = \frac{1}{2} \pi. \quad \square
 \end{aligned}$$

11.2 Dreifachintegrale

Der Begriff des *Dreifachintegrals* für Funktionen $f(x, y, z)$ über einem dreidimensionalen Gebiet $G \subset \mathbb{R}^3$ wird auf ähnliche Weise eingeführt und berechnet wie das Doppelintegral. Da sich eine Funktion von drei Variablen nicht mehr graphisch darstellen lässt, besitzt das Dreifachintegral über eine Funktion $f(x, y, z)$ zunächst keine geometrische Bedeutung. Nur für den Fall $f(x, y, z) = 1$ entspricht das Dreifachintegral dem Volumen des Gebietes G . Auch Dreifachintegrale reduziert man auf jetzt 3 aufeinander folgende, gewöhnliche Integrationen, wobei nun $3! = 6$ verschiedene Integrationsreihenfolgen möglich sind!

➤ 11.2.1 Definition und Berechnung von Dreifachintegralen

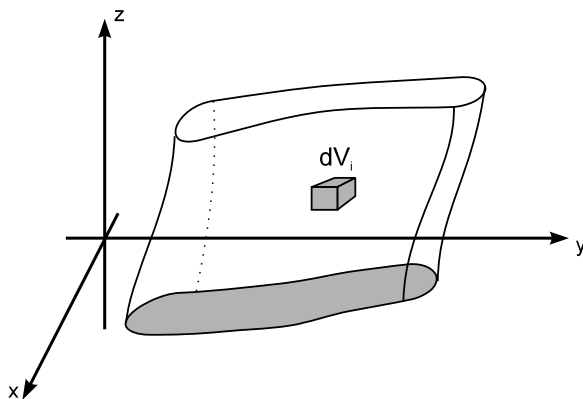


Abb. 11.15. Zur Definition von Dreifachintegralen

Zur Definition des Dreifachintegrals zerlegen wir den Körper in kleine Teilvolumen dV_i ($i = 1, \dots, n$) und wählen in jedem Volumen einen Punkt $P(x_i, y_i, z_i)$ aus, auf dem wir die Funktion f auswerten. Schließlich bilden wir das Produkt von Funktionswert und Volumenelement

$$f(x_i, y_i, z_i) dV_i$$

und summieren über alle Volumina auf

$$Z_n = \sum_{i=1}^n f(x_i, y_i, z_i) dV_i. \quad (\text{Zwischensumme})$$

Wir lassen nun die Anzahl der Teilvolumina anwachsen (dies bedeutet gleichzeitig, dass $dV_i \rightarrow 0$ geht). Strebt die Zwischensumme Z_n für $n \rightarrow \infty$ gegen einen Grenzwert, dann bezeichnen wir diesen als *Dreifachintegral* oder als *dreidimensionales Gebietsintegral*.

Definition: (Dreifachintegral; dreidimensionales Gebietsintegral). Sei $G \subset \mathbb{R}^3$ und $f: G \rightarrow \mathbb{R}$ stetig. Der Grenzwert

$$\iiint\limits_{(G)} f(x, y, z) dG = \lim_{n \rightarrow \infty} \sum_{i=1}^n f(x_i, y_i, z_i) dV_i$$

bezeichnet man als **Dreifachintegral** bzw. **dreidimensionales Gebietsintegral**.

Die Berechnung von Dreifachintegralen wird auf einfache Integrationen zurückgeführt, wobei natürlich die Beschreibung der Integrationsgrenzen komplizierter wird als für Doppelintegrale. Die Vorgehensweise bei der Berechnung werden wir an dem folgenden Beispiel verdeutlichen.

Beispiel 11.11.

$$I = \iiint\limits_{(G)} f(x, y, z) dG = \int_0^1 \int_0^x \int_{-y^2}^{x^2} (1+x) dz dy dx .$$

Die Integrationsreihenfolge wird durch die Reihenfolge der Differenziale dz , dy , dx festgelegt und zwar von Innen nach Außen. Das innerste Integral wird integriert über die Variable z . Die anderen Variablen x und y sind dabei konstant.

$$I_1 = \int_{z=-y^2}^{z=x^2} (1+x) dz = [(1+x) z]_{z=-y^2}^{z=x^2} = (1+x) x^2 - (1+x) (-y^2) .$$

I_1 enthält als Ergebnis nicht mehr die Variable z , sondern nur noch x und y . Das verbleibende Doppelintegral

$$I = \int_0^1 \int_0^x [(1+x) x^2 + (1+x) y^2] dy dx$$

wird berechnet, indem zunächst wieder das innerste Integral über y ausgeführt wird

$$\begin{aligned} I_2 &= \int_{y=0}^{y=x} [(1+x) x^2 + (1+x) y^2] dy \\ &= \left[(1+x) x^2 y + \frac{1}{3} (1+x) y^3 \right]_{y=0}^{y=x} = \frac{4}{3} x^3 + \frac{4}{3} x^4 . \end{aligned}$$

I_2 enthält als Variable nur noch x , über die zuletzt integriert wird

$$I = \int_{x=0}^{x=1} \left(\frac{4}{3} x^3 + \frac{4}{3} x^4 \right) dx = \left[\frac{1}{3} x^4 + \frac{4}{15} x^5 \right]_0^1 = \frac{3}{5} .$$

□

➤ 11.2.2 Anwendungen

Die Anwendungsbeispiele sollen die Berechnungen von Dreifachintegralen verdeutlichen. In der Regel sind die Volumen, Masse, Schwerpunktskoordinaten und Trägheitsmomente starrer Körper über Dreifachintegrale zu berechnen. Für Spezialfälle kann durch Einführung eines angepassten Koordinatensystems die Rechnung erheblich vereinfacht werden; in manchen Fällen wird sie durch spezielle Koordinatensysteme erst möglich.

Die für die Anwendungen wichtigsten Systeme sind Polar-, Zylinder- und Kugelkoordinaten. Auf der CD-Rom sind in Abschnitt 11.3 die Formeln für die [Integration in transformierten Koordinatensystemen](#) hergeleitet. Zusammenfassend gilt

- (1) **Polarkoordinaten:** Bei Polarkoordinaten wird ein Punkt in der (x, y) -Ebene eindeutig durch die Angabe des Winkel φ , $0 \leq \varphi < 2\pi$, und des Radius $r \geq 0$ angegeben. Die Transformationsgleichungen lauten

$$x = r \cos \varphi, \quad y = r \sin \varphi.$$

Ein Doppelintegral lautet in Polarkoordinaten

$$\iint_{(x,y)} f(x, y) \, dx \, dy = \iint_{(r,\varphi)} f(r \cos \varphi, r \sin \varphi) \, r \, dr \, d\varphi.$$

- (2) **Zylinderkoordinaten:** Ein Punkt im (x, y, z) -Raum wird eindeutig durch die Angabe seiner Polarkoordinaten (r, φ) in der (x, y) -Ebene und zusätzlich seiner z -Komponente festgelegt. Ein Dreifachintegral lautet in Zylinderkoordinaten

$$\iiint_{(x,y,z)} f(x, y, z) \, dx \, dy \, dz = \iiint_{(r,\varphi,z)} f(r \cos \varphi, r \sin \varphi, z) \, r \, dr \, d\varphi \, dz.$$

- (3) **Kugelkoordinaten:** Durch die Angabe zweier Winkel φ und ϑ sowie dem Abstand zum Ursprung lässt sich jeder Punkt im $\mathbb{R}^3$ eindeutig festlegen:

$$x = r \cos \varphi \cos \vartheta, \quad y = r \sin \varphi \cos \vartheta, \quad z = r \sin \vartheta$$

Ein Dreifachintegral lautet in Kugelkoordinaten

$$\iiint_{(x,y,z)} f(x, y, z) \, dx \, dy \, dz = \iiint_{(r,\varphi,\vartheta)} f(r \cos \varphi \cos \vartheta, r \sin \varphi \cos \vartheta, r \sin \vartheta) \cdot r^2 \cos \vartheta \, dr \, d\varphi \, d\vartheta.$$

Wir geben im Folgenden eine Zusammenstellung der wichtigsten Formeln aus der Physik starrer Körper für Volumen, Masse, Schwerpunktskoordinaten und Trägheitsmomente an.

Ist $K \subset \mathbb{R}^3$ ein starrer Körper mit der ortsabhängigen Dichte $\rho = \rho(x, y, z)$, dann ist

$$V = \iiint_{(K)} dx \, dy \, dz \quad \text{das Volumen des Körpers und}$$

$$M = \iiint_{(K)} \rho(x, y, z) \, dx \, dy \, dz \quad \text{die Masse des Körpers.}$$

Die **Koordinaten des Schwerpunktes** $S(x_s, y_s, z_s)$ lauten

$$x_s = \frac{1}{M} \iiint_{(K)} x \cdot \rho(x, y, z) \, dx \, dy \, dz$$

$$y_s = \frac{1}{M} \iiint_{(K)} y \cdot \rho(x, y, z) \, dx \, dy \, dz$$

$$z_s = \frac{1}{M} \iiint_{(K)} z \cdot \rho(x, y, z) \, dx \, dy \, dz.$$

Rotiert ein starrer Körper um die Drehachse x , dann heißt das Integral

$$I_x = \iiint_{(K)} \rho(x, y, z) (y^2 + z^2) \, dx \, dy \, dz$$

das **Trägheitsmoment bezüglich der x -Achse**. Die Trägheitsmomente bezüglich der y - und z -Achse sind entsprechend definiert.

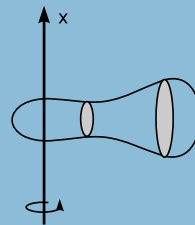


Abb. 11.16.

Zur Bestimmung des Trägheitsmomentes eines Körpers K bezüglich einer beliebigen Achse A wendet man den sog. *Steinerschen Satz* an:

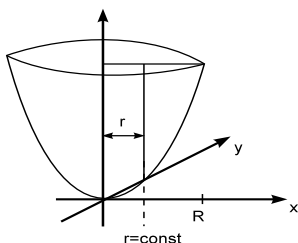
Steinerscher Satz: Für eine zur Schwerpunktsachse S im Abstand d parallel verlaufenden Rotationsachse A gilt

$$I_A = I_S + M d^2,$$

wenn I_S das Trägheitsmoment bezüglich der Schwerpunktsachse und M die Masse des Körpers ist.

Nach dem Steinerschen Satz genügt es jeweils nur parallele Achsen durch den Schwerpunkt zu berücksichtigen.

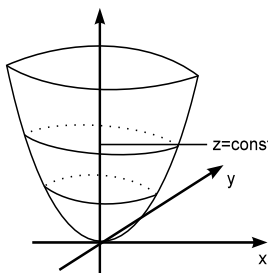
Anwendungsbeispiel 11.12 (Volumenberechnung). Gesucht ist das Volumen des Rotationskörpers, der durch Rotation von x^2 an der y -Achse entsteht. Zur Berechnung führen wir Zylinderkoordinaten ein:



- φ -Integration: $\varphi = 0$ bis $\varphi = 2\pi$
unabhängig von r und z .
 z -Integration: Bei konstantem r
variiert z von $z = r^2$ bis $z = R^2$.
 r -Integration: $r = 0$ bis $r = R$.

$$\begin{aligned} V &= \iiint_{(K)} dx dy dz = \iiint_{(K)} r d\varphi dr dz = \int_{r=0}^R \int_{z=r^2}^{R^2} \int_{\varphi=0}^{2\pi} r d\varphi dz dr \\ &= \int_{r=0}^R \int_{z=r^2}^{R^2} r 2\pi dz dr = \int_{r=0}^R [r 2\pi z]_{z=r^2}^{z=R^2} dr \\ &= 2\pi \int_{r=0}^R (R^2 r - r^3) dr = 2\pi \left[\frac{1}{2} R^2 r^2 - \frac{1}{4} r^4 \right]_{r=0}^R = \frac{\pi}{2} R^4. \end{aligned}$$

Alternativ kann die Integration auch "scheibenweise" durchgeführt werden:



- φ -Integration: $\varphi = 0$ bis $\varphi = 2\pi$
unabhängig von r und z .
 r -Integration: Bei konstantem z
variiert r von $r = 0$ bis $r = \sqrt{z}$.
 z -Integration: $z = 0$ bis $z = R^2$.

$$\Rightarrow V = \int_{z=0}^{R^2} \int_{r=0}^{\sqrt{z}} \int_{\varphi=0}^{2\pi} r d\varphi dr dz = \frac{\pi}{2} R^4.$$

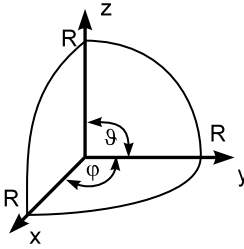
□

Anwendungsbeispiel 11.13 (Schwerpunktskoordinaten).

Gesucht sind die Schwerpunktskoordinaten eines Kugelausschnitts mit Dichte $\rho = 1$ und Radius R .

(i) Berechnung des Volumens in Kugelkoordinaten:

$$V = \iiint_{(K)} r^2 \cos \vartheta \, dr \, d\varphi \, d\vartheta$$



r -Integration: $r = 0$ bis $r = R$

unabhängig von ϑ und φ .

φ -Integration: $\varphi = 0$ bis $\varphi = \frac{\pi}{2}$

unabhängig von r und ϑ .

ϑ -Integration: $\vartheta = 0$ bis $\vartheta = \frac{\pi}{2}$

unabhängig von r und φ .

$$\begin{aligned} V &= \int_{\vartheta=0}^{\pi/2} \int_{\varphi=0}^{\pi/2} \left(\int_{r=0}^R r^2 \cos \vartheta \, dr \right) d\varphi \, d\vartheta \\ &= \int_{\vartheta=0}^{\pi/2} \cos \vartheta \int_{\varphi=0}^{\pi/2} \left[\frac{r^3}{3} \right]_{r=0}^{r=R} d\varphi \, d\vartheta = \frac{1}{3} R^3 \int_{\vartheta=0}^{\pi/2} \cos \vartheta \int_{\varphi=0}^{\pi/2} d\varphi \, d\vartheta \\ &= \frac{1}{3} R^3 \int_{\vartheta=0}^{\pi/2} \cos \vartheta \cdot \frac{\pi}{2} d\vartheta = \frac{1}{6} R^3 \pi \int_{\vartheta=0}^{\pi/2} \cos \vartheta \, d\vartheta = \frac{1}{6} R^3 \pi. \end{aligned}$$

(ii) Berechnung der Schwerpunktskoordinate x_s :

$$\begin{aligned} x_s &= \frac{1}{V} \iiint_{(K)} x r^2 \cos \vartheta \, dr \, d\varphi \, d\vartheta \\ &= \frac{1}{V} \int_{\vartheta=0}^{\pi/2} \int_{\varphi=0}^{\pi/2} \int_{r=0}^R r \cos \varphi \cos \vartheta r^2 \cos \vartheta \, dr \, d\varphi \, d\vartheta \\ &= \frac{1}{V} \int_{\vartheta=0}^{\pi/2} \cos^2 \vartheta \int_{\varphi=0}^{\pi/2} \cos \varphi \int_{r=0}^R r^3 \, dr \, d\varphi \, d\vartheta \\ &= \frac{1}{V} \int_{\vartheta=0}^{\pi/2} \cos^2 \vartheta \int_{\varphi=0}^{\pi/2} \cos \varphi \frac{1}{4} R^4 \, d\varphi \, d\vartheta = \frac{1}{V} \int_{\vartheta=0}^{\pi/2} \cos^2 \vartheta \frac{1}{4} R^4 \, d\vartheta \\ &= \frac{1}{V} \frac{1}{4} R^4 \left[\frac{1}{2} \vartheta + \frac{1}{2} \cos \vartheta \sin \vartheta \right]_0^{\pi/2} = \frac{1}{V} \frac{1}{16} R^4 \pi = \frac{3}{8} R. \end{aligned}$$

(iii) analog berechnen sich $y_s = z_s = \frac{3}{8} R$.

□

Anwendungsbeispiel 11.14 (Massenträgheitsmoment eines Würfels).

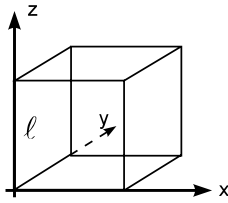


Abb. 11.17. Würfel

Gesucht ist das Massenträgheitsmoment eines homogenen Würfels bezüglich der z -Achse.

$$I_z = \iiint_{(K)} (x^2 + y^2) \rho \, dx \, dy \, dz$$

$$= \rho \int_{x=0}^{x=l} \left(\int_{y=0}^{y=l} \left(\int_{z=0}^{z=l} (x^2 + y^2) \, dz \right) dy \right) dx$$

$$= \rho \int_{x=0}^l \int_{y=0}^l (x^2 + y^2) \left(\int_0^l dz \right) dy \, dx$$

$$= \rho l \int_{x=0}^l \int_{y=0}^l (x^2 + y^2) \, dy \, dx$$

$$= \rho l \int_{x=0}^l \left[x^2 y + \frac{y^3}{3} \right]_{y=0}^l dx = \rho l \int_{x=0}^l (x^2 l + \frac{1}{3} l^3) \, dx$$

$$= \rho l \left[\frac{1}{3} x^3 l + \frac{1}{3} l^3 x \right]_0^l = \rho l \left(\frac{1}{3} l^4 + \frac{1}{3} l^4 \right) = \frac{2}{3} \rho l^5.$$

Mit der Masse $M = \rho \cdot V = \rho \cdot l^3$ folgt $I_z = \frac{2}{3} M l^2$.

□

Anwendungsbeispiel 11.15 (Massenträgheitsmoment eines Zylinders).

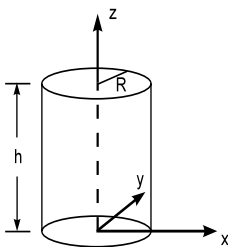


Abb. 11.18. Zylinder

Gesucht ist das Trägheitsmoment eines Zylinders der Höhe H und Grundfläche πR^2 bezüglich der z -Achse. Zur Berechnung des Trägheitsmomentes verwenden wir Zylinderkoordinaten

$$I_z = \rho \iiint_{(K)} (x^2 + y^2) \, dx \, dy \, dz$$

$$= \rho \iiint_{(K)} r^2 \, r \, dr \, d\varphi \, dz$$

$$= \rho \int_{z=0}^H \int_{\varphi=0}^{2\pi} \int_{r=0}^R r^3 \, dr \, d\varphi \, dz = \rho \frac{R^4}{4} \cdot 2\pi \cdot H.$$

Mit $\rho = \frac{M}{V} = \frac{M}{\pi R^2 H}$ folgt $I_z = M \frac{1}{2} R^2$.

□

MAPLE-Worksheets zu Kapitel 11



Auf der CD-Rom ist die MAPLE-Prozedur **starr** enthalten, die alle Formeln zur Bestimmung von Volumen, Schwerpunktskoordinaten und Trägheitsmomenten starrer Körper sowie die Berechnung von Dreifachintegralen über dreidimensionalen Gebieten zusammenfasst. Wahlweise kann mit kartesischen Koordinaten, Zylinder- oder Kugelkoordinaten gearbeitet werden.

Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 11 mit MAPLE zur Verfügung.

- [Doppelintegrale mit MAPLE](#)
- [Dreifachintegrale mit MAPLE](#)
- [Linienintegrale mit MAPLE](#)
- [Oberflächenintegrale mit MAPLE](#)

11.3 Aufgaben zur Integralrechnung

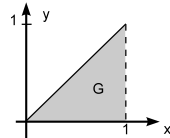
11.3

11.1 Berechnen Sie die folgenden Doppelintegrale

$$\begin{array}{ll} \text{a) } \int_{x=0}^1 \int_{y=1}^l \frac{x^2}{y} dy dx & \text{b) } \int_{x=0}^3 \int_{y=0}^{1-x} (25 - x^2 - y^2) dy dx \\ \text{c) } \int_{y=0}^{\pi} \int_{x=\pi/2}^{\pi} \sin(x+y) dx dy & \text{d) } \int_{y=0}^{\pi} \int_{x=\pi}^y x \cdot \cos(x+y) dx dy \end{array}$$

11.2 Bestimmen Sie das Doppelintegral über das schraffierte Gebiet G für die Funktion $z = x - y$, indem Sie sowohl Integralformel (D1) als auch (D2) anwenden:

$$\begin{aligned} I = \iint_G (x - y) dG &= \int_{x=0}^1 \int_{y=0}^x (x - y) dy dx \\ &= \int_{y=0}^1 \int_{x=y}^1 (x - y) dx dy \end{aligned}$$



11.3 Zeigen Sie, dass der Wert der beiden Gebietsintegrale I_1 und I_2 gleich ist

$$I_1 = \int_{x=0}^2 \int_{y=0}^{x^2} 2xy dy dx \quad I_2 = \int_{y=0}^4 \int_{x=\sqrt{y}}^2 2xy dx dy$$

11.4 Bestimmen Sie den Flächeninhalt des Halbkreises mit Radius $R = 2$ und Mittelpunkt $(2, 0)$ in der oberen Halbebene, indem Sie in Polarkoordinaten

das folgende Integral berechnen

$$\iint_{(G)} 1 \, dx \, dy = \int_{r=0}^2 \int_{\varphi=0}^{\pi} r \, d\varphi \, dr .$$

- 11.5 Gegeben sind die Kurven von $y = -x(x-3)$ und $y = -2x$.
 a) Welche Fläche schließen sie ein?
 b) Wie lauten die Koordinaten des Flächenschwerpunktes?
- 11.6 Bestimmen Sie die axialen Flächenmomente I_x und I_y sowie das polare Flächenmoment I_p eines Viertelkreises mit Radius R .
- 11.7 Bestimmen Sie den Schwerpunkt der Dreiecksfläche (Abb. a).
- 11.8 Bestimmen Sie den Schwerpunkt des Halbkreises mit Radius R (Abb. b).

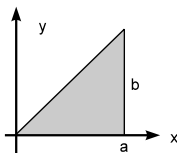


Abb. a

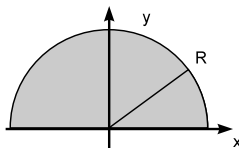


Abb. b

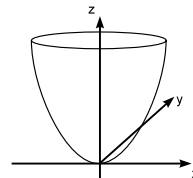


Abb. c

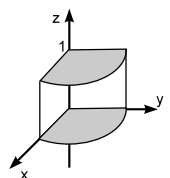


Abb. d

- 11.9 Berechnen Sie die Dreifachintegrale

a) $\int_{z=0}^1 \int_{y=z-1}^z \int_{x=y}^{y+1} x^2 \, dx \, dy \, dz$

b) $\int_{z=-l}^l \int_{x=z}^{z^2} \int_{y=x+z}^{x+z} x y z \, dy \, dx \, dz$

c) $\int_{\varphi=0}^{\pi} \int_{\vartheta=-\pi/2}^{\pi/2} \int_{r=0}^R r^2 \cos \vartheta \sin \varphi \, dr \, d\vartheta \, d\varphi$

d) $\int_{x=-R}^R \int_{y=0}^R \int_{r=0}^{\sqrt{x^2+y^2}} r \, dr \, dy \, dx$

- 11.10 Bestimmen Sie die Schwerpunktskoordinate z_s sowie die Massenträgheitsmomente des Rotationskörpers, der durch Rotation von x^2 an der z -Achse entsteht. Führen Sie zur Beschreibung des Körpers Zylinderkoordinaten ein (vgl. Abb. c).
- 11.11 Bestimmen Sie die Massenträgheitsmomente einer Halbkugel ($z > 0$).
- 11.12 Gesucht ist das Integral

$$I = \iiint_G x^2 y \, dx \, dy \, dz$$

wobei

$$G = \{(x, y, z) : x \geq 0, y \geq 0, x^2 + y^2 \leq 1, 0 \leq z \leq 1\} .$$

(Zur Berechnung führe man Zylinderkoordinaten ein; vgl. Abb. d.)

- 11.13 Gesucht ist der Schwerpunkt des Zylinderstücks aus Aufgabe 11.12.

Kapitel 12

Gewöhnliche Differenzialgleichungen

12

12	Gewöhnliche Differenzialgleichungen	473
12.1	Differenzialgleichungen erster Ordnung	476
12.1.1	Einleitende Problemstellungen	476
12.1.2	Lösen der homogenen Differenzialgleichung	480
12.1.3	Lösen der inhomogenen Differenzialgleichung	482
12.1.4	Lineare Differenzialgleichungen mit konstantem Koeffizient	490
12.1.5	Nichtlineare Differenzialgleichungen 1. Ordnung	494
12.1.6	Numerisches Lösen von DG 1. Ordnung	497
12.1.7	Nichtlinearer DG 1. Ordnung (Ergänzungen)	cd
12.1.8	Lösen von DG 1. Ordnung mit MAPLE	cd
12.2	Lineare Differenzialgleichungssysteme	501
12.2.1	Einführung	501
12.2.2	Homogene lineare Differenzialgleichungssysteme	504
12.2.3	Eigenwerte und Eigenvektoren	508
12.2.4	Lösen von homogenen LDGS mit konstanten Koeffizienten	513
12.2.5	Eigenwerte und Eigenvektoren mit MAPLE	cd
12.2.6	Berechnung spezieller Lösungen für inhomogene LDGS	cd
12.3	Lineare Differenzialgleichungen n -ter Ordnung	523
12.3.1	Einleitende Beispiele	523
12.3.2	Reduktion einer linearen DG n -ter Ordnung auf ein System	526
12.3.3	Homogene DG n -ter Ordnung mit konst. Koeffizienten	531
12.3.4	Inhomogene DG n -ter Ord. mit konstanten Koeffizienten	541
12.3.5	Lösen von DG n -ter Ordnung mit MAPLE	cd
12.4	Aufgaben zu Differenzialgleichungen	555
Zusätzliche Abschnitte auf der CD-Rom		
12.5	Numerisches Lösen von Differenzialgleichungen 1. Ordnung	cd
12.5.1	Streckenzugverfahren von Euler	cd
12.5.2	Verfahren höherer Ordnung	cd
12.5.3	Vergleich der numerischen Verfahren mit MAPLE	cd
12.5.4	Numerisches Lösen von DG 1. Ordnung mit MAPLE: dsolve	cd
12.6	Beschreibung elektrischer Filterschaltungen	cd
12.6.1	Physikalische Gesetzmäßigkeiten der Bauelemente	cd
12.6.2	Aufstellen der DG für elektrische Schaltungen	cd
12.6.3	Aufstellen und Lösen der DG für Filterschaltungen	cd

12 Gewöhnliche Differenzialgleichungen

Differenzialgleichungen sind für die Natur- und Ingenieurwissenschaften unentbehrlich, da durch sie viele Naturgesetze ausgedrückt werden. Differenzialgleichungen sind das Ergebnis einer mathematisch-physikalischen Modellierung, welche die auftretenden Phänomene möglichst gut beschreibt. Aber nicht nur das Lösen von Differenzialgleichungen innerhalb der Mathematik ist für den Ingenieur wichtig, sondern gerade auch das Aufstellen dieser Modellgleichungen. Daher werden wir in jedem Abschnitt zunächst für anwendungsrelevante Beispiele die Herleitung von Differenzialgleichungen beschreiben und anschließend auf das systematische Lösen der unterschiedlichen Typen von Differenzialgleichungen eingehen.

Begriffsbestimmung: Eine **Differenzialgleichung** ist eine Gleichung, in der neben der gesuchten Funktion (oder mehreren Funktionen) auch Ableitungen dieser Funktion (bzw. Funktionen) vorkommen. Eine **gewöhnliche Differenzialgleichung** ist eine Gleichung, in der nur Funktionen und deren *gewöhnliche* Ableitungen auftreten, im Gegensatz zu **partiellen Differenzialgleichungen**, bei denen auch *partielle* Ableitungen in der Bestimmungsgleichung enthalten sind.

Wir werden in diesem Kapitel nur gewöhnliche Differenzialgleichungen behandeln und daher im Folgenden den Zusatz *gewöhnlich* unterdrücken. Die **partiellen Differenzialgleichungen** werden exemplarisch in einem separaten CD-Kapitel auf der behandelt. Beim Auftreten von nur einer unbekannten Funktion hat eine gewöhnliche Differenzialgleichung die allgemeine Form

$$F\left(x, y(x), y'(x), \dots, y^{(n)}(x)\right) = 0.$$

Die Ordnung der höchsten in einer Differenzialgleichung auftretenden Ableitung heißt **Ordnung der Differenzialgleichung**. Eine Differenzialgleichung heißt **linear**, wenn alle Ableitungen der Funktion sowie die Funktion selbst linear (d.h. proportional) vorkommen. Ansonsten heißt die Differenzialgleichung **nichtlinear**. Im Folgenden gehen wir davon aus, dass die Differenzialgleichung die Form

$$y^{(n)}(x) + a_{n-1}(x) y^{(n-1)}(x) + \dots + a_1(x) y'(x) + a_0(x) y(x) = f(x)$$

besitzen, wobei die Koeffizienten $a_i(x)$ und die rechte Seite $f(x)$ bekannte, gegebene, stetige Funktionen sind. Obige Differenzialgleichung ist linear, weil die gesuchte Funktion $y(x)$ und all ihre Ableitungen nur in linearer Form auftreten.

In Abschnitt 12.1 behandeln wir lineare Differenzialgleichungen erster Ordnung sowie den Spezialfall mit konstanten Koeffizienten. In Abschnitt 12.2 werden Systeme von linearen Differenzialgleichungen mit Hilfe der Eigenwerte und Eigenvektoren der zugehörigen Systemmatrizen gelöst, während in 12.3 lineare Differenzialgleichungen n -ter Ordnung mit konstanten Koeffizienten behandelt werden.

Hinweis: Auf der CD-Rom befindet sich zusätzlich ein separater Abschnitt 12.4 über das [numerische Lösen von Differenzialgleichungen](#) erster Ordnung mit einem qualitativen Vergleich der Euler-, Prädiktor-Korrektor- und Runge-Kutta-Verfahren sowie die [Anwendung der numerischen Verfahren](#) mit MAPLE auf zahlreiche Anwendungsbeispiele. In Abschnitt 12.5 werden Differenzialgleichungssysteme 1. Ordnung am Beispiel von [elektrischen Filterschaltungen](#) mit dem Euler-Verfahren numerisch gelöst.

12.1 Differenzialgleichungen erster Ordnung

Wir beschäftigen uns in diesem Abschnitt hauptsächlich mit linearen Differenzialgleichungen erster Ordnung. Wichtig bei der mathematischen Darstellung der Lösung ist, dass die allgemeine Lösung der Differenzialgleichung sich zerlegen lässt in die allgemeine Lösung des *homogenen* Problems plus einer *partikulären* Lösung der inhomogenen Differenzialgleichung. Diese spezielle Struktur der Lösung überträgt sich sowohl auf die linearen Differenzialgleichungssysteme wie auch auf die Differenzialgleichungen n -ter Ordnung.

Die Methoden, die wir zum Lösen der Differenzialgleichungen erster Ordnung einführen sind Trennung der Variablen für das homogene und Variation der Konstanten für das inhomogene Problem. Die Lösungsmethoden werden anhand zahlreicher Anwendungsbeispiele geübt. Dieser Abschnitt beinhaltet auch einen Ausblick über das numerische Lösen von Differenzialgleichungen.

Hinweis: In zusätzlichen Abschnitten auf der CD-Rom werden [numerische Verfahren](#) zum Lösen von Differenzialgleichungen eingeführt, [Differenzialgleichungen mit MAPLE](#) gelöst sowie viele weitere Anwendungsbeispiele diskutiert.

➤ 12.1.1 Einleitende Problemstellungen

Im Folgenden werden einige Problemstellungen aus Physik, Elektronik und anderen Anwendungen modelliert. Diese Beispiele zeigen auf, wie man vom physikalischen Problem zu einer mathematischen Modellgleichung, nämlich in diesem Fall zu einer Differenzialgleichung erster Ordnung, kommt.

Anwendungsbeispiel 12.1 (RL-Kreis).

Ein Widerstand R , eine Spule mit Induktivität L und eine Batterie mit Spannung U_B sind mit einem Schalter S in Reihe geschaltet. Der Schalter ist zunächst offen und wird zur Zeit $t = 0$ geschlossen. Zum Zeitpunkt $t = 0$ ist der Strom Null: $I(0) = 0$. Wie verhält sich der Strom $I(t)$ als Funktion der Zeit für $t > 0$?

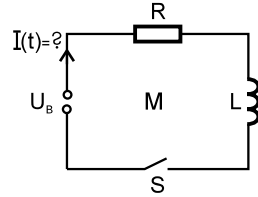


Abb. 12.1. RL-Kreis

Zur Lösung des Problems stellen wir zunächst die Bestimmungsgleichung für den Strom auf. Nach dem Maschensatz gilt für die Masche M , dass der Spannungsabfall entlang R plus dem Spannungsabfall entlang L gleich der angelegten Spannung U_B ist:

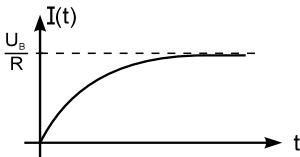
$$U_R + U_L = U_B$$

Mit dem Ohmschen Gesetz ($U_R = R \cdot I(t)$) und dem Induktionsgesetz ($U_L = L \frac{dI(t)}{dt}$) folgt weiter

$$RI(t) + L \frac{dI(t)}{dt} = U_B$$

$$\Rightarrow \frac{d}{dt} I(t) + \frac{R}{L} I(t) = \frac{1}{L} U_B \quad \text{mit } I(0) = 0.$$

Dies ist eine gewöhnliche, lineare Differenzialgleichung erster Ordnung für den Strom $I(t)$. Gesucht ist eine Funktion $I(t)$, welche die obige Differenzialgleichung mit der Anfangsbedingung erfüllt.

Abb. 12.2. Stromverlauf $I(t)$

Wie man durch Nachrechnen bestätigt, ist die Lösung gegeben durch

$$I(t) = \frac{U_B}{R} \left(1 - e^{-\frac{R}{L} t} \right).$$

Denn setzen wir die Ableitung von $I(t)$

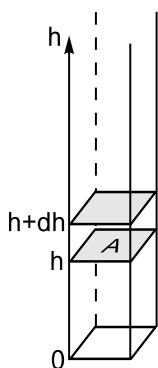
$$\dot{I}(t) = \frac{U_B}{R} \cdot \frac{R}{L} e^{-\frac{R}{L} t}$$

und die Funktion in die Differenzialgleichung ein, folgt

$$\Rightarrow \dot{I}(t) + \frac{R}{L} I(t) = \frac{U_B}{L} e^{-\frac{R}{L} t} + \frac{R}{L} \frac{U_B}{R} \left(1 - e^{-\frac{R}{L} t} \right) = \frac{U_B}{L}.$$

Die Funktion $I(t)$ erfüllt also die Differenzialgleichung und besitzt die geforderte Anfangsbedingung $I(0) = \frac{U_B}{R} (1 - e^0) = 0$. $\square$

Anwendungsbeispiel 12.2 (Barometrische Höhenformel).



Der Luftdruck $p(h)$ in der Höhe h über dem Meeresspiegel wird verursacht durch das Gewicht der über der Fläche A lastenden Luftsäule. Der Druckunterschied $p(h) - p(h + dh)$ ist gleich dem Gewicht G der vertikalen Luftsäule mit Querschnitt A , die sich zwischen h und $h + dh$ befindet. Für kleine dh nehmen wir an, dass die Dichte $\rho(h)$ in dieser Luftsäule konstant ist:

$$A (p(h) - p(h + dh)) = G = dm g = \rho(h) A \cdot dh g .$$

Die Division durch dh und der anschließende Grenzübergang $dh \rightarrow 0$ liefert

Abb. 12.3.

Höhenformel

$$p'(h) = \lim_{dh \rightarrow 0} \frac{p(h + dh) - p(h)}{dh} = -g \rho(h) .$$

Wird die Luft als ideales Gas betrachtet, so gilt folgender Zusammenhang zwischen ρ , p und der Temperatur T : $\rho = \alpha \frac{p}{T}$ mit einer Konstanten α .

- (i) Betrachtet man die Temperatur T als konstant und unabhängig von der Höhe und führt die Konstante $\beta = g \frac{\alpha}{T}$ ein, erhält man die lineare Differenzialgleichung 1. Ordnung mit Anfangsbedingung $p(h_0) = p_0$

$$p'(h) = -\alpha g \frac{p(h)}{T} = -\beta p(h) .$$

Die Lösung dieser Differenzialgleichung lautet

$$p(h) = p_0 e^{-\beta(h-h_0)}$$

(Barometrische Höhenformel).

- (ii) Die obige Differenzialgleichung ist wegen der Annahme $T = \text{const}$ nur in einem kleinen Bereich gültig. In Realität fällt die Temperatur mit zunehmender Höhe. Die einfachste Modellannahme ist, dass T einen linearen Temperaturabfall besitzt:

$$T(h) = T_0 - b(h - h_0) .$$

Dies führt auf die lineare Differenzialgleichung erster Ordnung

$$p'(h) = \frac{-\alpha g}{T_0 - b(h - h_0)} p(h) \quad \text{mit } p(h_0) = p_0 .$$

Wie man wieder durch Nachrechnen bestätigt, ist

$$p(h) = p_0 \left(1 - \frac{b}{T_0} (h - h_0) \right)^{\frac{\alpha g}{b}}$$

die Lösung der Differenzialgleichung. $\square$

Anwendungsbeispiel 12.3 (Radioaktiver Zerfall).

Sei $n(t)$ die Anzahl der zum Zeitpunkt t gegebenen Atome einer radioaktiven Substanz. Die Menge dieser Substanz, die in einer Zeitspanne dt zerfällt, ist proportional zur vorhandenen Substanz und zur Zeitspanne dt :

$$n(t + dt) - n(t) \sim -dt \cdot n(t).$$

Führt man die Proportionalitätskonstante $\lambda > 0$ ein, gilt

$$n(t + dt) - n(t) = -\lambda dt n(t).$$

Da die Anzahl der radioaktiven Atome abnimmt, steht auf der rechten Seite der Gleichung ein Minus als Vorzeichen. Division durch dt und anschließender Grenzübergang $dt \rightarrow 0$ liefert

$$n'(t) = \lim_{dt \rightarrow 0} \frac{n(t + dt) - n(t)}{dt} = -\lambda n(t) \quad (\text{Zerfallsgesetz}).$$

Dies ist eine *lineare Differenzialgleichung erster Ordnung* mit der Anfangsbedingung $n(0) = N$. Diese Differenzialgleichung hat als Lösung

$$n(t) = N e^{-\lambda t},$$

was man durch direktes Einsetzen der Funktion $n(t)$ in die Differenzialgleichung wieder nachprüft. $\square$

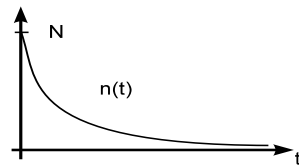


Abb. 12.4. Radioaktiver Zerfall

Im Folgenden werden wir klären, wie man systematisch Lösungen von linearen Differenzialgleichungen 1. Ordnung bestimmt. Mathematisch erhält man die folgende, allgemeine Problemstellung:

Allgemeine Problemstellung: Wir betrachten in einem Intervall I die lineare Differenzialgleichung 1. Ordnung

$$y'(x) = h(x)y(x) + f(x), \quad (D1)$$

wenn $h(x)$ und $f(x)$ gegebene, auf dem Intervall I stetige Funktionen sind. Für

$f(x) \neq 0$ heißt (D1) eine **inhomogene** Differenzialgleichung und für $f(x) = 0$ heißt (D1) eine **homogene** Differenzialgleichung.

Im Falle einer inhomogenen Differenzialgleichung nennt man $f(x) \neq 0$ die **Inhomogenität** oder die **Störfunktion**.

► 12.1.2 Lösen der homogenen Differenzialgleichung

Wir behandeln zunächst das homogene Problem

$$y'(x) = h(x)y(x)$$

mit der Anfangsbedingung $y(x_0) = y_0$. Die Lösung dieser Differenzialgleichung erfolgt durch die Methode der **Trennung der Variablen**. Dazu ersetzen wir $y'(x)$ durch $\frac{dy}{dx}$ und trennen die Variablen, indem wir formal die Gleichung mit dx multiplizieren und durch y dividieren:

$$\frac{dy}{dx} = h(x)y(x) \Rightarrow \frac{dy}{y} = h(x)dx.$$

Die anschließende Integration liefert

$$\int_{y_0}^y \frac{d\tilde{y}}{\tilde{y}} = \int_{x_0}^x h(\tilde{x}) d\tilde{x} \Rightarrow \ln \tilde{y}|_{y_0}^y = \ln \frac{y}{y_0} = \int_{x_0}^x h(\tilde{x}) d\tilde{x}.$$

Wendet man auf beiden Seiten der Gleichung die Exponentialfunktion an, erhält man als Lösung

$$y(x) = y_0 e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}. \quad (H)$$

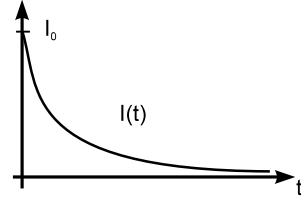
Mit Hilfe der Formel (H) ist man in der Lage durch die Berechnung des bestimmten Integrals $\int_{x_0}^x h(\tilde{x}) d\tilde{x}$ jede homogene lineare Differenzialgleichung erster Ordnung zu lösen. In den Anwendungen wird oftmals auch die Methode Trennung der Variablen statt dieser Lösungsformel verwendet.

Anwendungsbeispiel 12.4 (RL-Kreis).

Der unter Beispiel 12.1 diskutierte RL-Kreis ist zunächst geschlossen. Zum Zeitpunkt $t_0 = 0$ wird die Batterie überbrückt. Dann gilt

$$\frac{d}{dt} I(t) + \frac{R}{L} I(t) = 0 \quad \text{mit } I(0) = I_0.$$

$$\hookrightarrow \boxed{\dot{I}(t) = -\frac{R}{L} I(t)}.$$

Abb. 12.5. Stromverlauf $I(t)$

Die Lösungsformel (H) für homogene Differenzialgleichungen liefert

$$I(t) = I_0 e^{\int_{t_0}^t \left(-\frac{R}{L}\right) d\tau} = I_0 e^{-\frac{R}{L}(t-t_0)} = I_0 e^{-\frac{R}{L}t}. \quad \square$$

Anwendungsbeispiel 12.5 (Barometrische Höhenformel).

Die unter Beispiel 12.2 (i) und (ii) angegebenen Differenzialgleichungen werden ebenfalls mit der Formel (H) gelöst:

(i) $p'(h) = -\beta p(h)$:

$$\hookrightarrow p(h) = p(h_0) e^{\int_{h_0}^h -\beta d\tilde{h}} = p(h_0) e^{-\beta(h-h_0)}.$$

(ii) $p'(h) = -\frac{\alpha g}{T_0 - b(h-h_0)} p(h)$:

$$\hookrightarrow p(h) = p(h_0) e^{\int_{h_0}^h \frac{-\alpha g}{T_0 - b(\tilde{h}-h_0)} d\tilde{h}} = p(h_0) e^{\frac{\alpha g}{b} \ln(T_0 - b(\tilde{h}-h_0)) \Big|_{h_0}^h}.$$

Durch Einsetzen der oberen und unteren Grenze folgt für das Argument der Exponentialfunktion

$$\begin{aligned} \frac{\alpha g}{b} \ln \left(T_0 - b(\tilde{h} - h_0) \right) \Big|_{h_0}^h &= \frac{\alpha g}{b} [\ln(T_0 - b(h - h_0)) - \ln(T_0)] \\ &= \frac{\alpha g}{b} \ln \frac{T_0 - b(h - h_0)}{T_0} = \ln \left(1 - \frac{b}{T_0} (h - h_0) \right)^{\frac{\alpha g}{b}}. \end{aligned}$$

Wir erhalten damit insgesamt für die Funktion $p(h)$

$$p(h) = p(h_0) \left(1 - \frac{b}{T_0} (h - h_0) \right)^{\frac{\alpha g}{b}}. \quad \square$$

➤ 12.1.3 Lösen der inhomogenen Differenzialgleichung

Wir gehen nun zur inhomogenen Differenzialgleichung

$$y'(x) = h(x)y(x) + f(x)$$

mit der Anfangsbedingung $y(x_0) = y_0$ über und berechnen die Lösung der Differenzialgleichung mit der Methode der **Variation der Konstanten**.

Die Lösung des zugehörigen *homogenen* Problems $y'(x) = h(x)y(x)$ ist

$$y(x) = c e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}.$$

Um ausgehend von dieser Lösung eine Lösung der *inhomogenen* Differenzialgleichung zu erhalten, variieren wir die Konstante c , indem wir c als Funktion $c(x)$ zulassen und für $y(x)$ den **Produktansatz**

$$y(x) = c(x) \cdot \varphi(x) \quad \text{mit} \quad \varphi(x) = e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}$$

wählen. Um diesen Ansatz in die Differenzialgleichung einzusetzen, differenzieren wir $y(x)$ mit der Produktregel

$$\begin{aligned} y'(x) &= c'(x) \varphi(x) + c(x) \varphi'(x) \\ &= c'(x) \varphi(x) + c(x) \varphi(x) \cdot h(x), \end{aligned}$$

da $\varphi'(x) = e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}} \cdot h(x) = \varphi(x) h(x)$. Ersetzen wir $c(x) \cdot \varphi(x) = y(x)$ und setzen den Ansatz nun in die Differenzialgleichung ein, folgt weiter

$$y'(x) = c'(x) \varphi(x) + h(x) \cdot y(x) \stackrel{!}{=} f(x) + h(x) \cdot y(x).$$

Damit ist $y(x)$ Lösung der inhomogenen Differenzialgleichung, wenn

$$c'(x) \varphi(x) = f(x) \Rightarrow c'(x) = \frac{f(x)}{\varphi(x)}.$$

Die anschließende Integration liefert

$$c(x) = c_0 + \int_{x_0}^x \frac{f(\tilde{x})}{\varphi(\tilde{x})} d\tilde{x}.$$

Die Lösung der inhomogenen Differenzialgleichung $y(x) = c(x) \cdot \varphi(x)$ ergibt sich damit zu

$$y(x) = \varphi(x) \cdot \left(c_0 + \int_{x_0}^x \frac{f(\tilde{x})}{\varphi(\tilde{x})} d\tilde{x} \right).$$

Die Konstante c_0 muss so gewählt werden, dass $y(x_0) = y_0 \Rightarrow \boxed{c_0 = y_0}$.

Satz über lineare Differenzialgleichungen 1. Ordnung:

Seien $h, f : I \rightarrow \mathbb{R}$ gegebene, stetige Funktionen. Die lineare Differenzialgleichung 1. Ordnung

$$y'(x) = h(x) y(x) + f(x)$$

$$y(x_0) = y_0$$

besitzt auf dem Intervall I **genau eine** Lösung. Diese einzige Lösung ist gegeben durch

$$y(x) = \varphi(x) \left(y_0 + \int_{x_0}^x \frac{f(\tilde{x})}{\varphi(\tilde{x})} d\tilde{x} \right) \quad (I)$$

mit

$$\varphi(x) = e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}.$$

Bemerkung zur Methode: Immer dann, wenn Teillösungen einer Differenzialgleichung bekannt sind, versucht man weitere bzw. andere Lösungen der Differenzialgleichung zu konstruieren, indem man die Teilinformation in einem speziellen Ansatz berücksichtigt. Im Falle der inhomogenen Differenzialgleichung ist die Teilinformation die Kenntnis der homogenen Lösung $\varphi(x)$. Die Idee der Variation der Konstanten ist, dass durch die Inhomogenität der Differenzialgleichung die Amplitude der homogenen Lösung sich variabel ändert. Daher multipliziert man die homogene Lösung $\varphi(x)$ mit einer ortsabhängigen Amplitude $c(x)$. Diese unbekannte Amplitude wird durch Einsetzen des Ansatzes in die Differenzialgleichungen bestimmt.

Bemerkung: Es ist nicht unbedingt notwendig, diese fertige Lösungsformel auswendig zu lernen. In den Anwendungen werden oftmals die Lösungsverfahren auf den konkret gegebenen Fall angewendet:

- (1) Lösung der homogenen Differenzialgleichung $y'(x) = h(x) y(x)$ durch Trennung der Variablen bzw.

$$y(x) = c e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}.$$

- (2) Variation der Konstanten mit dem Ansatz

$$y(x) = c(x) \cdot e^{\int_{x_0}^x h(\tilde{x}) d\tilde{x}}.$$

□

Interpretation der Lösungsformel: Durch Ausmultiplizieren der Lösungsformel (I) erhält man die Darstellung

$$y(x) = \underbrace{y_0 \varphi(x)} + \underbrace{\varphi(x) \int_{x_0}^x \frac{f(\tilde{x})}{\varphi(\tilde{x})} d\tilde{x}}.$$

Die allgemeine Lösung der inhomogenen Differenzialgleichung lässt sich schreiben als Summe der **allgemeinen Lösung der homogenen Differenzialgleichung** und **einer speziellen Lösung der inhomogenen Differenzialgleichung**. Man nennt eine spezielle Lösung der inhomogenen Differenzialgleichung auch **partikuläre Lösung**.

Beweis des Satzes über lineare Differenzialgleichungen 1. Ordnung: Dass die angegebene Formel eine Lösung des Anfangswertproblems liefert, d.h. die Differenzialgleichung und die Anfangsbedingung $y(x_0) = y_0$ erfüllt, rechnet man nach, indem man die Funktion in die Differenzialgleichung einsetzt.

Dass $y(x)$ die einzige Lösung der Differenzialgleichung mit Anfangsbedingung ist, sieht man folgendermaßen leicht ein: Sei $y_2(x)$ ebenfalls eine Lösung. Dann gilt für die Differenz

$$d(x) = y(x) - y_2(x)$$

durch Differenzieren

$$\begin{aligned} d'(x) = y'(x) - y_2'(x) &= h(x) y(x) + f(x) - [h(x) y_2(x) + f(x)] \\ &= h(x) (y(x) - y_2(x)) = h(x) \cdot d(x) \end{aligned}$$

mit $d(x_0) = y(x_0) - y_2(x_0) = 0$. Also ist $d(x)$ Lösung der homogenen Differenzialgleichung

$$d'(x) = h(x) d(x) \quad \text{mit } d(x_0) = 0. \quad (*)$$

Wir zeigen nun, dass die Differenz $d(x) = 0$ für alle $x \in I$. Damit folgt dann, dass $y(x) = y_2(x)$ für alle $x \in I$ und $y(x)$ die einzige Lösung ist. Dazu definieren wir

$$u(x) := d(x) e^{-\int_{x_0}^x h(\tilde{x}) d\tilde{x}}.$$

Mit der Produkt- und Kettenregel gilt für die Ableitung

$$\begin{aligned} u'(x) &= d'(x) e^{-\int_{x_0}^x h(\tilde{x}) d\tilde{x}} + d(x) e^{-\int_{x_0}^x h(\tilde{x}) d\tilde{x}} (-h(x)) \\ &= \underbrace{[d'(x) - h(x) d(x)]}_{=0} \cdot e^{-\int_{x_0}^x h(\tilde{x}) d\tilde{x}} = 0, \end{aligned}$$

da $d(x)$ Lösung der homogenen Differenzialgleichung (*).

Also ist $u(x)$ eine konstante Funktion. Die Konstante kann man z.B. durch Auswerten von $u(x)$ an der Stelle x_0 bestimmen: Mit $u(x_0) = \text{const} = d(x_0) = 0 \Rightarrow u(x) = 0$. Damit ist d die Nullfunktion: $d(x) = 0$, so dass $y(x) = y_2(x)$ für alle $x \in I$. $\square$

Beispiel 12.6. Gegeben ist die Differenzialgleichung

$$y'(x) = 2xy(x) + x^3 \quad \text{mit } y(0) = y_0$$

Gesucht ist die Lösung $y(x)$, welche die Anfangsbedingung erfüllt.

Wir lösen diese Differenzialgleichung in drei Schritten. Zuerst bestimmen wir die allgemeine Lösung der homogenen DG, berechnen dann eine spezielle Lösung der inhomogenen DG und setzen anschließend beide Teile zusammen:

- (1) Lösung der **homogenen** Differenzialgleichung $y'(x) = 2xy(x)$ durch Formel (H)

$$\varphi(x) = e^{\int_0^x 2\bar{x} d\bar{x}} = e^{x^2}.$$

- (2) Lösung der **inhomogenen** Differenzialgleichung $y'(x) = 2xy(x) + x^3$ mit Formel (I):

$$y(x) = e^{x^2} \left(y_0 + \int_0^x \frac{t^3}{e^{t^2}} dt \right).$$

Die Berechnung des unbestimmten Integrals erfolgt zunächst durch eine Substitution ($\xi = t^2$, $d\xi = 2t dt$)

$$\int t^3 e^{-t^2} dt = \int t^3 e^{-\xi} \frac{d\xi}{2t} = \frac{1}{2} \int \xi e^{-\xi} d\xi$$

und anschließender partieller Integration

$$\frac{1}{2} \int \xi e^{-\xi} d\xi = \frac{1}{2} \left[-\xi e^{-\xi} + \int e^{-\xi} d\xi \right] = \frac{1}{2} [-\xi e^{-\xi} - e^{-\xi}] + C.$$

Durch Rücksubstitution ($\xi = t^2$) und Einsetzen der Grenzen gilt für das bestimmte Integral

$$\int_0^x t^3 e^{-t^2} dt = \frac{1}{2} \left[-t^2 e^{-t^2} - e^{-t^2} \right]_0^x = \frac{1}{2} - \frac{1}{2} e^{-x^2} (x^2 + 1).$$

- (3) Die **allgemeine Lösung** der Differenzialgleichung lautet damit

$$y(x) = y_0 e^{x^2} + e^{x^2} \left[\frac{1}{2} - \frac{1}{2} e^{-x^2} (x^2 + 1) \right] = y_0 e^{x^2} + \frac{1}{2} e^{x^2} - \frac{1}{2} (x^2 + 1) e^{x^2}. \quad \square$$

Anwendungsbeispiel 12.7 (RL-Kreis).

Mit der Lösungsformel (I) für inhomogene lineare Differenzialgleichung behandeln wir das in Beispiel 12.1 gestellte Problem des RL-Kreises:

$$\dot{I}(t) = -\frac{R}{L} I(t) + \frac{1}{L} U_B \quad \text{mit} \quad I(0) = 0.$$

- (1) Die Lösung der **homogenen** Differenzialgleichung $\dot{I}(t) = -\frac{R}{L} I(t)$ ist nach Formel (H)

$$I_h(t) = c e^{-\frac{R}{L} t}.$$

- (2) Damit erhält man mit $I_0 = 0$ die Lösung der **inhomogenen** Differenzialgleichung über Formel (I)

$$\begin{aligned} I(t) &= e^{-\frac{R}{L} t} \left(I_0 + \int_{t_0}^t \frac{U_B}{L} \frac{1}{e^{-\frac{R}{L} \tau}} d\tau \right) \\ &= e^{-\frac{R}{L} t} \frac{U_B}{L} \int_0^t e^{\frac{R}{L} \tau} d\tau = e^{-\frac{R}{L} t} \frac{U_B}{L} \frac{L}{R} e^{\frac{R}{L} \tau} \Big|_0^t \end{aligned}$$

$$\Rightarrow I(t) = e^{-\frac{R}{L} t} \frac{U_B}{R} \left(e^{\frac{R}{L} t} - 1 \right) = \frac{U_B}{R} \left(1 - e^{-\frac{R}{L} t} \right).$$

Dies ist genau der Stromverlauf, der in Beispiel 12.1 diskutiert wurde. $\square$

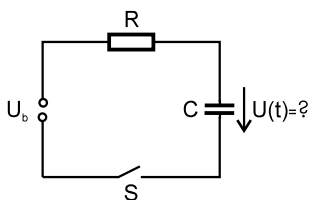
Anwendungsbeispiel 12.8 (Modellierung des RC-Kreis).

Abb. 12.6. RC-Kreis

Ein Widerstand R , ein Kondensator mit Kapazität C und eine Spannungsquelle $U_b(t)$ sind mit einem Schalter in Reihe geschaltet. Der Schalter ist zunächst offen und wird zur Zeit $t = 0$ geschlossen. Wie verhält sich die Spannung $U(t)$ am Kondensator als Funktion der Zeit für $t > 0$?

Nach dem Maschensatz gilt für die Spannungen

$$U_R(t) + U(t) = U_b(t). \quad (*)$$

Am Ohmschen Widerstand ist $U_R(t) = R \cdot I(t)$. An der Kapazität gilt

$$U(t) = \frac{1}{C} Q(t) = \frac{1}{C} \int I(t) dt \Rightarrow \dot{U}(t) = \frac{1}{C} I(t) \Rightarrow I(t) = C \cdot \dot{U}(t),$$

womit

$$U_R(t) = R \cdot I(t) = RC \cdot \dot{U}(t).$$

Eingesetzt in die Gleichung (*) folgt

$$RC \dot{U}(t) + U(t) = U_b(t) \quad \text{mit } U(0) = 0$$

$$\hookrightarrow \dot{U}(t) = -\frac{1}{RC} U(t) + \frac{1}{RC} U_b(t) \quad \text{mit } U(0) = 0.$$

□

Anwendungsbeispiel 12.9 (RC-Kreis bei Gleichspannung).

Für eine **konstante Batteriespannung** $U_b(t) = \hat{U}_b$ erhält man analog dem Vorgehen von Beispiel 12.7 als Lösung

$$U(t) = \hat{U}_b \left(1 - e^{-\frac{1}{RC} t}\right).$$

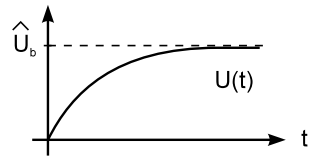


Abb. 12.7. Ladekurve eines Kondensators

Die Spannung am Kondensator $U(t)$ und damit die Ladung $Q(t) = C \cdot U(t)$ wächst mit der Zeit asymptotisch auf $\hat{U}_b$ bzw. $Q_0 = C \cdot \hat{U}_b$ an. □

Beispiel 12.10 (Musterbeispiel/Anwendungsbeispiel).

Ist die Eingangsspannung in Abb. 12.6 eine **Wechselspannung**

$$U_b(t) = \hat{U}_b \sin(\omega t)$$

mit dem Scheitelwert $\hat{U}_b$ und der Frequenz ω , so erhält die Differenzialgleichung die Form

$$\dot{U}(t) = -\frac{1}{RC} U(t) + \frac{1}{RC} \hat{U}_b \sin(\omega t) \quad \text{mit } U(0) = 0.$$

- (1) Gemäß dem Formel (H) lösen wir zunächst die homogene Differenzialgleichung

$$\dot{U}(t) = -\frac{1}{RC} U(t) \quad \text{durch} \quad U_h(t) = e^{\int_0^t -\frac{1}{RC} dt} = e^{-\frac{1}{RC} t}.$$

- (2) Anschließend verwenden wir die Lösungsformel (I)

$$U(t) = e^{-\frac{1}{RC} t} \left(U(0) + \frac{\hat{U}_b}{RC} \int_0^t \frac{\sin(\omega\tau)}{e^{-\frac{1}{RC} \tau}} d\tau \right).$$

Zur Berechnung des Integrals integrieren wir zweimal partiell

$$\begin{aligned} \int_0^t \sin(\omega\tau) e^{\frac{1}{RC} \tau} d\tau &= \\ &= \left[\sin(\omega\tau) \cdot RC \cdot e^{\frac{1}{RC} \tau} \right]_0^t - \int_0^t \omega \cos(\omega\tau) RC e^{\frac{1}{RC} \tau} d\tau \\ &= RC \sin(\omega t) e^{\frac{1}{RC} t} - \omega RC \left\{ \left[\cos(\omega\tau) e^{\frac{1}{RC} \tau} RC \right]_0^t \right. \\ &\quad \left. + \int_0^t \omega \sin(\omega\tau) e^{\frac{1}{RC} \tau} RC d\tau \right\} \\ &= RC \sin(\omega t) e^{\frac{1}{RC} t} - \omega (RC)^2 \left(\cos(\omega t) e^{\frac{1}{RC} t} - 1 \right) \\ &\quad - \omega^2 (RC)^2 \int_0^t \sin(\omega\tau) e^{\frac{1}{RC} \tau} d\tau. \end{aligned}$$

Da das verbleibende Integral auf der rechten Seite mit dem zu berechnenden Integral übereinstimmt, addieren wir auf beiden Seiten der Gleichung den Term $\omega^2 (RC)^2 \int_0^t \sin(\omega\tau) e^{\frac{1}{RC} \tau} d\tau$ und dividieren das Ergebnis durch $1 + (\omega RC)^2$:

$$\begin{aligned} \int_0^t \sin(\omega\tau) e^{\frac{1}{RC} \tau} d\tau &= \frac{1}{1 + (\omega RC)^2} \left\{ RC \sin(\omega t) e^{\frac{1}{RC} t} \right. \\ &\quad \left. - \omega (RC)^2 \cos(\omega t) e^{\frac{1}{RC} t} + \omega (RC)^2 \right\}. \end{aligned}$$

- (3) Setzen wir dieses Ergebnis in die Lösungsformel ein, erhalten wir mit der Anfangsbedingung $U(0) = 0$

$$\begin{aligned} U(t) &= e^{-\frac{1}{RC} t} \frac{\hat{U}_b}{1 + (\omega RC)^2} \left\{ \sin(\omega t) e^{\frac{1}{RC} t} - \omega RC \cos(\omega t) e^{\frac{1}{RC} t} + \omega RC \right\} \\ &= \frac{\hat{U}_b}{1 + (\omega RC)^2} \left\{ \sin(\omega t) - \omega RC \cos(\omega t) + \omega RC e^{-\frac{1}{RC} t} \right\}. \end{aligned}$$

Interpretation: Die Lösung $U(t)$ setzt sich zusammen aus einem exponentiell abklingenden und einem periodischen Term.

$$U(t) = \underbrace{\frac{\hat{U}_b \omega RC}{1 + (\omega RC)^2} e^{-\frac{1}{RC} t}}_{\text{Einschwingverhalten}} + \underbrace{\frac{\hat{U}_b}{1 + (\omega RC)^2} \{\sin(\omega t) - \omega RC \cos(\omega t)\}}_{\text{Verhalten für große } t}.$$

Der exponentiell abklingende Term spiegelt den **Einschwingvorgang** wider. Das **Langzeitverhalten** der Lösung ist jedoch durch den periodischen Anteil bestimmt. In Abb. 12.8 ist die Lösung für die Parameter $RC = 10$, $\hat{U}_b = 1$ und $\omega = 1$ gezeichnet. Daran erkennt man gut den Einschwingvorgang sowie das Langzeitverhalten.

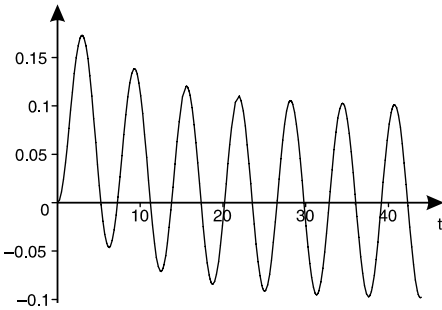


Abb. 12.8. Lösung der DG

Physikalische Interpretation: Wir stellen diesen periodischen Anteil der Lösung als reine harmonische Schwingung mit Amplitude A , Frequenz ω und Phase φ dar:

$$\sin(\omega t) - \omega RC \cos(\omega t) = A \sin(\omega t + \varphi) \quad (*)$$

Um A und φ zu bestimmen benötigen wir zwei Gleichungen. Dazu wenden wir das Additionstheorem für die rechte Seite an

$$A \sin(\alpha + \beta) = A \sin(\alpha) \cos(\beta) + A \cos(\alpha) \sin(\beta).$$

Mit $\alpha = \omega t$ und $\beta = \varphi$ erhalten wir durch Koeffizientenvergleich

$$-\omega RC = A \sin \varphi . \quad (1)$$

$$1 = A \cos \varphi . \quad (2)$$

Dividiert man Gleichung (1) durch (2), folgt

$$\boxed{\tan \varphi = -RC \omega} \Rightarrow \varphi . \quad (3)$$

Addieren wir die quadrierten Gleichungen (1) und (2) folgt

$$1 + (RC\omega)^2 = A^2 \cos^2 \varphi + A^2 \sin^2 \varphi = A^2 (\cos^2 \varphi + \sin^2 \varphi) = A^2$$

$$\Rightarrow A = \sqrt{1 + (RC\omega)^2}. \quad (4)$$

Aus (3) und (4) folgt insgesamt für die Lösung $U(t)$

$$U(t) = \text{Einschwingvorgang} + \frac{\hat{U}_b}{\sqrt{1 + (\omega RC)^2}} \sin(\omega t + \varphi). \quad \square$$

➤ 12.1.4 Lineare Differenzialgleichungen mit konstantem Koeffizient

Durch die beiden Formeln (H) und (I) ist jede lineare Differenzialgleichung 1. Ordnung im Prinzip lösbar. Die Auswertung der Integrale kann aber sehr aufwändig werden. In vielen Fällen braucht man nicht auf diese Integraldarstellung der Lösung zurückgreifen, sondern ermittelt eine *partikuläre* Lösung durch einen speziellen Lösungsansatz. Dies ist insbesondere für die Differenzialgleichungen mit **konstanten** Koeffizienten der Fall. Beim konkreten Problem wird entsprechend dem Typ der *Inhomogenität* ein Ansatz für die partikuläre Lösung mit freien Parametern gewählt. Diese Parameter bestimmt man anschließend durch Einsetzen des Ansatzes in die Differenzialgleichung. Bei dieser Methode des Ratens der partikulären Lösung ist zu beachten, dass sie nur in einfachen Fällen zum Ziele führt.

⊗ Homogene Lösung

Zum Lösen der *homogenen* linearen Differenzialgleichungen mit konstantem Koeffizient α ,

$$y'(x) = \alpha \cdot y(x),$$

wählen wir den Ansatz

$$y_h(x) = c e^{\lambda x}$$

und setzen diese Funktion in die Differenzialgleichung ein:

$$c \lambda e^{\lambda x} = \alpha \cdot c e^{\lambda x} \hookrightarrow \lambda = \alpha.$$

Die allgemeine Lösung der homogenen Differenzialgleichung lautet somit

$$y_h(x) = c e^{\alpha x}.$$

⊙ **Partikuläre Lösung**

Wir betrachten die *inhomogene* lineare Differenzialgleichung

$$y'(x) = \alpha y(x) + f(x)$$

mit $\alpha \neq 0$. Je nach Inhomogenität (*Störfunktion*) f lässt sich eine *partikuläre Lösung* durch einen einfachen Ansatz finden. Für häufig auftretende Störfunktionen geben wir die Ansatzfunktionen in Tab. 12.1 an.

Tabelle 12.1: Partikuläre Lösungen für spezielle Inhomogenitäten.

Inhomogenität	Ansatzfunktion	Parameter
$f(x) = \sum_{i=0}^n a_i x^i$ (Polynom vom Grad n)	$y_p(x) = \sum_{i=0}^n A_i x^i$	$A_0, \dots, A_n$
$f(x) = a \cdot \sin(\omega x)$ (Sinusanregung)	$y_p(x) = A \sin(\omega x) + B \cos(\omega x)$	A, B
$f(x) = a \cdot \cos(\omega x)$ (Kosinusanregung)	$y_p(x) = A \sin(\omega x) + B \cos(\omega x)$	A, B
$f(x) = a e^{\mu x} \ (\mu \neq \alpha)$ (Exponentialfunktion)	$y_p(x) = A e^{\mu x}$	A

Bemerkung: Besteht die Inhomogenität $f(x)$ aus einer Summe von mehreren Störgliedern, wählt man für jedes Störglied einzeln die entsprechende Ansatzfunktion und bestimmt durch Einsetzen der Ansatzfunktion in die Differenzialgleichung die zugehörigen Parameter. Zum Schluss addiert man alle partikulären Lösungen auf und erhält eine spezielle Lösung des gestellten, inhomogenen Problems.

Beispiel 12.11 (Musterbeispiel).

Gesucht ist die allgemeine Lösung der linearen Differenzialgleichung

$$y'(x) = 4 \cdot y(x) + x^3 \ .$$

- (1) Die **homogene** Differenzialgleichung $y'(x) = 4 \cdot y(x)$ lösen wir mit dem Ansatz

$$y_h(x) = c e^{\lambda x} \ .$$

In die Differenzialgleichung eingesetzt, ergibt sich

$$c \lambda e^{\lambda x} = 4 c e^{\lambda x} \hookrightarrow \lambda = 4 \quad \Rightarrow \quad y_h(x) = c e^{4x} \ .$$

- (2) Eine partikuläre Lösung der
- inhomogenen**
- Differenzialgleichung

$$y'(x) = 4 \cdot y(x) + x^3$$

bestimmen wir nach Tabelle 12.1 durch den Ansatz

$$y_p(x) = a x^3 + b x^2 + c x + d,$$

da die Inhomogenität $f(x) = x^3$ ein Polynom vom Grad 3 ist. Die Ansatzfunktion in die inhomogene Differenzialgleichung eingesetzt, ergibt

$$\begin{aligned} 3 a x^2 + 2 b x + c &\stackrel{!}{=} 4 a x^3 + 4 b x^2 + 4 c x + 4 d + x^3 \\ &= (4 a + 1) x^3 + 4 b x^2 + 4 c x + 4 d. \end{aligned}$$

Ein Koeffizientenvergleich liefert für absteigende Potenzen in x

$$\begin{array}{lll} x^3: & 4 a + 1 = 0 & \hookrightarrow a = -\frac{1}{4} \\ x^2: & 4 b = 3 a = -\frac{3}{4} & \hookrightarrow b = -\frac{3}{16} \\ x^1: & 4 c = 2 b = -\frac{3}{8} & \hookrightarrow c = -\frac{3}{32} \\ x^0: & 4 d = c = -\frac{3}{32} & \hookrightarrow d = -\frac{3}{128}. \end{array}$$

Eine partikuläre Lösung ist also

$$y_p(x) = -\frac{1}{4} x^3 - \frac{3}{16} x^2 - \frac{3}{32} x - \frac{3}{128}.$$

- (3) Somit ist die
- allgemeine Lösung**
- der inhomogenen Differenzialgleichung

$$y(x) = y_h(x) + y_p(x) = c e^{4x} - \frac{1}{4} x^3 - \frac{3}{16} x^2 - \frac{3}{32} x - \frac{3}{128}. \quad \square$$

Beispiel 12.12 (Musterbeispiel, mit MAPLE-Worksheet). Die Differenzialgleichung aus Beispiel 12.10 lautet für eine Wechselspannung $U_b(t) = \hat{U}_b \sin(\omega t)$ mit $R \cdot C = 1$

$$\dot{U}(t) = -U(t) + \hat{U}_b \sin(\omega t); \quad U(0) = 0.$$

- (1) Die
- homogene**
- Differenzialgleichung
- $\dot{U}(t) = -U(t)$
- lösen wir durch den Ansatz
- $U_h(t) = c e^{\lambda t}$
- . In die Differenzialgleichung eingesetzt, folgt
- $\lambda = -1$
- .

$$\Rightarrow U_h(t) = c e^{-t}.$$

- (2) Zum Lösen der
- inhomogenen**
- Differenzialgleichung wählen wir für eine partikuläre Lösung gemäß Tabelle 12.1 den Ansatz:

$$U_p(t) = A \sin(\omega t) + B \cos(\omega t) \quad (*)$$

und setzen $U_p(t)$ in die inhomogene Differenzialgleichung ein, um die noch unbekannten Konstanten A und B zu bestimmen:

$$A \omega \cos(\omega t) - B \omega \sin(\omega t) = -A \sin(\omega t) - B \cos(\omega t) + \hat{U}_b \sin(\omega t).$$

Wir ordnen die rechte Seite der Gleichung nach Vielfachen von $\cos(\omega t)$ und $\sin(\omega t)$

$$(-B \omega) \sin(\omega t) + (A \omega) \cos(\omega t) = (\hat{U}_b - A) \sin(\omega t) - B \cos(\omega t).$$

Die Gleichung ist für alle t nur dann erfüllt, wenn die Koeffizienten sowohl der Sinus- als auch der Kosinusfunktionen auf beiden Seiten der Gleichung übereinstimmen. Ein Koeffizientenvergleich in $\cos(\omega t)$ und $\sin(\omega t)$ auf beiden Seiten dieser Gleichung führt zu dem linearen Gleichungssystem

$$\cos(\omega t): \quad A \omega \quad = \quad -B \quad (1)$$

$$\sin(\omega t): \quad -B \omega \quad = \quad \hat{U}_b - A \quad (2)$$

aus welchem A und B zu bestimmen sind. Setzt man (1) in (2) ein, gilt

$$-B \omega = \hat{U}_b - \frac{1}{\omega} B \Rightarrow B = \frac{\omega}{\omega^2 + 1} \hat{U}_b.$$

Mit (1) folgt dann

$$A = \frac{\hat{U}_b}{\omega^2 + 1}.$$

Somit ist eine partikuläre Lösung nach (*)

$$U_p(t) = \frac{\hat{U}_b}{\omega^2 + 1} [\sin(\omega t) - \omega \cos(\omega t)].$$

(3) Die **allgemeine Lösung** lautet

$$U(t) = U_h(t) + U_p(t) = c e^{-t} + \frac{\hat{U}_b}{\omega^2 + 1} [\sin(\omega t) - \omega \cos(\omega t)].$$

(4) Zum Schluss bestimmen wir die Konstante c durch die Anfangsbedingung $U(0) = 0$:

$$0 = c + \frac{\hat{U}_b}{\omega^2 + 1} [0 - \omega] \Rightarrow c = \frac{\omega \hat{U}_b}{\omega^2 + 1}.$$

Damit ist die Lösung des Problems gegeben durch

$$U(t) = \frac{\hat{U}_b}{\omega^2 + 1} [\omega e^{-t} + \sin(\omega t) - \omega \cos(\omega t)],$$

welche für $RC = 1$ mit der Endformel aus Beispiel 12.10 übereinstimmt. $\square$

➤ 12.1.5 Nichtlineare Differenzialgleichungen 1. Ordnung

Die Methode Trennung der Variablen wird nicht nur beim Lösen von homogenen linearen Differenzialgleichungen eingesetzt, sondern einfache *nichtlineare* Differenzialgleichungen lassen sich ebenfalls mit dieser Methode lösen.

Hinweis: Auf der CD-Rom befindet sich ein erweiterter Abschnitt, in dem neben der Trennung der Variablen auch **Substitutionsmethoden** eingeführt werden, um nichtlineare Differenzialgleichungen zu lösen.

⊗ Differenzialgleichung mit trennbaren Variablen

Wir betrachten die Differenzialgleichung

$$y'(x) = f(x) g(y)$$

mit einer gegebenen stetigen Funktion g . Sofern g nicht eine lineare Funktion darstellt, ist die Differenzialgleichung nichtlinear! Sie hat aber die gleiche Bauart, wie eine homogene *lineare* Differenzialgleichung 1. Ordnung und lässt sich durch **Trennung der Variablen** lösen. Man bezeichnet diesen Typ daher auch als Differenzialgleichung *mit trennbaren Variablen*. Zum Lösen wird die Differenzialgleichung wie folgt umgestellt

$$\begin{aligned} \frac{dy}{dx} &= f(x) \cdot g(y) \quad | : g(y) \cdot dx \\ \hookrightarrow \frac{dy}{g(y)} &= f(x) \cdot dx. \end{aligned}$$

Die linke Seite der Gleichung enthält nur noch die Variable y und die rechte Seite nur noch die Variable x . Anschließende Integration liefert

$$G(y) = \int \frac{dy}{g(y)} = \int f(x) dx.$$

Die Stammfunktion des linken Integrals $G(y)$ wird anschließend nach y aufgelöst, was in vielen Fällen möglich ist.

Beispiel 12.13. $y'(x) = e^{y(x)} \cos x$ mit $y(0) = y_0$.

Diese Differenzialgleichung lässt sich durch Trennung der Variablen lösen:

$$\begin{aligned} \frac{dy}{dx} &= e^y \cos x \quad | : e^y \cdot dx \\ \frac{dy}{e^y} &= \cos x dx. \end{aligned}$$

Da hier ein Anfangswertproblem vorliegt, wählen wir auf beiden Seiten der Gleichung das bestimmte Integral. Die Integration über y erfolgt von y_0 ab und die Integration über x von x_0 ab:

$$\int_{y_0}^y e^{-\tilde{y}} d\tilde{y} = \int_0^x \cos \tilde{x} d\tilde{x}.$$

Wir führen die Integrationen auf beiden Seiten aus und lösen anschließend nach y auf

$$-e^{-\tilde{y}} \Big|_{y_0}^y = \sin \tilde{x} \Big|_0^x \quad \hookrightarrow \quad -e^{-y} + e^{-y_0} = \sin x$$

$$\hookrightarrow e^{-y} = e^{-y_0} - \sin x \quad \Rightarrow \quad y(x) = -\ln(e^{-y_0} - \sin x).$$

Bemerkung: Wählt man statt dem bestimmten Integral die unbestimmte Form $\int e^{-y} dy = \int \cos x dx$, ist bei der Integration eine Integrationskonstante C zu berücksichtigen. Diese Konstante C wird zum Schluss durch die Anfangsbedingung $y(0) = y_0$ festgelegt. $\square$

Anwendungsbeispiel 12.14 (Freier Fall mit Luftwiderstand).

Wir untersuchen die Sinkgeschwindigkeit $v(t)$ eines Körpers der Masse m unter Berücksichtigung der Luftreibung. Die am Körper angreifenden Kräfte sind

- (1) die Schwerkraft mg
- (2) der Luftwiderstand $-kv^2$, wenn eine quadratische Abhängigkeit der Reibungskraft von der Geschwindigkeit angenommen wird.

Dabei ist g die Erdbeschleunigung und k der Reibungskoeffizient.

Nach dem Newtonschen Bewegungsgesetz ist die Beschleunigungskraft $m \frac{dv}{dt}$ gleich der Summe aller angreifenden Kräfte

$$\Rightarrow m \frac{dv(t)}{dt} = mg - kv^2(t).$$

Unter der Annahme, dass der freie Fall aus der Ruhe erfolgt, gilt zusätzlich die Anfangsbedingung $v(0) = 0$.

Diese *nichtlineare* Differenzialgleichung 1. Ordnung lösen wir durch Trennung der Variablen:

$$\frac{dv}{dt} = g - \frac{k}{m} v^2 \quad \Bigg| : \left(g - \frac{k}{m} v^2 \right) \cdot dt$$

$$\hookrightarrow \frac{dv}{g - \frac{k}{m} v^2} = dt.$$

Anschließend integrieren wir

$$\int_0^t d\tilde{t} = \int_0^v \frac{d\tilde{v}}{g - \frac{k}{m} \tilde{v}^2} = \frac{1}{g} \int_0^v \frac{d\tilde{v}}{1 - \frac{k}{m g} \tilde{v}^2}.$$

Das linke Integral ergibt sich zu $\int_0^t d\tilde{t} = t$. Zur Berechnung des rechten Integrals substituieren wir $\xi = \sqrt{\frac{k}{m g}} v$. Damit ist $d\xi = \sqrt{\frac{k}{m g}} dv$ und es gilt

$$\int \frac{d\tilde{v}}{1 - \frac{k}{m g} \tilde{v}^2} = \sqrt{\frac{m g}{k}} \int \frac{d\xi}{1 - \xi^2} = \sqrt{\frac{m g}{k}} \operatorname{artanh}(\xi) + C.$$

Nach der Rücksubstitution erhalten wir für das bestimmte Integral den Ausdruck

$$\frac{1}{g} \int_0^v \frac{d\tilde{v}}{1 - \frac{k}{m g} \tilde{v}^2} = \frac{1}{g} \sqrt{\frac{m g}{k}} \operatorname{artanh} \left(\sqrt{\frac{k}{m g}} \tilde{v} \right) \Big|_0^v = \sqrt{\frac{m}{k g}} \operatorname{artanh} \left(\sqrt{\frac{k}{m g}} v \right).$$

Insgesamt folgt also

$$t = \sqrt{\frac{m}{k g}} \operatorname{artanh} \left(\sqrt{\frac{k}{m g}} v \right).$$

Diese Gleichung lösen wir durch Anwenden der Funktion $\tanh$ nach v auf:

$$\sqrt{\frac{k g}{m}} t = \operatorname{artanh} \left(\sqrt{\frac{k}{m g}} v \right) \hookrightarrow \tanh \left(\sqrt{\frac{k g}{m}} t \right) = \sqrt{\frac{k}{m g}} v$$

$$\Rightarrow v(t) = \sqrt{\frac{m g}{k}} \tanh \left(\sqrt{\frac{k g}{m}} t \right).$$

Dieses Geschwindigkeitsgesetz für die Fallgeschwindigkeit $v(t)$ eines Körpers der Masse m mit Luftwiderstand ist in Abb. 12.9 skizziert:

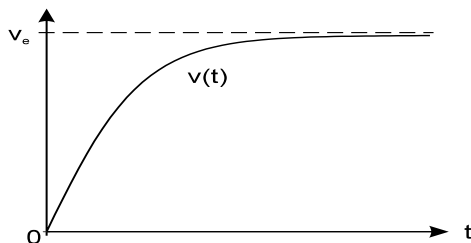


Abb. 12.9. Fallgeschwindigkeit mit Luftwiderstand

Für $t \rightarrow \infty$ wird die Endgeschwindigkeit $v_E = \lim_{t \rightarrow \infty} v(t) = \sqrt{\frac{m g}{k}}$ erreicht, da $\lim_{x \rightarrow \infty} \tanh(x) = 1$. Der Körper fällt dann kräftefrei mit konstanter Geschwindigkeit, da sich Reibungskraft und Gewichtskraft gegenseitig aufheben.

□

➤ 12.1.6 Numerisches Lösen von DG 1. Ordnung

➤ Richtungsfelder, mit [MAPLE-Worksheet](#)

Charakteristisch für die diskutierten Differenzialgleichungen ist, dass die Ableitung der Funktion

$$y'(x) = f(x, y(x))$$

für jeden Punkt (x, y) der Ebene durch die rechte Seite der Differenzialgleichung $f(x, y(x))$ gegeben ist. Dies führt auf die Darstellung der Differenzialgleichung in Form eines *Richtungsfeldes*, in der in jedem Punkt der Ebene die Steigung der Funktion $y(x)$ (also $f(x, y(x))$) als Vektor aufgetragen wird. Wählen wir beispielsweise die Differenzialgleichung

$$y'(x) = -y(x) + 1,$$

so ist das Richtungsfeld gegeben durch

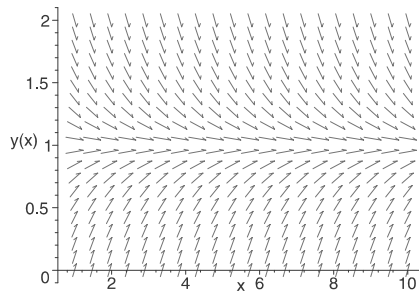


Abb. 12.10. Richtungsfeld der Differenzialgleichung $y'(x) = -y(x) + 1$

Hat man zur Differenzialgleichung noch einen Startwert (Anfangswert) $y(x_0)$ gegeben, dann kennt man den Punkt der Ebene $(x_0, y(x_0))$, in dem die Lösung beginnt. Durch diese zusätzliche Information kann man nun die Lösung konstruieren, indem man ausgehend vom Startpunkt $(x_0, y(x_0))$ über die Steigung der Funktion in diesem Punkt $f(x_0, y(x_0))$ zum nächsten Punkt an der Stelle $x_0 + dx$ kommt und damit y an der Stelle $x_0 + dx$ erhält. Dann ist sowohl $y(x_0 + dx)$ als auch die Steigung $y'(x_0 + dx) = f(x_0 + dx, y(x_0 + dx))$ bekannt und man konstruiert hieraus $y(x_0 + 2dx)$ usw.

In Abb. 12.11 ist das Richtungsfeld der Differenzialgleichung zusammen mit der Lösung zum Anfangswert $y(4) = 0$ eingezeichnet.

Dieses Vorgehen führt auf das Euler-Verfahren zum numerischen Lösen von Differenzialgleichungen erster Ordnung.

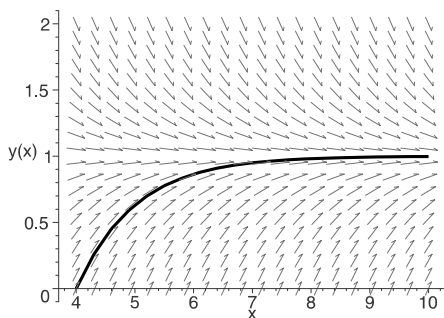


Abb. 12.11. Konstruktion der Lösung über das Richtungsfeld

② Streckenzugverfahren von Euler

Wir gehen von dem Anfangswertproblem (AWP)

$$y'(t) = f(t, y(t)) \quad \text{mit } y(t_0) = y_0 \quad (1)$$

aus und werden dieses AWP für Zeiten $t_0 \leq t \leq T$ numerisch lösen. Dazu zerlegen wir das Intervall $[t_0, T]$ in N Teilintervalle der Länge

$$h = dt = \frac{T - t_0}{N} .$$

Die Größen h und dt werden als **Schrittweite** bzw. **Zeitschritt** bezeichnet. Wir erhalten als Zwischenzeiten

$$t_j = t_0 + j \cdot dt \quad j = 0, \dots, N$$

und werden die Lösung nur zu diesen diskreten Zeiten $t_0, t_1, t_2, \dots, t_N$ berechnen: Ausgehend vom Startwert y_0 bestimmen wir der Reihe nach Näherungen $y_1, y_2, \dots, y_N$ für die Funktionswerte $y(t_1), y(t_2), \dots, y(t_N)$ der Lösung von (1). Man nennt dieses Vorgehen die **Diskretisierung** des AWP.

Für den Startwert (t_0, y_0) kennen wir nach Gl. (1) die exakte Steigung $\tan \alpha$ der Lösungsfunktion

$$y'(t_0) = \tan \alpha = f(t_0, y_0) .$$

Für eine kleine Schrittweite h wird die Funktion y im Intervall $[t_0, t_0 + h]$ durch ihre Tangente angenähert (Linearisierung) (siehe Abb. 12.12 a). Für den Funktionswert $y(t_1)$ gilt dann näherungsweise

$$y(t_1) = y(t_0 + h) \approx y(t_0) + y'(t_0) \cdot h .$$

Wir setzen

$$y_1 := y_0 + f(t_0, y_0) h .$$

Damit hat man den Funktionswert $y(t_1)$ zum Zeitpunkt t_1 durch y_1 angenähert. Ausgehend von diesem fehlerhaften Wert y_1 berechnet man mit (1)

die fehlerhafte Steigung $y'_1 = f(t_1, y_1)$. Nun benutzt man y_1 und y'_1 für die Berechnung eines Näherungswertes für den nächsten Zeitpunkt $t_2 = t_1 + h$:

$$y(t_2) = y(t_1 + h) \approx y(t_1) + y'(t_1) h. \quad (\text{Linearisierung})$$

Wir setzen

$$y_2 := y_1 + f(t_1, y_1) h.$$

y_2 ist eine Näherung für den exakten Wert der Lösung $y(t_2)$ zur Zeit t_2 .

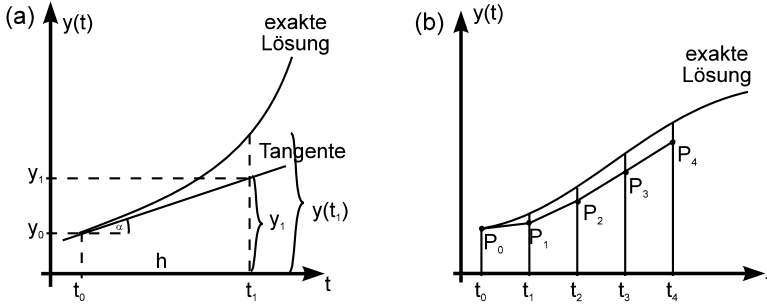


Abb. 12.12. a) 1. Schritt des Polygonzugverfahrens b) Polygonzugverfahren nach Euler

Das beschriebene Verfahren wiederholt man für den neuen Punkt $P_2 = (t_2, y_2)$, der in der Regel nicht auf der exakten Lösungskurve liegt. Allgemein kann man das Näherungsverfahren beschreiben durch:

$$\text{neuer Wert} = \text{alter Wert} + \text{Änderung der Lösung}$$

bzw.

$$y_{\text{neu}} = y_{\text{alt}} + f(t_{\text{neu}}, y_{\text{alt}}) \cdot h.$$

Man berechnet also, ausgehend vom Punkt $P_0 = (t_0, y_0)$ sukzessiv die Werte

$$y_{i+1} = y_i + h f(t_i, y_i) \quad i = 0, 1, 2, \dots, N - 1.$$

Die Lösungskurve setzt sich aus geradlinigen Strecken zusammen, so dass die Näherung in Form eines Streckenzugs vorliegt (siehe Abb. 12.12b). Dieses Verfahren heißt gemäß seiner geometrischen Bedeutung das **Polygonzugverfahren** bzw. nach seinem Erfinder auch das **Euler-Verfahren**.

Bemerkung: Für einfache Beispiele liefert dieses einfache Verfahren für genügend kleine Schrittweiten h ausreichend genaue Näherungswerte $y_1, y_2, \dots, y_N$ für die gesuchten Funktionswerte $y(t_1), y(t_2), \dots, y(t_N)$.

Beispiel 12.15 (Mit MAPLE-Worksheet). Das Anfangswertproblem

$$y'(t) = y(t) + e^t \quad \text{mit } y(0) = 1$$

besitzt die exakte Lösung

$$y(t) = (t+1)e^t.$$

Wir berechnen mit dem Euler-Verfahren Näherungslösungen dieser Differenzialgleichungen im Intervall $0 \leq t \leq 0.2$ für Schrittweiten $h = 0.05$ und $h = 0.025$ und vergleichen die Ergebnisse mit der exakten Lösung.

Wendet man das Euler-Verfahren auf die Differenzialgleichung an, lautet die Iterationsvorschrift

$$y_{i+1} = y_i + h (y_i + e^{t_i}).$$

Speziell für die Schrittweite $h = 0.05$ gilt mit dem Anfangswert $y(0) = 1$

$$\begin{aligned} y_0 &= 1 \\ y_1 &= y_0 + h (y_0 + e^{0.}) = 1.1 \\ y_2 &= y_1 + h (y_1 + e^{0.05}) = 1.207564 \\ y_3 &= y_2 + h (y_2 + e^{0.1}) = 1.323201 \\ y_4 &= y_3 + h (y_3 + e^{0.15}) = 1.447453 \end{aligned}$$

In der unten angegebenen Tabelle sind diese Näherungswerte für $h = 0.05$ und $h = 0.025$ zusammen mit den exakten Werten angegeben. Der Vergleich zeigt, dass die Näherungswerte bei kleineren Schrittweiten (3. Spalte) sich verbessern.

Ergebnisse:

t	$y(h = 0.05)$	$y(h = 0.025)$	y exakt
0.00	1.000 000	1.000 000	1.000 000
0.05	1.100 000	1.101 883	1.103 835
0.10	1.207 564	1.211 552	1.215 688
0.15	1.323 201	1.329 535	1.336 109
0.20	1.447 453	1.456 396	1.465 683

Hinweis: Auf der CD-Rom befindet sich ein zusätzliches Kapitel über das **numerische Lösen von Differenzialgleichungen** erster Ordnung mit einem qualitativen Vergleich der Euler-, Prädiktor-Korrektor- und Runge-Kutta-Verfahren. Neben der ausführlichen Beschreibung der jeweiligen Formeln und Algorithmen wird dort die MAPLE-Prozedur **DGsolve** vorgestellt, welche die drei Verfahren umsetzt und die Lösung graphisch ausgibt.

12.2 Lineare Differenzialgleichungssysteme

In vielen Anwendungen sind zeitlich veränderliche Größen $x_1(t)$, $x_2(t)$, $\dots$, $x_n(t)$ gekoppelt, so dass die Änderung einer Größe $\dot{x}_i(t)$ nicht nur von t und $x_i(t)$, sondern auch von den restlichen Größen und deren Ableitungen abhängt. Dies führt auf Differenzialgleichungssysteme. Wir behandeln in diesem Abschnitt *lineare Differenzialgleichungssysteme* erster Ordnung mit konstanten Koeffizienten. Für diese Systeme verwenden wir im Folgenden die Abkürzung *LDGS*.

Die homogenen LDGS werden mit Hilfe der Eigenvektoren und Eigenwerten der zum Problem zugeordneten Systemmatrix gelöst. Mit dem damit berechneten Fundamentalsystem des homogenen Problems bestimmen wir mit der Methode Variation der Konstanten eine Lösungsformel für das inhomogene Problem.

Hinweis: Auf der CD-Rom befinden sich das Kapitel über *inhomogene lineare Differenzialgleichungssysteme* sowie zusätzliche Anwendungsbeispiele, deren Lösung mit Unterstützung von MAPLE berechnet und in Form von Animationen visualisiert werden. Insbesondere bei der physikalischen Interpretation der mathematischen Begriffe Eigenwert und Eigenvektor als Eigenfrequenz und Eigenschwingung werden wir diese Animationen heranziehen.

12.2.1 Einführung

Einfachste Beispiele für LDGS erhält man bei elektrischen Filterschaltungen bestehend aus RCL-Elementen oder bei mechanischen Koppelschwingungen, bei denen mehrere Feder-Masse-Systeme gekoppelt werden. Zur Einführung betrachten wir ein gekoppeltes System aus Doppelpendel.

Anwendungsbeispiel 12.16 (Gekoppelte Pendel).

Zwei Fadenpendel der Länge l , an deren Ende jeweils eine Masse m_1 bzw. m_2 hängt, werden durch eine Feder mit Federkonstanten D gekoppelt (siehe Abb. 12.13). Die beiden Massen werden um den Winkel φ_1 bzw. φ_2 ausgelenkt. Berücksichtigt man, dass während der Bewegung auf die Massen eine Reibungskraft proportional zur Geschwindigkeit wirkt

$$F_R = -\gamma l \dot{\varphi}_i(t) \quad \text{mit Reibungskoeffizient } \gamma,$$

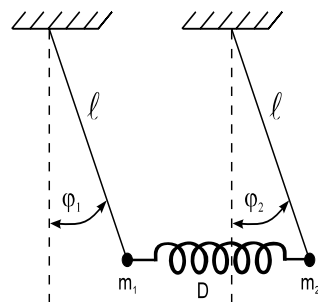


Abb. 12.13. Gekoppelte Pendel

so lauten die Bewegungsgleichungen für kleine Auslenkungen φ_1 und φ_2 :

$$\begin{aligned}\ddot{\varphi}_1(t) &= -\frac{g}{l} \varphi_1(t) - \frac{\gamma}{m_1} \dot{\varphi}_1(t) + \frac{D}{m_1} (\varphi_2(t) - \varphi_1(t)) \\ \ddot{\varphi}_2(t) &= -\frac{g}{l} \varphi_2(t) - \frac{\gamma}{m_2} \dot{\varphi}_2(t) + \frac{D}{m_2} (\varphi_1(t) - \varphi_2(t)),\end{aligned}\quad (*)$$

wenn $\varphi_1(t)$ und $\varphi_2(t)$ die Auslenkungen der Massen m_1 und m_2 zum Zeitpunkt t sind.

Dies ist ein LDGS **zweiter** Ordnung für die Winkelauslenkungen $\varphi_1(t)$ und $\varphi_2(t)$. Wir reduzieren dieses System von zwei Differenzialgleichungen 2. Ordnung auf ein System von vier Differenzialgleichungen 1. Ordnung. Dazu führen wir die zu $\varphi_1(t)$ und $\varphi_2(t)$ gehörenden Winkelgeschwindigkeiten $\dot{\varphi}_1(t)$ und $\dot{\varphi}_2(t)$ als zusätzliche Größen ein. Zur übersichtlicheren Darstellung setzen wir

$$\begin{aligned}y_1(t) &= \varphi_1(t) \\ y_2(t) &= \dot{\varphi}_1(t) \\ y_3(t) &= \varphi_2(t) \\ y_4(t) &= \dot{\varphi}_2(t).\end{aligned}$$

Differenzieren wir jede dieser vier unbekannten Funktionen $y_i(t)$, so gilt mit (*):

$$\begin{aligned}\dot{y}_1(t) &= \dot{\varphi}_1(t) = y_2(t) \\ \dot{y}_2(t) &= \ddot{\varphi}_1(t) = -\frac{g}{l} \varphi_1(t) - \frac{\gamma}{m_1} \dot{\varphi}_1(t) + \frac{D}{m_1} (\varphi_2(t) - \varphi_1(t)) \\ &= -\frac{g}{l} y_1(t) - \frac{\gamma}{m_1} y_2(t) + \frac{D}{m_1} (y_3(t) - y_1(t)) \\ \dot{y}_3(t) &= \dot{\varphi}_2(t) = y_4(t) \\ \dot{y}_4(t) &= \ddot{\varphi}_2(t) = -\frac{g}{l} \varphi_2(t) - \frac{\gamma}{m_2} \dot{\varphi}_2(t) + \frac{D}{m_2} (\varphi_1(t) - \varphi_2(t)) \\ &= -\frac{g}{l} y_3(t) - \frac{\gamma}{m_2} y_4(t) + \frac{D}{m_2} (y_1(t) - y_3(t)).\end{aligned}$$

Definieren wir die Vektorfunktion $\vec{y}(t) := \begin{pmatrix} y_1(t) \\ y_2(t) \\ y_3(t) \\ y_4(t) \end{pmatrix}$ und $\vec{y}'(t) := \begin{pmatrix} \dot{y}_1(t) \\ \dot{y}_2(t) \\ \dot{y}_3(t) \\ \dot{y}_4(t) \end{pmatrix}$

als die Ableitung, so lässt sich obiges LDGS in Vektornotation schreiben als

$$\vec{y}'(t) = \begin{pmatrix} y_2(t) \\ \left(-\frac{g}{l} - \frac{D}{m_1}\right) y_1(t) - \frac{\gamma}{m_1} y_2(t) + \frac{D}{m_1} y_3(t) & y_4(t) \\ \frac{D}{m_2} y_1(t) & \left(-\frac{g}{l} - \frac{D}{m_2}\right) y_3(t) - \frac{\gamma}{m_2} y_4(t) \end{pmatrix}$$

$$= \begin{pmatrix} 0 & 1 & 0 & 0 \\ -\frac{g}{l} - \frac{D}{m_1} & -\frac{\gamma}{m_1} & \frac{D}{m_1} & 0 \\ 0 & 0 & 0 & 1 \\ \frac{D}{m_2} & 0 & -\frac{g}{l} - \frac{D}{m_2} & -\frac{\gamma}{m_2} \end{pmatrix} \begin{pmatrix} y_1(t) \\ y_2(t) \\ y_3(t) \\ y_4(t) \end{pmatrix}.$$

Mit der Matrix

$$A := \begin{pmatrix} 0 & 1 & 0 & 0 \\ -\frac{g}{l} - \frac{D}{m_1} & -\frac{\gamma}{m_1} & \frac{D}{m_1} & 0 \\ 0 & 0 & 0 & 1 \\ \frac{D}{m_2} & 0 & -\frac{g}{l} - \frac{D}{m_2} & -\frac{\gamma}{m_2} \end{pmatrix}$$

stellt sich obiges Problem abgekürzt dar als:

$$\vec{y}'(t) = A \vec{y}(t).$$

Da sich mit obigem Vorgehen jedes LDGS mit Differenzialgleichungen höherer Ordnung in ein erweitertes LDGS 1. Ordnung überführen lässt, betrachten wir im Folgenden nur Systeme 1. Ordnung:

Allgemeine Formulierung:

Sei I ein Intervall, $\vec{f}(t) = \begin{pmatrix} f_1(t) \\ \vdots \\ f_n(t) \end{pmatrix} : I \rightarrow \mathbb{R}^n$ eine Vektorfunktion mit stetigen Komponenten $f_i(t)$ ($i = 1, \dots, n$) und A eine $(n \times n)$ -Matrix.

Dann betrachten wir das LDGS

$$\vec{y}'(t) = A \vec{y}(t) + \vec{f}(t). \quad (1)$$

Für $\vec{f}(t) \neq 0$ nennt man (1) ein **inhomogenes** System;
für $\vec{f}(t) = 0$ ein **homogenes** System.

Gesucht ist eine differenzierbare Vektorfunktion $\vec{y} : I \rightarrow \mathbb{C}^n$, welche das LDGS (1) erfüllt.

Die gesuchte Vektorfunktion $\vec{y}(t)$ besteht aus n Funktion

$$\vec{y}(t) = (y_1(t), \dots, y_n(t))^t;$$

jede dieser Funktionen ist differenzierbar und $\vec{y}'(t) = (y_1'(t), \dots, y_n'(t))^t$. Wir diskutieren wie bei linearen Differenzialgleichungen 1. Ordnung das homogene Problem:

➤ 12.2.2 Homogene lineare Differenzialgleichungssysteme

Wir betrachten das homogene LDGS

$$\vec{y}'(t) = A \vec{y}(t) \quad (2)$$

mit einer $(n \times n)$ -Matrix A . Obwohl die Lösungen zunächst nicht bekannt sind, ist man in der Lage, Aussagen darüber zu treffen, welche Eigenschaften die Lösungen des LDGS besitzen:

Satz 12.1: (Homogene LDGS). Die Menge aller Lösungen $\mathbb{L}_h$ eines homogenen LDGS

$$\vec{y}'(t) = A \vec{y}(t)$$

mit einer $(n \times n)$ -Matrix A ist ein n -dimensionaler Vektorraum.

Diese zentrale Aussage über homogene LDGS werden wir anhand des Eingangsbeispiels verdeutlichen. Dass die Lösungsmenge einen Vektorraum bildet, spiegelt die Gültigkeit der Superpositionsgesetzes wider. Es besagt im Falle von schwingfähigen Systemen, dass mit zwei Schwingungsformen $\vec{y}_1(t)$ und $\vec{y}_2(t)$ auch deren Überlagerung (Superposition) $\vec{y}_1(t) + \vec{y}_2(t)$ eine mögliche Schwingungsform darstellt. Außerdem ist mit jeder Schwingung $\vec{y}(t)$ auch ein Vielfaches $\alpha \vec{y}(t)$ eine Schwingungsform. Zusätzlich gilt die triviale Aussage, dass die Ruhelage $\vec{0}$ auch einen Zustand des Systems darstellt. Formal ist dies allgemein für homogene LDGS nachprüfbar:

- (1) Der Nullvektor $\vec{y}(t) = \vec{0}$ ist immer eine Lösung:
 Für $\vec{y}(t) = \vec{0}$ folgt $\vec{y}'(t) = \vec{0}' = \vec{0}$. Außerdem ist $A\vec{0} = \vec{0} \Rightarrow \vec{0}' = A\vec{0}$.
 $\Rightarrow$ Der Nullvektor $\vec{0}$ ist immer eine Lösung des homogenen LDGS.
- (2) Sind $\vec{y}_1(t)$ und $\vec{y}_2(t)$ Lösungen. Dann gilt $\vec{y}_1'(t) = A\vec{y}_1(t)$ und $\vec{y}_2'(t) = A\vec{y}_2(t)$. Mit diesen beiden Lösungen ist auch die Überlagerung $\vec{y}_1(t) + \vec{y}_2(t)$ eine Lösung, denn
 $(\vec{y}_1(t) + \vec{y}_2(t))' = \vec{y}_1'(t) + \vec{y}_2'(t) = A\vec{y}_1(t) + A\vec{y}_2(t) = A(\vec{y}_1(t) + \vec{y}_2(t)).$
- (3) Ist $\vec{y}(t)$ eine Lösung, d.h. $\vec{y}'(t) = A\vec{y}(t)$, dann ist $\alpha \vec{y}(t)$ ebenfalls eine Lösung: $(\alpha \vec{y}(t))' = \alpha \vec{y}'(t) = \alpha A\vec{y}(t) = A(\alpha \vec{y}(t)).$ $\square$

Damit ist allgemein gezeigt, dass **für alle physikalischen Systeme, die sich durch homogene LDGS beschreiben lassen, immer das Superpositionsgesetz gültig ist!** Nach dem Unterraumkriterium aus der Linearen Algebra (2.4.2) ist (1) - (3) gleichbedeutend, dass $\mathbb{L}_h$ einen Vektorraum bildet. Da jeder endlich dimensionale Vektorraum eine Basis besitzt, lässt sich jede Lösung des homogenen LDGS darstellen als Linearkombination von Ba-

sisfunktionen:

$$\vec{y}(t) = c_1 \vec{\varphi}_1(t) + c_2 \vec{\varphi}_2(t) + \dots + c_k \vec{\varphi}_k(t).$$

Die Frage ist, wieviele Basisfunktionen es gibt bzw. wieviele freie Parameter c_i die Lösungsdarstellung enthalten muss.

Kehren wir zum Pendelproblem aus Beispiel 12.16 zurück: Es gibt 4 unabhängige Möglichkeiten, das System anzuregen: Auslenkung φ_1 , Auslenkung φ_2 , Anfangsgeschwindigkeit $\dot{\varphi}_1$, Anfangsgeschwindigkeit $\dot{\varphi}_2$. Somit muss die Lösungsdarstellung für das Pendelproblem mindestens 4 freie Parameter enthalten, die unabhängig gewählt werden können. Folglich ist die Dimension von $\mathbb{L}_h \geq 4$. Da für das Pendelproblem Vektorfunktionen $\vec{y}(t) = (y_1(t), \dots, y_4(t))$ gesucht sind, ist die Dimension von $\mathbb{L}_h \leq 4$.

$$\Rightarrow \dim(\mathbb{L}_h) = 4.$$

Der folgende Satz klärt, welche Bedingungen erfüllt sein müssen, damit die Lösungen des LDGS linear unabhängig sind:

Satz 12.2: (Linear unabhängige Funktionen). Sei $\mathbb{L}_h$ die Lösungsmenge des homogenen LDGS $\vec{y}'(t) = A \vec{y}(t)$ mit einer $(n \times n)$ -Matrix A . Für n verschiedene Lösungen $\vec{\varphi}_1(t), \vec{\varphi}_2(t), \dots, \vec{\varphi}_n(t)$ sind die nachfolgenden Aussagen gleichbedeutend:

- (1) $\vec{\varphi}_1, \dots, \vec{\varphi}_n$ sind **linear unabhängige Funktionen**.
- (2) Für jedes t sind die Vektoren $\vec{\varphi}_1(t), \vec{\varphi}_2(t), \dots, \vec{\varphi}_n(t)$ linear unabhängig.
- (3) Für ein t_0 sind die Vektoren $\vec{\varphi}_1(t_0), \vec{\varphi}_2(t_0), \dots, \vec{\varphi}_n(t_0)$ linear unabhängig.

Konsequenzen der Sätze:

- (1) Sind $\vec{\varphi}_1, \vec{\varphi}_2, \dots, \vec{\varphi}_k$ Lösungsfunktionen des homogenen LDGS, dann ist jede beliebige Linearkombination

$$\vec{y}(t) = c_1 \vec{\varphi}_1(t) + c_2 \vec{\varphi}_2(t) + \dots + c_k \vec{\varphi}_k(t) \quad (c_k \in \mathbb{C}, k \in \mathbb{N})$$

ebenfalls eine Lösung.

- (2) Da $\mathbb{L}_h$ einen n -dimensionalen Vektorraum bildet, gibt es eine Basis aus n Funktionen, so dass sich die allgemeine Lösung des homogenen LDGS darstellen lässt als Linearkombination dieser Basisfunktionen:

$$\vec{y}(t) = c_1 \vec{\varphi}_1(t) + c_2 \vec{\varphi}_2(t) + \dots + c_n \vec{\varphi}_n(t).$$

- (3) Zwar ist noch nicht geklärt, wie man die Basisfunktionen berechnet, aber aufgrund von Satz 12.2 kann man bei gegebenen Lösungen entscheiden, ob eine Basis von $\mathbb{L}_h$ vorliegt oder nicht.

Da die Basisfunktionen des Vektorraumes $\mathbb{L}_h$ für die Beschreibung aller Lösungen des homogenen LDGS eine entscheidende Rolle spielen, erhalten sie eine eigene Bezeichnung:

Definition: Unter einem **Lösungs-Fundamentalsystem** des homogenen LDGS

$$\vec{y}'(t) = A\vec{y}(t) \quad (2)$$

versteht man eine Basis von Vektorfunktionen $(\vec{\varphi}_1(t), \dots, \vec{\varphi}_n(t))$ des Vektorraums $\mathbb{L}_h$ aller Lösungen.

Im n -dimensionalen Vektorraum $\mathbb{R}^n$ sind n Vektoren genau dann linear unabhängig, wenn die Determinante dieser Vektoren nicht verschwindet:

Satz 12.3:

n Lösungen $(\vec{\varphi}_1, \vec{\varphi}_2, \dots, \vec{\varphi}_n)$ von (2) bilden ein **Fundamentalsystem**

$$\Leftrightarrow \det(\vec{\varphi}_1(t_0), \vec{\varphi}_2(t_0), \dots, \vec{\varphi}_n(t_0)) \neq 0 \quad \text{für ein } t_0.$$

Anwendungsbeispiel 12.17 (Bewegung von Ladungen im Magnetfeld).

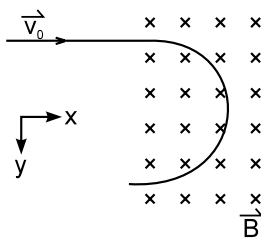


Abb. 12.14. Elektron im Magnetfeld

Die Newtonsche Bewegungsgleichung eines geladenen Teilchens $q = -e$ in einem homogenen Magnet-

feld $\vec{B} = \begin{pmatrix} 0 \\ 0 \\ B_z \end{pmatrix}$ lautet

$$\begin{aligned} m \frac{d}{dt} \vec{v}(t) &= q (\vec{v} \times \vec{B}) = q \begin{vmatrix} \vec{e}_x & v_x & 0 \\ \vec{e}_y & v_y & 0 \\ \vec{e}_z & v_z & B_z \end{vmatrix} \\ &= -e \begin{pmatrix} v_y B_z \\ -v_x B_z \\ 0 \end{pmatrix}. \end{aligned}$$

In Komponentenschreibweise gilt mit $\omega = \frac{e}{m} B_z$ für

$$\begin{aligned} \text{die 1. Komponente: } \dot{v}_x(t) &= -\frac{e}{m} B_z v_y(t) = -\omega v_y(t), \\ \text{die 2. Komponente: } \dot{v}_y(t) &= \frac{e}{m} B_z v_x(t) = \omega v_x(t), \\ \text{die 3. Komponente: } \dot{v}_z(t) &= 0. \end{aligned}$$

Aus der dritten Komponente folgt $v_z(t) = \text{const.} \Rightarrow v_z(t) = 0$, wenn keine Anfangsgeschwindigkeit in z -Richtung vorliegt.

Für die ersten beiden Komponenten $v_x(t), v_y(t)$ erhalten wir ein LDGS der Form

$$\begin{pmatrix} v_x(t) \\ v_y(t) \end{pmatrix}' = \begin{pmatrix} 0 & -\omega \\ \omega & 0 \end{pmatrix} \begin{pmatrix} v_x(t) \\ v_y(t) \end{pmatrix} \Rightarrow \vec{v}'(t) = A \vec{v}(t) \quad (*)$$

mit der (2×2) -Matrix $A = \begin{pmatrix} 0 & -\omega \\ \omega & 0 \end{pmatrix}$. Wie man durch Nachrechnen bestätigt,

sind $\vec{v}_1(t) = \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix}$ und $\vec{v}_2(t) = \begin{pmatrix} -\sin(\omega t) \\ \cos(\omega t) \end{pmatrix}$ Lösungen von $(*)$:

$$\begin{aligned} \vec{v}_1'(t) &= \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix}' = \begin{pmatrix} -\omega \sin(\omega t) \\ \omega \cos(\omega t) \end{pmatrix} \\ A \vec{v}_1(t) &= \begin{pmatrix} 0 & -\omega \\ \omega & 0 \end{pmatrix} \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix} = \begin{pmatrix} -\omega \sin(\omega t) \\ \omega \cos(\omega t) \end{pmatrix} \end{aligned}$$

$\Rightarrow \vec{v}_1'(t) = A \vec{v}_1(t)$. Analog prüft man dies für $\vec{v}_2(t)$ nach.

$\vec{v}_1(t)$ und $\vec{v}_2(t)$ sind linear unabhängig: Nach Satz 12.3 genügt es zu prüfen, dass $\det(\vec{v}_1(0), \vec{v}_2(0)) \neq 0$:

$$\det(\vec{v}_1(0), \vec{v}_2(0)) = \begin{vmatrix} 1 & 0 \\ 0 & 1 \end{vmatrix} = 1 \neq 0.$$

$\Rightarrow (\vec{v}_1, \vec{v}_2)$ bilden ein Fundamentalsystem und jede Lösung des Problems lässt sich schreiben als Linearkombination von $\vec{v}_1$ und $\vec{v}_2$

$$\vec{v}(t) = \begin{pmatrix} v_x(t) \\ v_y(t) \end{pmatrix} = c_1 \vec{v}_1(t) + c_2 \vec{v}_2(t) = c_1 \begin{pmatrix} \cos(\omega t) \\ \sin(\omega t) \end{pmatrix} + c_2 \begin{pmatrix} -\sin(\omega t) \\ \cos(\omega t) \end{pmatrix}$$

bzw. in Komponenten

$$\begin{aligned} v_x(t) &= c_1 \cos(\omega t) - c_2 \sin(\omega t) \\ v_y(t) &= c_1 \sin(\omega t) + c_2 \cos(\omega t). \end{aligned}$$

Die Konstanten c_1 und c_2 werden durch die Anfangsbedingungen des Problems festgelegt. In unserem Beispiel ist $v_x(0) = v_0$ und $v_y(0) = 0$:

$$\begin{aligned} v_x(0) &= c_1 = v_0 \\ v_y(0) &= c_2 = 0 \end{aligned} \quad \Rightarrow \quad \boxed{\begin{aligned} v_x(t) &= v_0 \cos(\omega t) \\ v_y(t) &= v_0 \sin(\omega t) \end{aligned}} \quad \square$$

② **Lösung des homogenen LDGS mit konstanten Koeffizienten**

Die Lösung von homogenen LDGS reduziert sich vollständig auf die Analyse der Matrix A . Grundlage hierfür bildet der folgende Satz:

Satz 12.4: (Lösungen von homogenen LDGS)

Sei A eine $(n \times n)$ -Matrix und $\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n$ ein Vektor, zu dem es

ein $\lambda \in \mathbb{C}$ gibt, so dass $A\vec{x} = \lambda\vec{x}$. Dann ist die Funktion

$$\vec{\varphi}(t) = \vec{x} e^{\lambda t}$$

eine Lösung des homogenen LDGS $\vec{y}'(t) = A\vec{y}(t)$.

Begründung: Sei $\vec{x} \in \mathbb{R}^n$ ein Vektor, zu dem es ein $\lambda \in \mathbb{C}$ gibt mit $A\vec{x} = \lambda\vec{x}$. Dann gilt für die Ableitung der Vektorfunktion $\vec{\varphi}(t) = \vec{x} e^{\lambda t}$:

$$\vec{\varphi}'(t) = (\vec{x} e^{\lambda t})' = \vec{x} \lambda e^{\lambda t} = (\lambda \vec{x}) e^{\lambda t} = (A\vec{x}) e^{\lambda t} = A(\vec{x} e^{\lambda t}) = A\vec{\varphi}(t) \quad . \quad \square$$

Die Frage ist also, wie verschafft man sich Vektoren $\vec{x}$ mit der Eigenschaft $A\vec{x} = \lambda\vec{x}$? Dies ist Inhalt des folgenden Abschnitts.

➤ **12.2.3 Eigenwerte und Eigenvektoren**

Definition: Sei A eine $(n \times n)$ -Matrix und $\vec{x}$ ein Vektor $\vec{x} \neq \vec{0}$. Dann heißt $\vec{x}$ **Eigenvektor** von A , wenn es eine komplexe Zahl λ gibt mit

$$A\vec{x} = \lambda\vec{x}.$$

λ heißt dann **Eigenwert** von A zum Eigenvektor $\vec{x}$.



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich eine MAPLE-Prozedur, um Eigenwerte und Eigenvektoren für (2×2) -Matrizen graphisch zu finden.

Beispiel 12.18. Gegeben ist die (3×3) -Matrix $A = \begin{pmatrix} 4 & 7 & -5 \\ 0 & 3 & -1 \\ 2 & 8 & -4 \end{pmatrix}$ und die

Vektoren $\vec{x}_1 = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$, $\vec{x}_2 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}$, $\vec{x}_3 = \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}$. Dann gilt:

$$A \vec{x}_1 = \begin{pmatrix} 4 & 7 & -5 \\ 0 & 3 & -1 \\ 2 & 8 & -4 \end{pmatrix} \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix} = \begin{pmatrix} 8+7-15 \\ 0+3-3 \\ 4+8-12 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} = 0 \cdot \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$$

$\hookrightarrow \vec{x}_1$ ist Eigenvektor zum Eigenwert $\lambda_1 = 0$.

$$A \vec{x}_2 = \begin{pmatrix} 4 & 7 & -5 \\ 0 & 3 & -1 \\ 2 & 8 & -4 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} = \begin{pmatrix} 4+7-10 \\ 0+3-2 \\ 2+8-8 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} = 1 \cdot \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}$$

$\hookrightarrow \vec{x}_2$ ist Eigenvektor zum Eigenwert $\lambda_2 = 1$.

$$A \vec{x}_3 = \begin{pmatrix} 4 & 7 & -5 \\ 0 & 3 & -1 \\ 2 & 8 & -4 \end{pmatrix} \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} -4+7-5 \\ 0+3-1 \\ -2+8-4 \end{pmatrix} = \begin{pmatrix} -2 \\ 2 \\ 2 \end{pmatrix} = 2 \cdot \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}$$

$\hookrightarrow \vec{x}_3$ ist Eigenvektor zum Eigenwert $\lambda_3 = 2$.

Um die Eigenvektoren einer Matrix A zu berechnen, bestimmt man zunächst sämtliche Eigenwerte von A und dann zu jedem Eigenwert die zugehörigen Eigenvektoren:

➤ Berechnung der Eigenwerte einer Matrix

Sei A eine $(n \times n)$ -Matrix und $\lambda \in \mathbb{C}$ sei ein Eigenwert von A . Dann gibt es einen Vektor $\vec{x} \neq 0$, so dass

$$A \vec{x} = \lambda \vec{x}. \quad (3)$$

Schreiben wir mit der Einheitsmatrix $I_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}$ den Vektor

$$\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = I_n \vec{x},$$

so gilt äquivalent zu Gl. (3)

$$A \vec{x} = \lambda I_n \vec{x} \Leftrightarrow A \vec{x} - \lambda I_n \vec{x} = \vec{0} \Leftrightarrow (A - \lambda I_n) \vec{x} = \vec{0}.$$

Dies ist ein lineares Gleichungssystem mit der Matrix $B = (A - \lambda I_n)$. Ein lineares Gleichungssystem ist nach dem Fundamentalsatz über LGS (siehe

3.3.1) genau dann eindeutig lösbar, wenn $\det(B) \neq 0$. Im Falle der eindeutigen Lösbarkeit ist $\vec{x} = \vec{0}$ die einzige Lösung. Diese Lösung liefert dann aber keinen Eigenvektor, da die Eigenvektoren von Null verschieden sein müssen. Damit wir also einen Eigenvektor erhalten, darf $\vec{x} = \vec{0}$ nicht die einzige Lösung sein, d.h. das LGS darf nicht eindeutig lösbar sein! Somit muss $\det(B) = 0$ sein bzw.

$$\det(A - \lambda I_n) = 0.$$

Es gilt allgemein

Satz 12.5: (Eigenwerte einer Matrix).

λ ist Eigenwert der Matrix $A \Leftrightarrow \det(A - \lambda I_n) = 0$.

Beispiel 12.19. Gegeben ist die (3×3) -Matrix $A = \begin{pmatrix} 5 & 7 & -5 \\ 0 & 4 & -1 \\ 2 & 8 & -3 \end{pmatrix}$. Gesucht sind alle Eigenwerte dieser Matrix.

Zur Bestimmung der Eigenwerte gehen wir zur Matrix $A - \lambda I_3$ über

$$A - \lambda I_3 = \begin{pmatrix} 5 & 7 & -5 \\ 0 & 4 & -1 \\ 2 & 8 & -3 \end{pmatrix} - \lambda \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 5 - \lambda & 7 & -5 \\ 0 & 4 - \lambda & -1 \\ 2 & 8 & -3 - \lambda \end{pmatrix}$$

und berechnen die Determinante

$$\begin{aligned} \det(A - \lambda I_3) &= \begin{vmatrix} 5 - \lambda & 7 & -5 \\ 0 & 4 - \lambda & -1 \\ 2 & 8 & -3 - \lambda \end{vmatrix} \\ &= (5 - \lambda) \begin{vmatrix} 4 - \lambda & -1 \\ 8 & -3 - \lambda \end{vmatrix} + 2 \begin{vmatrix} 7 & -5 \\ 4 - \lambda & -1 \end{vmatrix} \\ &= -\lambda^3 + 6\lambda^2 - 11\lambda + 6 = -(\lambda - 1)(\lambda - 2)(\lambda - 3). \end{aligned}$$

Aus $\det(A - \lambda I_3) \stackrel{!}{=} 0$ erhalten wir die Eigenwerte

$$\lambda_1 = 1, \lambda_2 = 2 \quad \text{und} \quad \lambda_3 = 3.$$

□

Zur (3×3) -Matrix A ist $\det(A - \lambda I_3) = -\lambda^3 + 6\lambda^2 - 11\lambda + 6$ ein Polynom in λ vom Grade 3. Allgemein gilt

Satz 12.6: (Charakteristisches Polynom). Ist A eine $(n \times n)$ -Matrix, dann ist

$$P(\lambda) := \det(A - \lambda I_n)$$

ein Polynom n -ten Grades in λ . $P(\lambda)$ heißt das **charakteristische Polynom**. Die Nullstellen des charakteristischen Polynoms sind die **Eigenwerte der Matrix A** .

Da nach dem Fundamentalsatz der Algebra (siehe 5.2.7) jedes komplexe Polynom vom Grad n genau n Nullstellen besitzt, zerfällt das charakteristische Polynom im Komplexen immer in Linearfaktoren. Die Eigenwerte mit entsprechender Vielfachheit sind dadurch festgelegt. In der Praxis ist man allerdings bei Systemen mit mehr als 3 Gleichungen auf numerische Verfahren wie z.B. das Newton-Verfahren in 7.9 angewiesen, um die Eigenwerte als Nullstellen des charakteristischen Polynoms zu bestimmen. Auch MAPLE verwendet numerische Algorithmen, um die Nullstellen näherungsweise zu bestimmen.

⊗ Berechnung der Eigenvektoren

Nachdem die Eigenwerte einer Matrix bestimmt sind, berechnet man zu jedem Eigenwert die zugehörigen Eigenvektoren, indem das lineare Gleichungssystem $(A - \lambda I_n) \vec{x} = \vec{0}$ gelöst wird:

Ist λ ein Eigenwert der Matrix A , so sind alle Eigenvektoren zum Eigenwert λ gegeben als Lösung des linearen Gleichungssystems

$$(A - \lambda I_n) \vec{x} = \vec{0}.$$

Beispiel 12.20 (Eigenvektoren zu gegebenen Eigenwerten). Gegeben

ist die Matrix $A = \begin{pmatrix} 5 & 7 & -5 \\ 0 & 4 & -1 \\ 2 & 8 & -3 \end{pmatrix}$ aus Beispiel 12.19 mit den Eigenwerten

$\lambda_1 = 1, \lambda_2 = 2, \lambda_3 = 3$. Gesucht sind die zu den Eigenwerten gehörenden Eigenvektoren.

i) Berechnung der Eigenvektoren zum Eigenwert $\lambda_1 = 1$: Gesucht sind Vektoren $\vec{x} \neq \vec{0}$, so dass $(A - \lambda_1 I_3) \vec{x} = \vec{0}$. Zu lösen ist das lineare Gleichungssystem

$$\begin{pmatrix} 5-1 & 7 & -5 \\ 0 & 4-1 & -1 \\ 2 & 8 & -3-1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} :$$

$$\hookrightarrow \left(\begin{array}{ccc|c} 4 & 7 & -5 & 0 \\ 0 & 3 & -1 & 0 \\ 2 & 8 & -4 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 4 & 7 & -5 & 0 \\ 0 & 3 & -1 & 0 \\ 0 & -9 & 3 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 4 & 7 & -5 & 0 \\ 0 & 3 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right).$$

Damit ist $x_3 = t$; $3x_2 - t = 0 \hookrightarrow x_2 = \frac{1}{3}t$; $4x_1 + \frac{7}{3}t - 5t = 0 \hookrightarrow x_1 = \frac{2}{3}t$.

$$\Rightarrow \mathbb{L}_1 = \left\{ \vec{x} \in \mathbb{R}^3 : \vec{x} = t \begin{pmatrix} \frac{2}{3} \\ \frac{1}{3} \\ 1 \end{pmatrix}, \quad t \in \mathbb{R} \right\}$$

ist die Lösungsmenge. Setzt man z.B. $t = 3$, so erhält man $\vec{x}_1 = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$ als einen Eigenvektor zum Eigenwert $\lambda_1 = 1$.

ii) Berechnung der Eigenvektoren zum Eigenwert $\lambda_2 = 2$: Zu lösen ist das lineare Gleichungssystem $(A - \lambda_2 I_3) \vec{x} = \vec{0}$:

$$\left(\begin{array}{ccc|c} 5-2 & 7 & -5 & 0 \\ 0 & 4-2 & -1 & 0 \\ 2 & 8 & -3-2 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 3 & 7 & -5 & 0 \\ 0 & 2 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right).$$

Damit ist $x_3 = t$; $2x_2 - t = 0 \hookrightarrow x_2 = \frac{1}{2}t$; $3x_1 + \frac{7}{2}t - 5t = 0 \hookrightarrow x_1 = \frac{1}{2}t$.

$$\Rightarrow \mathbb{L}_2 = \left\{ \vec{x} \in \mathbb{R}^3 : \vec{x} = t \begin{pmatrix} \frac{1}{2} \\ \frac{1}{2} \\ 1 \end{pmatrix}, \quad t \in \mathbb{R} \right\}$$

ist die Lösungsmenge des linearen Gleichungssystems. Einen Eigenvektor zum Eigenwert $\lambda_2 = 2$ erhält man z.B. für $t = 2$ durch $\vec{x}_2 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}$.

iii) Berechnung der Eigenvektoren zum Eigenwert $\lambda_3 = 3$: Durch Lösen des linearen Gleichungssystems $(A - \lambda_3 I_3) \vec{x} = \vec{0}$ erhält man z.B. $\vec{x}_3 = \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}$ als einen Eigenvektor zum Eigenwert $\lambda_3 = 3$. $\square$

Da zu gegebenem Eigenwert die zugehörigen Eigenvektoren $\vec{x}$ immer Lösungen eines homogenen, linearen Gleichungssystems sind, gilt für zwei Eigenvektoren $\vec{x}_1$ und $\vec{x}_2$ zum selben Eigenwert λ :

$$\begin{aligned} A(t_1 \vec{x}_1 + t_2 \vec{x}_2) &= t_1 A \vec{x}_1 + t_2 A \vec{x}_2 = t_1 \lambda \vec{x}_1 + t_2 \lambda \vec{x}_2 \\ &= \lambda(t_1 \vec{x}_1 + t_2 \vec{x}_2). \end{aligned}$$

Folglich ist $t_1 \vec{x}_1 + t_2 \vec{x}_2$ für beliebige $t_1, t_2 \in \mathbb{C}$ ebenfalls ein Eigenvektor zum Eigenwert λ . Mit dem UVR-Kriterium (siehe 2.4.2) gilt damit der folgende Satz:

Satz 12.7: Ist A eine $(n \times n)$ -Matrix und λ ein Eigenwert von A . Dann bildet die Menge der Eigenvektoren zum Eigenwert λ einen Vektorraum.

Bezeichnung:

$$\begin{aligned} \text{Eig}(A, \lambda) &= \text{Menge aller Eigenvektoren zum Eigenwert } \lambda \\ &= \text{Eigenraum von } A \text{ zum Eigenwert } \lambda. \end{aligned}$$

► 12.2.4 Lösen von homogenen LDGS mit konstanten Koeffizienten

Kommen wir nun zu unserem ursprünglich gestellten Problem, dem Lösen von homogenen LDGS, zurück. Zusammenfassend können wir mit den Begriffen aus dem vorigen Abschnitt formulieren:

Folgerung/Zusammenfassung: Besitzt die $(n \times n)$ -Matrix A eine Basis von Eigenvektoren $\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n$ zu den Eigenwerten $\lambda_1, \dots, \lambda_n \in \mathbb{C}$, so bilden die Vektorfunktionen

$$\vec{\varphi}_k(t) = \vec{x}_k e^{\lambda_k t} \quad (k = 1, \dots, n)$$

ein Lösungs-Fundamentalsystem des homogenen LDGS

$$\vec{y}'(t) = A \vec{y}(t) .$$

Die in 12.2.3 beschriebene Vorgehensweise zur Bestimmung der Eigenwerte und zugehörigen Eigenvektoren ist ausreichend, um ein Fundamentalsystem des LDGS $\vec{y}'(t) = A \vec{y}(t)$ zu berechnen, wenn man eine Basis aus Eigenvektoren findet. Dies ist aber nur unter gewissen Voraussetzungen der Fall, welche der folgende Satz zusammenfasst. Auf die weiterführende Theorie werden wir nicht näher eingehen.

Satz 12.8: Sei A eine $(n \times n)$ -Matrix.

- (1) Besitzt das charakteristische Polynom $P(\lambda) = \det(A - \lambda I_n)$ n verschiedene Nullstellen, dann gibt es eine Basis aus Eigenvektoren.
- (2) Existieren zu jedem Eigenwert der Vielfachheit m , m linear unabhängige Eigenvektoren, dann gibt es eine Basis aus Eigenvektoren.
- (3) Ist A eine reelle, symmetrische Matrix (d.h. $A = A^t$), dann gibt es eine Basis aus Eigenvektoren.
- (4) Ist A eine komplexe, hermitesche Matrix (d.h. $A = \overline{A^t}$), dann gibt es eine Basis aus Eigenvektoren.

Bemerkung: Stimmt die Dimension des Eigenraumes $\text{Eig}(A, \lambda)$ mit der Vielfachheit des Eigenwerts überein, dann existiert eine Basis aus Eigenvektoren. Man kann allgemeiner sogar zeigen, dass diese Bedingung nicht nur notwendig, sondern auch hinreichend ist: **Es existiert eine Basis aus Eigenvektoren genau dann, wenn für alle Eigenwerte die Vielfachheit mit der Dimension des Eigenraumes übereinstimmt.** Da solche Matrizen eine besondere Rolle spielen, bezeichnet man sie als *diagonalisierbare* Matrizen.

Beispiel 12.21 (Musterbeispiel: Eigenwerte und Eigenvektoren).

Gesucht ist ein Lösungs-Fundamentalsystem des LDGS

$$\vec{y}'(t) = A \vec{y}(t) \quad \text{mit} \quad A = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}.$$

(i) **Bestimmung der Eigenwerte** von A :

$$\begin{aligned} P(\lambda) &= \det(A - \lambda I_3) = \begin{vmatrix} 1-\lambda & 1 & 1 \\ 1 & 1-\lambda & 1 \\ 1 & 1 & 1-\lambda \end{vmatrix} = \\ &= (1-\lambda) \begin{vmatrix} 1-\lambda & 1 \\ 1 & 1-\lambda \end{vmatrix} - \begin{vmatrix} 1 & 1 \\ 1 & 1-\lambda \end{vmatrix} + \begin{vmatrix} 1 & 1 \\ 1-\lambda & 1 \end{vmatrix} \\ &= -\lambda^2 (\lambda - 3). \end{aligned}$$

Die Eigenwerte sind die Nullstellen des charakteristischen Polynoms:

$$\begin{aligned} P(\lambda) \stackrel{!}{=} 0 &\hookrightarrow \lambda_1 = 0 \quad \text{Eigenwert mit Vielfachheit 2.} \\ &\lambda_2 = 3 \quad \text{Eigenwert mit Vielfachheit 1.} \end{aligned}$$

(ii) **Bestimmung der Eigenvektoren** von A :

Eigenvektoren zum Eigenwert $\lambda_1 = 0$:

$$(A - 0 \cdot I_3) \vec{x} = 0 \hookrightarrow \left(\begin{array}{ccc|c} 1 & 1 & 1 & 0 \\ 1 & 1 & 1 & 0 \\ 1 & 1 & 1 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right).$$

$\hookrightarrow x_3 = r; x_2 = t; x_1 = -r - t$. Damit folgt

$$\text{Eig}(A, 0) = \left\{ \vec{x} \in \mathbb{R}^3 : \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = r \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} + t \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} ; r, t \in \mathbb{R} \right\}.$$

Die Dimension des Eigenraums $\text{Eig}(A, 0)$ ist 2 und gleich der Vielfachheit des Eigenwerts. Zwei linear unabhängige Eigenvektoren sind z.B.

$$\vec{x}_1 = \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} \quad \text{und} \quad \vec{x}_2 = \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix}.$$

Eigenvektoren zum Eigenwert $\lambda_2 = 3$:

$$(A - 3 I_3) \vec{x} = 0 \hookrightarrow \left(\begin{array}{ccc|c} -2 & 1 & 1 & 0 \\ 1 & -2 & 1 & 0 \\ 1 & 1 & -2 & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} -2 & 1 & 1 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right).$$

$\hookrightarrow x_3 = r; x_2 = r; -2x_1 + r + r = 0 \hookrightarrow x_1 = r$.

$$\Rightarrow \text{Eig}(A, 3) = \left\{ \vec{x} \in \mathbb{R}^3 : \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = r \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}; r \in \mathbb{R} \right\}.$$

Die Dimension des Eigenraumes $\text{Eig}(A, 3)$ ist 1 und gleich der Vielfachheit des Eigenwertes. Ein Eigenvektor ist z.B. $\vec{x}_3 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$.

(iii) Ein **Fundamentalsystem** von $\vec{y}'(t) = A \vec{y}(t)$ ist damit

$$\begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} e^{0t}, \quad \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} e^{0t}, \quad \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} e^{3t}$$

und die allgemeine Lösung lautet

$$\vec{y}(t) = c_1 \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} e^{0t} + c_2 \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} e^{0t} + c_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} e^{3t}$$

mit frei wählbaren Konstanten c_1, c_2, c_3 . □

Beispiel 12.22 (Musterbeispiel: Lösen von LDGS).

Das Lösen von LDGS erster Ordnung soll am Beispiel der folgenden drei gekoppelten Differenzialgleichungen zusammengefasst werden. Gegeben ist das System von Differenzialgleichungen:

$$\begin{aligned} 4y_2(t) &= y_2'(t) + y_3(t) \\ 5y_1(t) + 7y_2(t) &= y_1'(t) + 5y_3(t) \\ y_3'(t) &= 2y_1(t) + 8y_2(t) - 3y_3(t). \end{aligned} \tag{*}$$

Gesucht sind die Lösungen $y_1(t)$, $y_2(t)$ und $y_3(t)$ zu den Anfangsbedingungen

$$y_1(0) = 3, y_2(0) = 2, y_3(0) = 1.$$

1. Aufstellen des LDGS: Alle Terme mit Ableitungen werden auf die linke und die gesuchten Funktionen auf die rechte Seite gebracht.

$$\begin{aligned} y_1'(t) &= 5y_1(t) + 7y_2(t) - 5y_3(t) \\ y_2'(t) &= 4y_2(t) - y_3(t) \\ y_3'(t) &= 2y_1(t) + 8y_2(t) - 3y_3(t) \end{aligned}$$

2. Aufstellen der Systemmatrix:

$$\begin{pmatrix} y_1(t) \\ y_2(t) \\ y_3(t) \end{pmatrix}' = \begin{pmatrix} 5 & 7 & -5 \\ 0 & 4 & -1 \\ 2 & 8 & -3 \end{pmatrix} \begin{pmatrix} y_1(t) \\ y_2(t) \\ y_3(t) \end{pmatrix} \Rightarrow A = \begin{pmatrix} 5 & 7 & -5 \\ 0 & 4 & -1 \\ 2 & 8 & -3 \end{pmatrix}.$$

3. Bestimmung der Eigenwerte und Eigenvektoren: Nach Beispiel 12.19

und 12.20 sind $\vec{x}_1 = \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}$, $\vec{x}_2 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}$ und $\vec{x}_3 = \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix}$ Eigenvektoren zu den Eigenwerten $\lambda_1 = 1$, $\lambda_2 = 2$ und $\lambda_3 = 3$.

4. Fundamentalsystem: Durch die Kenntnis der Eigenwerten und der zugehörigen Eigenvektoren ist man in der Lage, ein Fundamentalsystem durch $\vec{x}_1 e^{\lambda_1 t}$, $\vec{x}_2 e^{\lambda_2 t}$, $\vec{x}_3 e^{\lambda_3 t}$ anzugeben

$$\begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix} e^{1t}, \quad \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} e^{2t}, \quad \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix} e^{3t}.$$

5. Allgemeine Lösung: Die allgemeine Lösung $\vec{y}(t)$ ist dann eine Linearkombination der Fundamentallösungen

$$\vec{y}(t) = c_1 \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix} e^{1t} + c_2 \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} e^{2t} + c_3 \begin{pmatrix} -1 \\ 1 \\ 1 \end{pmatrix} e^{3t}.$$

In Komponentenschreibweise lauten die gesuchten Funktionen

$$\begin{aligned} y_1(t) &= 2c_1 e^{1t} + 1c_2 e^{2t} - c_3 e^{3t} \\ y_2(t) &= 1c_1 e^{1t} + 1c_2 e^{2t} + c_3 e^{3t} \\ y_3(t) &= 3c_1 e^{1t} + 2c_2 e^{2t} + c_3 e^{3t}. \end{aligned}$$

6. Bestimmung der Koeffizienten: Die Koeffizienten bestimmen sich über die Anfangsbedingungen:

$$\begin{aligned} y_1(0) &= 2c_1 + 1c_2 - c_3 = 3 \\ y_2(0) &= 1c_1 + 1c_2 + c_3 = 2 \\ y_3(0) &= 3c_1 + 2c_2 + c_3 = 1. \end{aligned}$$

Dies ist ein lineares Gleichungssystem für die Konstanten c_1 , c_2 und c_3 der Form

$$\left(\begin{array}{ccc|c} 2 & 1 & -1 & 3 \\ 1 & 1 & 1 & 2 \\ 3 & 2 & 1 & 1 \end{array} \right) \hookrightarrow \left(\begin{array}{ccc|c} 1 & 1 & 1 & 2 \\ 0 & -1 & -3 & -1 \\ 0 & 0 & -1 & 4 \end{array} \right),$$

das z.B. mit dem Gauß-Algorithmus gelöst wird. Es ergeben sich die Konstanten zu $c_1 = -7$, $c_2 = 13$ und $c_3 = -4$. Damit lauten die Lösungen

des LDGS (*)

$$\begin{aligned}y_1(t) &= -14 e^{1t} + 13 e^{2t} + 4 e^{3t} \\y_2(t) &= -7 e^{1t} + 13 e^{2t} - 4 e^{3t} \\y_3(t) &= -21 e^{1t} + 26 e^{2t} - 4 e^{3t} .\end{aligned}$$

□

Die bisher beschriebenen Methoden lassen sich direkt auf LDGS der Form

$$\vec{y}''(t) = A \vec{y}(t)$$

übertragen. Dies ist deshalb von besonderem Interesse, da Schwingungsprobleme ohne Reibung sich durch solche LDGS beschreiben lassen.

Satz 12.9: Ist A eine $(n \times n)$ -Matrix und $\vec{x}$ ein Eigenvektor zum Eigenwert λ . Dann sind die Funktionen

$$\vec{y}_1(t) = \vec{x} e^{+\sqrt{\lambda} t} \quad \text{und} \quad \vec{y}_2(t) = \vec{x} e^{-\sqrt{\lambda} t}$$

Lösungen des LDGS zweiter Ordnung

$$\vec{y}''(t) = A \vec{y}(t).$$

Begründung: Ist $\vec{x}$ ein Eigenvektor zum Eigenwert λ . Dann gilt für $\vec{y}(t) := \vec{x} e^{\sqrt{\lambda} t}$:

$$\begin{aligned}\vec{y}''(t) &= \left(\vec{x} e^{\sqrt{\lambda} t} \right)'' = \left(\vec{x} \sqrt{\lambda} e^{\sqrt{\lambda} t} \right)' = \vec{x} \sqrt{\lambda}^2 e^{\sqrt{\lambda} t} = \lambda \vec{x} e^{\sqrt{\lambda} t} \\&= A \vec{x} e^{\sqrt{\lambda} t} = A \vec{y}(t).\end{aligned}$$

D.h. $\vec{y}(t)$ ist eine Lösung des LDGS zweiter Ordnung $\vec{y}''(t) = A \vec{y}(t)$. Analog zeigt man, dass auch $\vec{x} e^{-\sqrt{\lambda} t}$ eine Lösung des LDGS ist. □

Bemerkung: Besitzt A eine Basis aus Eigenvektoren $(\vec{x}_1, \dots, \vec{x}_n)$ mit den zugehörigen Eigenwerten $\lambda_i \neq 0$ ($i = 1, \dots, n$), dann ist

$$\vec{x}_1 e^{\sqrt{\lambda_1} t}, \vec{x}_1 e^{-\sqrt{\lambda_1} t}, \dots, \vec{x}_n e^{\sqrt{\lambda_n} t}, \vec{x}_n e^{-\sqrt{\lambda_n} t}$$

ein Lösungs-Fundamentalsystem zu $\vec{y}''(t) = A \vec{y}(t)$.

Anwendungsbeispiel 12.23 (Gekoppelte Pendel ohne Reibung, mit MAPLE-Worksheet).

Kommen wir auf das Einführungsbeispiel 12.16 der gekoppelten Pendel zurück. Vernachlässigen wir Reibungskräfte, ist das System von Differenzialgleichungen für die Winkelauslenkungen $\varphi_1(t)$ und $\varphi_2(t)$ gegeben durch

$$\begin{aligned}\ddot{\varphi}_1(t) &= -\frac{g}{l} \varphi_1(t) + \frac{D}{m} (\varphi_2(t) - \varphi_1(t)) \\ \ddot{\varphi}_2(t) &= -\frac{g}{l} \varphi_2(t) + \frac{D}{m} (\varphi_1(t) - \varphi_2(t)).\end{aligned}$$

Für $\vec{\varphi}(t) := \begin{pmatrix} \varphi_1(t) \\ \varphi_2(t) \end{pmatrix}$ gilt mit $A = \begin{pmatrix} -\frac{g}{l} - \frac{D}{m} & \frac{D}{m} \\ \frac{D}{m} & -\frac{g}{l} - \frac{D}{m} \end{pmatrix}$

$$\vec{\varphi}''(t) = A \vec{\varphi}(t).$$

(i) Berechnung der Eigenwerte:

$$\begin{aligned}P(\lambda) &= \det(A - \lambda I_2) = \begin{vmatrix} -\frac{g}{l} - \frac{D}{m} - \lambda & \frac{D}{m} \\ \frac{D}{m} & -\frac{g}{l} - \frac{D}{m} - \lambda \end{vmatrix} \\ &= \left(-\frac{g}{l} - \frac{D}{m} - \lambda\right)^2 - \left(\frac{D}{m}\right)^2 \stackrel{!}{=} 0.\end{aligned}$$

Die Eigenwerte sind die Nullstellen des charakteristischen Polynoms

$$\begin{aligned}P(\lambda) &\stackrel{!}{=} 0 \Rightarrow \left(\frac{g}{l} + \frac{D}{m}\right) + \lambda_{1/2} = \pm \frac{D}{m} \\ &\hookrightarrow \boxed{\lambda_1 = -\frac{g}{l}} \quad \text{und} \quad \boxed{\lambda_2 = -\frac{g}{l} - 2\frac{D}{m}}.\end{aligned}$$

(ii) Berechnung der Eigenvektoren:

$$\lambda_1 = -\frac{g}{l} : (A - \lambda_1 I_2) \vec{x} = 0 \hookrightarrow \left(\begin{array}{cc|c} -\frac{D}{m} & \frac{D}{m} & 0 \\ \frac{D}{m} & -\frac{D}{m} & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cc|c} -\frac{D}{m} & \frac{D}{m} & 0 \\ 0 & 0 & 0 \end{array} \right) :$$

Eigenvektor zum Eigenwert λ_1 ist $\vec{x}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ (gleichphasig).

$$\lambda_2 = -\frac{g}{l} - 2\frac{D}{m} : (A - \lambda_2 I_2) \vec{x} = 0 \hookrightarrow \left(\begin{array}{cc|c} \frac{D}{m} & \frac{D}{m} & 0 \\ \frac{D}{m} & \frac{D}{m} & 0 \end{array} \right) \hookrightarrow \left(\begin{array}{cc|c} \frac{D}{m} & \frac{D}{m} & 0 \\ 0 & 0 & 0 \end{array} \right) :$$

Eigenvektor zum Eigenwert λ_2 ist $\vec{x}_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$ (gegenphasig).

(iii) **Aufstellen des Fundamentalsystems:** Mit den Eigenvektoren und zugehörigen Eigenwerten stellen wir das komplexe Fundamentalsystem auf:

$$\vec{x}_1 e^{\sqrt{\lambda_1} t}, \quad \vec{x}_1 e^{-\sqrt{\lambda_1} t}, \quad \vec{x}_2 e^{\sqrt{\lambda_2} t}, \quad \vec{x}_2 e^{-\sqrt{\lambda_2} t}$$

mit

$$\sqrt{\lambda_1} = \sqrt{-\frac{g}{l}} = i \sqrt{\frac{g}{l}} = i \omega_1$$

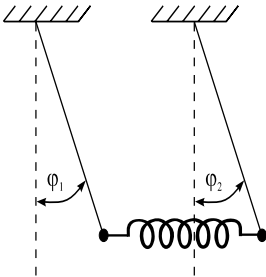
und

$$\sqrt{\lambda_2} = \sqrt{-\frac{g}{l} - 2 \frac{D}{m}} = i \sqrt{\frac{g}{l} + 2 \frac{D}{m}} = i \omega_2.$$

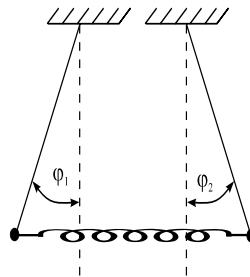
(iv) **Interpretation:**

$\omega_1 = \sqrt{\frac{g}{l}}$ ist die Eigenfrequenz des Pendels ohne Federkopplung. Zu dieser Frequenz gehört der Eigenvektor $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$, welches einem *gleichphasigen* Auslenken der Pendel entspricht (siehe linke Abb. (1)). Die Feder ist nicht bemerkbar und beide Pendel schwingen mit der Eigenfrequenz eines Einzelpendels ohne Kopplung.

$\omega_2 = \sqrt{\frac{g}{l} + 2 \frac{D}{m}}$ ist die Eigenfrequenz des Pendels mit Federkopplung, wenn die beiden Massen gegenphasig ausgelenkt werden. Der zugehörige Eigenvektor ist $\begin{pmatrix} 1 \\ -1 \end{pmatrix}$ (siehe rechte Abb. (2)). Durch die entgegengesetzte Auslenkung der Pendel macht sich die Federauslenkung doppelt bemerkbar, welches sich in dem Faktor $2 \frac{D}{m}$ bei der Frequenz widerspiegelt.



(1) Gleichphasige Auslenkung



(2) Gegenphasige Auslenkung

Die zu ω_1 und ω_2 gehörenden Schwingungen nennt man *Grundschiwingungen*. Regt man das System mit einem **Eigenvektor** an, so wird nur die zugehörige **Eigenfrequenz** angeregt. Alle anderen Schwingungsformen sind Überlagerungen dieser Grundschiwingungen. Die allgemeine Lösung ist mit beliebigen komplexen Konstanten c_1, c_2, c_3, c_4 gegeben durch

$$\vec{\varphi}(t) = c_1 \begin{pmatrix} 1 \\ 1 \end{pmatrix} e^{i \omega_1 t} + c_2 \begin{pmatrix} 1 \\ 1 \end{pmatrix} e^{-i \omega_1 t} + c_3 \begin{pmatrix} 1 \\ -1 \end{pmatrix} e^{i \omega_2 t} + c_4 \begin{pmatrix} 1 \\ -1 \end{pmatrix} e^{-i \omega_2 t}.$$

(v) **Übergang zu einem reellen Fundamentalsystem:** Aufgrund der Eulerschen Formel (siehe 5.1)

$$e^{i\omega t} = \cos \omega t + i \sin \omega t$$

gilt für beliebiges t

$$\begin{aligned}\cos \omega t &= \frac{1}{2} (e^{i\omega t} + e^{-i\omega t}) \\ \sin \omega t &= \frac{1}{2i} (e^{i\omega t} - e^{-i\omega t}).\end{aligned}$$

Mit den Lösungen $\vec{\psi}_1(t) = \vec{x}_1 e^{i\omega_1 t}$ und $\vec{\psi}_2(t) = \vec{x}_1 e^{-i\omega_1 t}$ erfüllen auch die beiden Überlagerungen

$$\begin{aligned}\frac{1}{2} \vec{\psi}_1(t) + \frac{1}{2} \vec{\psi}_2(t) &= \vec{x}_1 \frac{1}{2} (e^{i\omega_1 t} + e^{-i\omega_1 t}) = \vec{x}_1 \cos(\omega_1 t) \\ \frac{1}{2i} \vec{\psi}_1(t) - \frac{1}{2i} \vec{\psi}_2(t) &= \vec{x}_1 \frac{1}{2i} (e^{i\omega_1 t} - e^{-i\omega_1 t}) = \vec{x}_1 \sin(\omega_1 t)\end{aligned}$$

das LDGS. Analog erhält man $\vec{x}_2 \cos(\omega_2 t)$, $\vec{x}_2 \sin(\omega_2 t)$ als Lösungen. Insgesamt haben wir somit vier reelle Lösungen durch Linearkombination der vier komplexen Lösungen erhalten:

$$\vec{x}_1 \cos(\omega_1 t), \quad \vec{x}_1 \sin(\omega_1 t), \quad \vec{x}_2 \cos(\omega_2 t), \quad \vec{x}_2 \sin(\omega_2 t).$$

Die allgemeine reelle Lösung lautet daher

$$\begin{aligned}\vec{\varphi}(t) &= \begin{pmatrix} \varphi_1(t) \\ \varphi_2(t) \end{pmatrix} = c_1 \begin{pmatrix} 1 \\ 1 \end{pmatrix} \cos(\omega_1 t) + c_2 \begin{pmatrix} 1 \\ 1 \end{pmatrix} \sin(\omega_1 t) \\ &\quad + c_3 \begin{pmatrix} 1 \\ -1 \end{pmatrix} \cos(\omega_2 t) + c_4 \begin{pmatrix} 1 \\ -1 \end{pmatrix} \sin(\omega_2 t)\end{aligned}$$

bzw. in Komponentendarstellung

$$\begin{aligned}\varphi_1(t) &= c_1 \cos(\omega_1 t) + c_2 \sin(\omega_1 t) + c_3 \cos(\omega_2 t) + c_4 \sin(\omega_2 t) \\ \varphi_2(t) &= c_1 \cos(\omega_1 t) + c_2 \sin(\omega_1 t) - c_3 \cos(\omega_2 t) - c_4 \sin(\omega_2 t).\end{aligned}$$

Die Konstanten c_1, c_2, c_3, c_4 werden durch die Anfangsbedingungen festgelegt.

(vi) **Lösung für unterschiedliche Anfangsbedingungen:**

a) Mit $\varphi_1(0) = \varphi_0, \varphi_2(0) = \varphi_0, \dot{\varphi}_1(0) = 0, \dot{\varphi}_2(0) = 0$ regt man die gleichphasige Grundschwingung an. Aus den Anfangsbedingungen folgt $c_1 = \varphi_0, c_2 = c_3 = c_4 = 0$. Somit lautet die Lösung

$$\begin{aligned}\varphi_1(t) &= \varphi_0 \cos(\omega_1 t) \\ \varphi_2(t) &= \varphi_0 \cos(\omega_1 t).\end{aligned}$$

Die Pendel schwingen gleichphasig mit der Frequenz $\omega_1 = \sqrt{\frac{g}{l}}$.

b) Für $\varphi_1(0) = -\varphi_0$, $\varphi_2(0) = \varphi_0$, $\dot{\varphi}_1(0) = 0$, $\dot{\varphi}_2(0) = 0$ regt man die gegenphasige Grundschiwingung an. Aus den Anfangsbedingungen folgt $c_3 = -\varphi_0$, $c_1 = c_2 = c_4 = 0$

$$\Rightarrow \begin{aligned} \varphi_1(t) &= -\varphi_0 \cos(\omega_2 t) \\ \varphi_2(t) &= \varphi_0 \cos(\omega_2 t). \end{aligned}$$

Die Pendel schwingen gegenphasig mit der Frequenz $\omega_2 = \sqrt{\frac{g}{l} + 2\frac{D}{m}}$.

c) Wird nur das erste Pendel ausgelenkt, lauten die Anfangsbedingungen

$$\varphi_1(0) = -\varphi_0, \varphi_2(0) = 0, \dot{\varphi}_1(0) = 0, \dot{\varphi}_2(0) = 0.$$

Aus diesen Anfangsbedingungen erhalten wir vier lineare Gleichungen für die zu bestimmenden Koeffizienten:

$$\begin{aligned} c_1 + c_3 &= -\varphi_0 \\ c_1 - c_3 &= 0 \\ c_2 \omega_1 + c_4 \omega_2 &= 0 \\ c_2 \omega_1 - c_4 \omega_2 &= 0. \end{aligned}$$

Die Lösung des linearen Gleichungssystem ist

$$c_1 = \frac{1}{2}\varphi_0, c_2 = 0, c_3 = \frac{1}{2}\varphi_0, c_4 = 0.$$

In Abb. 12.15 wird für die Pendellänge $l = 2$, Federkonstante $D = 0.2$ und Masse $m = 1$ die Lösung für $\varphi_1(t)$ graphisch dargestellt. Die Anfangsauslenkung ist hierbei $\varphi_0 = -0.1$.

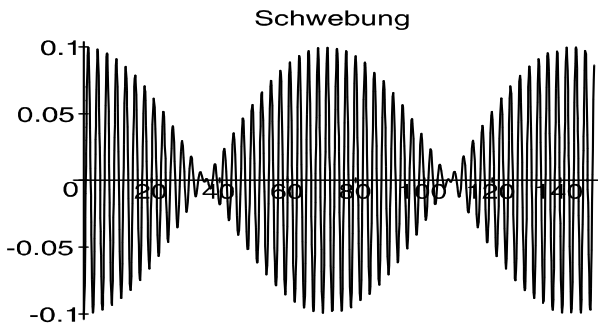


Abb. 12.15. Schwebung beim Doppelpendel ohne Reibung

Man kann die entstandene Schwingung sehr gut als Schwebung identifizieren. □



Visualisierung mit MAPLE: Auf der CD-Rom befindet sich eine Animation, in der diese Schwebung der Pendel visualisiert wird. Das [MAPLE-Worksheet](#) ist so aufgebaut, dass auch andere Anfangsbedingungen gewählt werden können. Die zugehörigen Koeffizienten werden bestimmt und die Animation dargestellt.



Hinweis: Auf der CD-Rom befindet sich eine [MAPLE-Worksheet](#), welches die Schwingungen der Pendel mit Reibung behandelt. Im zugehörigen Worksheet wird die dann auftretende Schwebung der Pendel in Form einer Animation visualisiert. Die Winkelauslenkung für das erste Pendel $\varphi_1(t)$ ist in Abb. 12.16 angegeben.

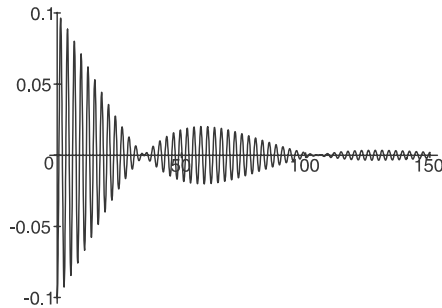


Abb. 12.16. Schwebung beim Doppelpendel mit Reibung

Zusammenfassung: Durch die Bestimmung der Eigenwerte und Eigenvektoren ist man in der Lage, ein Fundamentalsystem des LDGS

$$\vec{y}''(t) = A \vec{y}(t)$$

zu bestimmen: Ist $\vec{x}_k$ ein Eigenvektor zum Eigenwert λ_k , dann sind zwei Lösungen gegeben durch

$$\vec{\varphi}_{k,1}(t) = \vec{x}_k e^{\sqrt{\lambda_k} t} \quad \text{und} \quad \vec{\varphi}_{k,2}(t) = \vec{x}_k e^{-\sqrt{\lambda_k} t}.$$

Die Eigenwerte repräsentieren die **Eigenfrequenzen** des Systems und die Eigenvektoren die zu den Eigenfrequenzen gehörenden **Schwingungsformen** (Auslenkungen). Die allgemeine Schwingung ist eine Überlagerung dieser Grundschrwingungen. $\square$

12.3 Lineare Differenzialgleichungen n -ter Ordnung

12.3

Viele kinematische Probleme führen nach dem Newtonschen Bewegungsgesetz auf Differenzialgleichungen 2. Ordnung. Für ausgedehnte Körper sogar auf 4. oder höherer Ordnung. In diesem Abschnitt behandeln wir generell das systematische Lösen von linearen Differenzialgleichungen n -ter Ordnung.

Die theoretischen Grundlagen übertragen sich aus Abschnitt 12.2, da wir eine lineare Differenzialgleichung n -ter Ordnung auf ein System von n linearen Differenzialgleichungen 1. Ordnung reduzieren können. Gemäß der Lösungsstruktur lösen wir zuerst das homogene Problem und kommen dann auf die inhomogenen Differenzialgleichungen zu sprechen.

Wir beschränken uns in diesem Abschnitt generell auf lineare Differenzialgleichungen n -ter Ordnung mit *konstanten Koeffizienten*.

► 12.3.1 Einleitende Beispiele

Anwendungsbeispiel 12.24 (Fadenpendel).

An einem Faden der Länge l ist eine Masse m befestigt. Gesucht ist der Winkel $\varphi(t)$ als Funktion der Zeit, wenn die Masse um kleine Winkel φ_0 ausgelenkt wird. Die Kräfte, die auf die Masse m wirken, sind die Komponente der Gewichtskraft F_t senkrecht zum Faden

$$F_t = -F_G \sin \varphi = -m g \sin \varphi \approx -m g \varphi \quad (\text{kleine Winkel})$$

und die Reibungskraft F_R proportional zur Geschwindigkeit

$$F_R = -\gamma v = -\gamma (l \dot{\varphi}) .$$

Nach dem Newtonschen Bewegungsgesetz ist die Beschleunigungskraft,

$$F_B = m \ddot{x}(t) = m l \ddot{\varphi}(t) ,$$

gleich der Summe aller angreifenden Kräfte:

$$m l \ddot{\varphi}(t) = -m g \varphi(t) - \gamma l \dot{\varphi}(t)$$

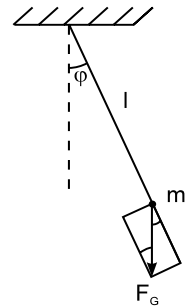


Abb. 12.17.
Fadenpendel

(homogene, lineare DG 2. Ordnung). $\square$

Anwendungsbeispiel 12.25 (Federpendel).

Am Ende einer senkrecht herabhängenden Feder mit der Federkonstanten D befindet sich eine Masse m . Es wirkt eine Reibungskraft proportional zur Momentangeschwindigkeit. Die Auslenkung der Masse m zur Zeit t wird mit $x(t)$ bezeichnet. Gesucht ist das Weg-Zeit-Gesetz $x(t)$, wenn die Masse m zum Zeitpunkt $t_0 = 0$ um x_0 aus der Ruhelage ausgelenkt wird.

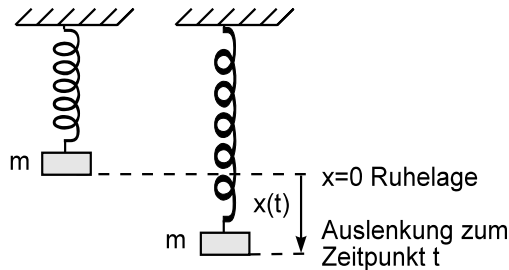


Abb. 12.18. Federpendel

Die Kräfte, die auf die Masse m wirken, sind die Federrückstellkraft $F_D = -D x(t)$ und die Reibungskraft $F_R = -\beta \dot{x}(t)$. Nach dem Newtonschen Bewegungsgesetz ist die Beschleunigungskraft $F_B = m \ddot{x}(t)$ gleich der Summe aller angreifenden Kräfte:

$$m \ddot{x}(t) = -\beta \dot{x}(t) - D x(t) \quad \text{mit } x(0) = x_0, \dot{x}(0) = 0$$

(homogene, lineare DG 2. Ordnung). $\square$

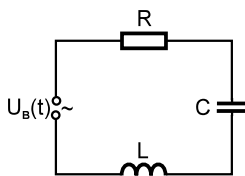
Anwendungsbeispiel 12.26 (RCL-Kreis).

Abb. 12.19. RCL-Kreis

Der RCL-Wechselstromkreis ist das elektromagnetische Analogon zu den Pendelbeispielen. Ein Stromkreis ist aufgebaut mit einer Induktivität L , einer Kapazität C und einem Ohmschen Widerstand R . Zur Zeit $t = 0$ wird der Stromkreis durch Anlegen einer äußeren Spannungsquelle

$$U_B(t) = U_0 \sin(\omega t)$$

geschlossen. Gesucht ist der Strom $I(t)$ als Funktion der Zeit.

Nach dem Maschensatz ist die Summe der Spannungsabfälle an R , C und L gleich der eingespeisten Spannung U_B

$$U_R(t) + U_L(t) + U_C(t) = U_B(t) .$$

Mit dem Ohmschen Gesetz ($U_R(t) = R \cdot I(t)$), dem Induktionsgesetz ($U_L(t) = L \frac{dI(t)}{dt}$) und der Spannung am Kondensator ($U_C(t) = \frac{1}{C} Q(t) = \frac{1}{C} \int_0^t I(\tau) d\tau$) gilt

$$R \cdot I(t) + L \frac{dI(t)}{dt} + \frac{1}{C} \int_0^t I(\tau) d\tau = U_0 \sin(\omega t)$$

bzw. nach Differenziation

$$L \ddot{I}(t) + R \dot{I}(t) + \frac{1}{C} I(t) = U_0 \omega \cos(\omega t)$$

(inhomogene, lineare DG 2. Ordnung). $\square$

⊙ Allgemeine Problemstellung:

Problemstellung: Sei $f(x)$ eine auf dem Intervall I stetige Funktion und seien $a_k \in \mathbb{R}$ ($k = 0, \dots, n-1$) reelle Koeffizienten. Gesucht ist eine n -mal stetig differenzierbare Funktion $y(x)$ mit der Eigenschaft, dass sie für alle $x \in I$ die lineare Differenzialgleichung

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = f(x) \quad (\text{DG } 2)$$

erfüllt.

Bezeichnung:

Für $f(x) \neq 0$ heißt (DG 2) eine **inhomogene DG n -ter Ordnung**;
für $f(x) = 0$ heißt (DG 2) eine **homogene DG n -ter Ordnung**.

Wir diskutieren im Folgenden die Problemstellungen: Wieviele Lösungen besitzt eine lineare Differenzialgleichung n -ter Ordnung? Wie löst man das homogene Problem, wie das inhomogene? Zur Beantwortung dieser Fragen übertragen wir zunächst die Ergebnisse aus dem Kapitel über LDGS 1. Ordnung auf lineare Differenzialgleichungen n -ter Ordnung:

➤ 12.3.2 Reduktion einer linearen DG n -ter Ordnung auf ein System

Ausgehend von der inhomogenen Differenzialgleichung n -ter Ordnung konstruieren wir ein System 1. Ordnung, indem wir - verallgemeinernd zum Vorgehen in Beispiel 12.16 - n Funktionen $y_0(x), y_1(x), \dots, y_{n-1}(x)$ einführen, die durch die Bestimmungsgleichungen

$$\begin{aligned} y_0(x) &:= y(x) \\ y_1(x) &:= y'(x) \\ y_2(x) &:= y''(x) \\ &\vdots \\ y_{n-1}(x) &:= y^{(n-1)}(x) \end{aligned}$$

festgelegt werden. Definieren wir die Vektorfunktion

$$\vec{Y}(x) := \begin{pmatrix} y_0(x) \\ y_1(x) \\ \vdots \\ y_{n-2}(x) \\ y_{n-1}(x) \end{pmatrix}$$

folgt für deren Ableitung

$$\begin{aligned} \vec{Y}'(x) &= \begin{pmatrix} y_0'(x) \\ y_1'(x) \\ \vdots \\ y_{n-2}'(x) \\ y_{n-1}'(x) \end{pmatrix} = \begin{pmatrix} y'(x) \\ y''(x) \\ \vdots \\ y^{(n-1)}(x) \\ y^{(n)}(x) \end{pmatrix} \\ &= \begin{pmatrix} y_1(x) \\ y_2(x) \\ \vdots \\ y_{n-1}(x) \\ -a_0 y(x) - a_1 y'(x) - \dots - a_{n-1} y^{(n-1)}(x) + f(x) \end{pmatrix}. \end{aligned}$$

Die Gleichheit der ersten $(n-1)$ Komponenten gilt nach Definition der Funktionen $y_i(x)$ und die der letzten Komponente aufgrund der Differenzialgleichung (DG2)

$$y^{(n)}(x) = -a_0 y(x) - a_1 y'(x) - \dots - a_{n-1} y^{(n-1)}(x) + f(x).$$

Ersetzen wir auch in der letzten Komponente die Ableitungen von $y(x)$ durch die entsprechenden Funktionen $y_i(x)$, ist

$$\vec{Y}'(x) = \begin{pmatrix} y_1(x) \\ y_2(x) \\ \vdots \\ y_{n-1}(x) \\ -a_0 y_0(x) - a_1 y_1(x) - \dots - a_{n-1} y_{n-1}(x) + f(x) \end{pmatrix}$$

$$= \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ 0 & 0 & 1 & \cdots & 0 \\ \vdots & \vdots & & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 1 \\ -a_0 & -a_1 & -a_2 & \cdots & -a_{n-1} \end{pmatrix} \vec{Y}(x) + \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ f(x) \end{pmatrix} \quad (\text{DGS 2})$$

Dies ist ein inhomogenes LDGS 1. Ordnung für $\vec{Y}(x)$.

Bemerkungen:

- (1) Ist $y(x)$ eine Lösung der inhomogenen, linearen Differenzialgleichung n -ter Ordnung (DG 2), dann ist $\vec{Y}(x) = (y(x), y'(x), \dots, y^{(n-1)}(x))^t$ eine Lösung des entsprechenden inhomogenen LDGS 1. Ordnung (DGS 2).
- (2) Es ist aber auch die Umkehrung gültig:
Ist $\vec{Y}(x) = (y_0(x), y_1(x), \dots, y_{n-1}(x))^t$ eine Lösung von (DGS 2), dann ist die erste Komponente des Vektors $\vec{Y}(x)$ nämlich $y(x) := y_0(x)$ eine Lösung der Differenzialgleichung n -ter Ordnung (DG 2).

Begründung: Ist $\vec{Y}(x)$ eine Lösung von (DGS 2), so gilt für die erste Komponente $y_0(x)$ nach (DGS 2):

$$\begin{aligned} y_0'(x) &= y_1(x) \\ y_0''(x) &= y_1'(x) = y_2(x) \\ &\vdots \\ y_0^{(n-1)}(x) &= y_{n-2}'(x) = y_{n-1}(x) \\ y_0^{(n)}(x) &= y_{n-1}'(x) \\ &= -a_0 y_0(x) - a_1 y_1(x) - \dots - a_{n-1} y_{n-1}(x) + f(x) \\ &= -a_0 y_0(x) - a_1 y_0'(x) - \dots - a_{n-1} y_0^{(n-1)}(x) + f(x). \end{aligned}$$

Somit ist $y_0(x)$ eine Lösung der Differenzialgleichung n -ter Ordnung. Die Umkehrung gilt aufgrund obiger Überlegungen und der Konstruktion des LDGS (DGS 2). $\square$

Beispiel 12.27. Gesucht ist das zur Differenzialgleichung 4. Ordnung

$$x''''(t) + 8x'''(t) + 22x''(t) + 24x'(t) + 9x(t) = 0$$

gehörende LDGS 1. Ordnung. Die Differenzialgleichung ist von der Ordnung 4, also führen wir 4 Funktionen

$$\begin{aligned} y_0(t) &= x(t) \\ y_1(t) &= x'(t) \\ y_2(t) &= x''(t) \\ y_3(t) &= x'''(t) \end{aligned}$$

ein. Dann gilt für die Ableitungen der Funktionen $y_0(t)$, $y_1(t)$, $y_2(t)$ und $y_3(t)$

$$\begin{aligned} y_0'(t) &= x'(t) = y_1(t) \\ y_1'(t) &= x''(t) = y_2(t) \\ y_2'(t) &= x'''(t) = y_3(t) \\ y_3'(t) &= x''''(t) = -9x(t) - 24x'(t) - 22x''(t) - 8x'''(t) \\ &= -9y_0(t) - 24y_1(t) - 22y_2(t) - 8y_3(t). \end{aligned}$$

Für den Vektor $\vec{Y}(t) := \begin{pmatrix} y_0(t) \\ y_1(t) \\ y_2(t) \\ y_3(t) \end{pmatrix}$ gilt dann

$$\begin{aligned} \vec{Y}'(t) &= \begin{pmatrix} y_0' \\ y_1' \\ y_2' \\ y_3' \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \\ y_3 \\ -9y_0 - 24y_1 - 22y_2 - 8y_3 \end{pmatrix} \\ &= \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ -9 & -24 & -22 & -8 \end{pmatrix} \vec{Y}(t). \end{aligned}$$

Dies ist das zur Differenzialgleichung gehörende LDGS 1. Ordnung. □

Satz 12.12: Das Lösen einer linearen Differenzialgleichung n -ter Ordnung (DG2) ist äquivalent zum Lösen des zugehörigen LDGS 1. Ordnung (DGS2).

Bemerkung: Nach Satz 12.12 ist es gleichgültig, ob die Differenzialgleichung n -ter Ordnung gelöst wird oder das zugehörige DG-System. Dieser Satz hat weitreichende Konsequenzen für das *numerische* Lösen von Differenzialgleichungen n -ter Ordnung: Statt eine Differenzialgleichung n -ter Ordnung zu lösen, geht man zum System 1. Ordnung über und löst jede dieser Differenzialgleichungen 1. Ordnung mit einem numerischen Verfahren wie z.B. dem Euler-Verfahren.

Hinweis: Auf der CD-Rom befindet sich ein separater Abschnitt über das [numerische Lösen von Differenzialgleichungen 1. Ordnung](#) sowie dem [Lösen von DG \$n\$ -ter Ordnung mit MAPLE](#).

Durch die in Satz 12.12 aufgestellte Äquivalenz übertragen sich auch die Aussagen über die Lösung von LDGS auf Differenzialgleichungen n -ter Ordnung.

Entsprechend Satz 12.1 und Satz 12.10 gilt

Satz 12.13: (Lösung linearer DG n -ter Ordnung)

- (1) Sei $\mathbb{L}_h$ die Menge aller Lösungen der *homogenen linearen Differenzialgleichung n -ter Ordnung*

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = 0.$$

$\mathbb{L}_h$ ist ein n -dimensionaler Vektorraum.

- (2) Sei $\mathbb{L}_i$ die Menge aller Lösungen der *inhomogenen linearen Differenzialgleichung n -ter Ordnung*

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = f(x).$$

Dann ist

$$\mathbb{L}_i = y_p(x) + \mathbb{L}_h,$$

wenn $y_p(x)$ eine beliebige *partikuläre* (= *spezielle*) Lösung der inhomogenen Differenzialgleichung ist.

- (3) n verschiedene Lösungen $\varphi_1(x), \varphi_2(x), \dots, \varphi_n(x)$ der homogenen Differenzialgleichung sind genau dann linear unabhängig, wenn für **ein**, und damit für alle $x \in I$, die sog. "**Wronski-Determinante**" $W(x)$ ungleich Null ist:

$$W(x) = \det \begin{pmatrix} \varphi_1(x) & \varphi_2(x) & \dots & \varphi_n(x) \\ \varphi_1'(x) & \varphi_2'(x) & \dots & \varphi_n'(x) \\ \vdots & \vdots & & \vdots \\ \varphi_1^{(n-1)}(x) & \varphi_2^{(n-1)}(x) & \dots & \varphi_n^{(n-1)}(x) \end{pmatrix} \neq 0.$$

Definition: Eine Basis $\varphi_1(x), \dots, \varphi_n(x)$ des Lösungsraumes $\mathbb{L}_h$ der homogenen Differenzialgleichung heißt **Lösungs-Fundamentalsystem**.

Folgerung: $\varphi_1(x), \dots, \varphi_n(x)$ ist ein **Lösungs-Fundamentalsystem**

genau dann, wenn

$$W(x_0) \neq 0$$

für ein $x_0 \in I$.

Beispiel 12.28. Gegeben ist für $x > 0$ die homogene DG 2. Ordnung

$$y''(x) - \frac{1}{2x} y'(x) + \frac{1}{2x^2} y(x) = 0 .$$

Zwei Lösungen sind

$$\varphi_1(x) = x, \quad \varphi_2(x) = \sqrt{x} ,$$

wie man durch Einsetzen in die Differenzialgleichung bestätigt. Die Wronski-Determinante zu φ_1, φ_2 lautet

$$W(x) = \det \begin{pmatrix} \varphi_1(x) & \varphi_2(x) \\ \varphi_1'(x) & \varphi_2'(x) \end{pmatrix} = \begin{vmatrix} x & \sqrt{x} \\ 1 & \frac{1}{2\sqrt{x}} \end{vmatrix} = -\frac{1}{2} \sqrt{x} .$$

Für $x > 0$ ist $W(x) \neq 0$ und $(\varphi_1(t), \varphi_2(t))$ bilden somit ein Fundamentalsystem. Die allgemeine Lösung der Differenzialgleichung ist daher

$$y(x) = c_1 x + c_2 \sqrt{x} .$$

Die Konstanten c_1 und c_2 bestimmen sich aus Anfangsbedingungen. □

Anwendungsbeispiel 12.29 (Elektron im Magnetfeld). Nach Beispiel 12.17 lauten die nicht-relativistischen Bewegungsgleichungen eines Elektrons in einem homogenen Magnetfeld, welches senkrecht zur Bewegungsrichtung steht,

$$\dot{v}_x(t) = -\omega v_y(t) \quad \text{und} \quad \dot{v}_y(t) = \omega v_x(t)$$

mit $\omega = \frac{e}{m} B \neq 0$. Differenziert man die erste Gleichung, $\ddot{v}_x(t) = -\omega \dot{v}_y(t)$, und setzt die zweite ein, erhält man eine Differenzialgleichung 2. Ordnung für die Geschwindigkeit $v_x(t)$:

$$\ddot{v}_x(t) + \omega^2 v_x(t) = 0 .$$

Zwei Lösungen dieser Differenzialgleichung kann man direkt angeben:

$$\varphi_1(t) = \cos(\omega t) \quad \text{und} \quad \varphi_2(t) = \sin(\omega t) ,$$

wie man durch Einsetzen in die Differenzialgleichung bestätigt! Diese beiden Lösungen bilden ein Fundamentalsystem, da die Wronski-Determinante

$$\begin{aligned} W(t) &= \det \begin{pmatrix} \varphi_1(t) & \varphi_2(t) \\ \varphi_1'(t) & \varphi_2'(t) \end{pmatrix} = \begin{vmatrix} \cos(\omega t) & \sin(\omega t) \\ -\omega \sin(\omega t) & \omega \cos(\omega t) \end{vmatrix} \\ &= \omega \cos^2(\omega t) + \omega \sin^2(\omega t) = \omega \neq 0 . \end{aligned}$$

Daher ist die allgemeine Lösung für $v_x(t)$

$$v_x(t) = c_1 \cos(\omega t) + c_2 \sin(\omega t) .$$

□

Im Folgenden werden wir uns mit der Frage beschäftigen, wie man alle Lösungen des homogenen Problems und eine spezielle Lösung des inhomogenen Problems berechnet: Die Lösung des homogenen Problems ist äquivalent zur Bestimmung von Nullstellen eines Polynoms n -ten Grades (*charakteristisches Polynom*) ($\rightarrow$ 12.3.3). Eine partikuläre Lösung der inhomogenen Differenzialgleichung kann man oftmals durch einen speziellen Ansatz gewinnen ($\rightarrow$ 12.3.4).

► 12.3.3 Homogene DG n -ter Ordnung mit konst. Koeffizienten

Beispiel 12.30. Gegeben ist die Differenzialgleichung 2. Ordnung

$$\ddot{x}(t) + \omega_0^2 x(t) = 0.$$

Die zugehörige physikalische Problemstellung kann z.B. das Fadenpendel ohne Reibung 12.24, das Federpendel 12.25, ein LC-Kreis 12.26 oder die Bewegungsgleichung eines Elektrons im Magnetfeld 12.29 sein.

Zur Lösung der Differenzialgleichung wählen wir den **Ansatz**:

$$x(t) = e^{\lambda t}. \quad (*)$$

Setzen wir diesen Ansatz in die Differenzialgleichung ein, folgt

$$\lambda^2 e^{\lambda t} + \omega_0^2 e^{\lambda t} = 0 \hookrightarrow \lambda^2 + \omega_0^2 = 0.$$

Man nennt

$$P(\lambda) = \lambda^2 + \omega_0^2$$

das zur Differenzialgleichung zugehörige *charakteristische Polynom*. Wenn die im Ansatz auftretende Größe λ eine Nullstelle des charakteristischen Polynoms ist, dann ist $e^{\lambda t}$ eine Lösung der Differenzialgleichung. Aus $P(\lambda) = 0$, folgt $\lambda = \pm \sqrt{-\omega_0^2} = \pm i\omega_0$.

$$\Rightarrow \varphi_1(t) = e^{i\omega_0 t} \quad \text{und} \quad \varphi_2(t) = e^{-i\omega_0 t}$$

sind Lösungen der Differenzialgleichung. Sie bilden gleichzeitig ein Fundamentalsystem, da die Wronski-Determinante

$$W(t) = \det \begin{pmatrix} \varphi_1(t) & \varphi_2(t) \\ \varphi_1'(t) & \varphi_2'(t) \end{pmatrix} = \begin{vmatrix} e^{i\omega_0 t} & e^{-i\omega_0 t} \\ i\omega_0 e^{i\omega_0 t} & -i\omega_0 e^{-i\omega_0 t} \end{vmatrix} = -2i\omega_0 \neq 0.$$

Da $\varphi_1(t)$, $\varphi_2(t)$ komplexe Funktionen sind, nennt man $(\varphi_1(t), \varphi_2(t))$ ein *komplexes Fundamentalsystem*.

② **Übergang zu einem reellen Fundamentalsystem.**

Zu diesem komplexen Fundamentalsystem konstruiert man ein reelles, indem man zu speziellen Linearkombinationen von $\varphi_1(t)$ und $\varphi_2(t)$ übergeht. Da für lineare DG das Superpositionsprinzip gilt, ist mit zwei Lösungen $\varphi_1(t)$ und $\varphi_2(t)$ jede Linearkombination $c_1 \varphi_1(t) + c_2 \varphi_2(t)$ ebenfalls eine Lösung der Differenzialgleichung. Mit $\varphi_1(t)$ und $\varphi_2(t)$ sind also auch die beiden Funktionen

$$x_1(t) = \frac{1}{2} \varphi_1(t) + \frac{1}{2} \varphi_2(t) = \frac{1}{2} (e^{i\omega_0 t} + e^{-i\omega_0 t}) = \cos(\omega_0 t)$$

$$x_2(t) = \frac{1}{2i} \varphi_1(t) - \frac{1}{2i} \varphi_2(t) = \frac{1}{2i} (e^{i\omega_0 t} - e^{-i\omega_0 t}) = \sin(\omega_0 t)$$

Lösungen der Differenzialgleichung. Da die Wronski-Determinante dieser beiden Funktionen $W(t) = \omega_0 \neq 0$, bilden $(\cos(\omega_0 t), \sin(\omega_0 t))$ ein *reelles Fundamentalsystem*. Die allgemeine Lösung lautet

$$x(t) = c_1 \cos(\omega_0 t) + c_2 \sin(\omega_0 t).$$

□

Wir übertragen die Lösungsmethode von Beispiel 12.30 auf den Fall einer allgemeinen, **homogenen linearen Differenzialgleichung n -ter Ordnung**:

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = 0. \quad (*)$$

Mit dem **Ansatz**

$$y(x) = e^{\lambda x}$$

für die gesuchte Funktion, lautet die k -te Ableitung von $y(x)$

$$y^{(k)}(x) = \lambda^k e^{\lambda x}.$$

Eingesetzt in die Differenzialgleichung (*) ergibt

$$\lambda^n e^{\lambda x} + a_{n-1} \lambda^{n-1} e^{\lambda x} + \dots + a_1 \lambda e^{\lambda x} + a_0 e^{\lambda x} = 0$$

$$\Rightarrow \lambda^n + a_{n-1} \lambda^{n-1} + \dots + a_1 \lambda + a_0 = 0.$$

Definition:

$$P(\lambda) := \lambda^n + a_{n-1} \lambda^{n-1} + \dots + a_1 \lambda + a_0$$

heißt das zur Differenzialgleichung (*) zugehörige **charakteristische Polynom**.

Ist λ_0 eine **Nullstelle** des charakteristischen Polynoms $P(\lambda)$, dann stellt

$$y(x) = e^{\lambda_0 x}$$

eine Lösung der Differenzialgleichung dar. Nach dem Fundamentalsatz der Algebra (5.2.7) besitzt jedes komplexe (also auch reelle) Polynom vom Grade n genau n komplexe Nullstellen $\lambda_1, \dots, \lambda_n$, die allerdings auch mehrfach vorkommen können. Hat das charakteristische Polynom n verschiedene Nullstellen, dann sind durch

$$y_k(x) = e^{\lambda_k x} \quad k = 1, \dots, n$$

n verschiedene Funktionen gegeben und es gilt

Satz 12.14: (Charakteristisches Polynom mit n verschiedenen Nullstellen). Gegeben ist die *homogene lineare Differenzialgleichung n -ter Ordnung*

$$y^{(n)}(x) + a_{n-1}y^{(n-1)}(x) + \dots + a_1y'(x) + a_0y(x) = 0.$$

Das zugehörige charakteristische Polynom $P(\lambda)$ habe n verschiedene Nullstellen $\lambda_1, \lambda_2, \dots, \lambda_n$. Dann bilden die n Lösungen der Differenzialgleichung

$$y_k(x) := e^{\lambda_k x} \quad (k = 1, \dots, n)$$

ein **Fundamentalsystem**.

Beweis: Aufgrund unserer Vorüberlegungen ist klar, dass $e^{\lambda_k x}$ ($k = 1, \dots, n$) Lösungen der DG sind, wenn die λ_k Nullstellen des charakteristischen Polynoms sind. Zu zeigen bleibt also nur noch die lineare Unabhängigkeit der Lösungen. Dazu gehen wir zur Wronski-Determinante über

$$W(x) = \det \begin{pmatrix} e^{\lambda_1 x} & e^{\lambda_2 x} & \dots & e^{\lambda_n x} \\ \lambda_1 e^{\lambda_1 x} & \lambda_2 e^{\lambda_2 x} & \dots & \lambda_n e^{\lambda_n x} \\ \vdots & \vdots & & \vdots \\ \lambda_1^{n-1} e^{\lambda_1 x} & \lambda_2^{n-1} e^{\lambda_2 x} & \dots & \lambda_n^{n-1} e^{\lambda_n x} \end{pmatrix}.$$

Mit vollständiger Induktion zeigt man, dass die sog. Vandermondesche Determinante

$$W(x=0) = \det \begin{pmatrix} 1 & 1 & \dots & 1 \\ \lambda_1 & \lambda_2 & \dots & \lambda_n \\ \vdots & \vdots & & \vdots \\ \lambda_1^{n-1} & \lambda_2^{n-1} & \dots & \lambda_n^{n-1} \end{pmatrix} = \prod_{i>j} (\lambda_i - \lambda_j) \neq 0.$$

Damit bilden die Lösungen auch ein Fundamentalsystem. □

Beispiel 12.31. Gesucht ist ein Fundamentalsystem der Differenzialgleichung

$$y^{(4)}(x) + 3y''(x) - 4y(x) = 0.$$

Ansatz: Wir setzen die Funktion $y(x) = e^{\lambda x}$ in die Differenzialgleichung ein. Dieses Vorgehen liefert das charakteristische Polynom

$$\lambda^4 e^{\lambda x} + 3\lambda^2 e^{\lambda x} - 4e^{\lambda x} = 0 \Rightarrow P(\lambda) = \lambda^4 + 3\lambda^2 - 4 \stackrel{!}{=} 0.$$

Die Nullstellen des charakteristischen Polynoms bestimmen wir mit der Substitution $Z := \lambda^2$. Dann ist $Z^2 + 3Z - 4 = 0 \hookrightarrow Z_1 = 1, Z_2 = -4$.

$$\Rightarrow \lambda_{1/2} = \pm\sqrt{Z_1} = \pm 1 \quad \text{und} \quad \lambda_{3/4} = \pm\sqrt{Z_2} = \pm\sqrt{-4} = \pm 2i.$$

$P(\lambda)$ hat somit 4 verschiedene Nullstellen $\pm 1, \pm 2i$.

$$\Rightarrow e^{1x}, \quad e^{-1x}, \quad e^{2ix}, \quad e^{-2ix}$$

ist ein komplexes Fundamentalsystem. Durch

$$\frac{1}{2} (e^{2ix} + e^{-2ix}) = \cos(2x)$$

$$\frac{1}{2i} (e^{2ix} - e^{-2ix}) = \sin(2x)$$

erhält man ein reelles Fundamentalsystem:

$$e^x, \quad e^{-x}, \quad \cos(2x), \quad \sin(2x). \quad \square$$

Nach Satz 12.14 ist eindeutig geklärt, wie man ein Fundamentalsystem bestimmt, wenn $P(\lambda)$ n verschiedene Nullstellen besitzt. Zur Klärung des Problems von doppelten bzw. mehrfachen Nullstellen betrachten wir das folgende Beispiel:

Beispiel 12.32. Gegeben ist die Differenzialgleichung

$$\ddot{x}(t) + 2\dot{x}(t) + x(t) = 0. \quad (*)$$

Ansatz: Setzen wir $x(t) = e^{\lambda t}$ in die Differenzialgleichung ein, liefert dies das charakteristische Polynom

$$P(\lambda) = \lambda^2 + 2\lambda + 1 \stackrel{!}{=} 0.$$

$$P(\lambda) = 0 \hookrightarrow \lambda_{1/2} = -1 \text{ ist doppelte Nullstelle}$$

$$\hookrightarrow x_1(t) = e^{-t} \text{ ist eine Lösung von } (*).$$

Mit dem Ansatz $e^{\lambda t}$ erhält man bei diesem Beispiel nur **eine** Lösung. Da (*) eine Differenzialgleichung 2. Ordnung, stellt $\mathbb{L}_h$ einen 2-dimensionalen Vektorraum dar und das Fundamentalsystem besteht aus **zwei** linear unabhängigen Funktionen! Eine weitere Lösung ist gegeben durch

$$x_2(t) = t \cdot e^{-t} ;$$

denn $x_2(t)$ und die Ableitungen

$$\begin{aligned}\dot{x}_2(t) &= e^{-t} - t e^{-t} \\ \ddot{x}_2(t) &= -e^{-t} - e^{-t} + t e^{-t}\end{aligned}$$

in die Differenzialgleichung eingesetzt

$$\Rightarrow \ddot{x}_2(t) + 2\dot{x}_2(t) + x_2(t) = 0.$$

Außerdem sind $x_1(t)$ und $x_2(t)$ linear unabhängig:

$$W(t) = \det \begin{pmatrix} x_1(t) & x_2(t) \\ x_1'(t) & x_2'(t) \end{pmatrix} = \begin{vmatrix} e^{-t} & t e^{-t} \\ -e^{-t} & e^{-t}(1-t) \end{vmatrix} = e^{-2t} \neq 0.$$

Daher ist ein Fundamentalsystem

$$e^{-t}, \quad t e^{-t}. \quad \square$$

Ist λ_0 eine m -fache Nullstelle des charakteristischen Polynoms $P(\lambda)$, dann sind

$$e^{\lambda_0 t}, \quad t e^{\lambda_0 t}, \quad t^2 e^{\lambda_0 t}, \dots, t^{m-1} e^{\lambda_0 t}$$

linear unabhängige Lösungen der Differenzialgleichung n -ter Ordnung:

Satz 12.15: (Charakteristisches Polynom mit Mehrfachnullstellen). Gegeben ist die *homogene lineare Differenzialgleichung n -ter Ordnung*

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = 0.$$

Das zugehörige charakteristische Polynom $P(\lambda)$ habe l verschiedene Nullstellen $\lambda_k \in \mathbb{C}$ ($k = 1, \dots, l$) mit der Vielfachheit m_k ($k = 1, \dots, l$). Dann sind

$$e^{\lambda_k x}, x e^{\lambda_k x}, \dots, x^{m_k-1} e^{\lambda_k x}$$

linear unabhängige Lösungen der Differenzialgleichung und bilden für $k = 1, \dots, l$ ein **Fundamentalsystem**.

Beispiel 12.33. Gesucht ist ein reelles Fundamentalsystem der Differenzialgleichung

$$y^{(4)}(x) + 8y''(x) + 16y(x) = 0.$$

Ansatz: $y(x) = e^{\lambda x}$ in die Differenzialgleichung eingesetzt liefert das charakteristische Polynom

$$P(\lambda) = \lambda^4 + 8\lambda^2 + 16 = 0.$$

Mit $Z := \lambda^2$ ist $Z^2 + 8Z + 16 = 0 \Leftrightarrow Z_{1/2} = -4$ doppelt. Damit sind

$$\lambda_{1/2} = \pm\sqrt{-4} = \pm 2i$$

doppelte Nullstellen.

$$\lambda_1 = 2i \quad \hookrightarrow \quad \varphi_1(x) = e^{2ix}, \quad \varphi_2(x) = x \cdot e^{2ix} \quad \text{sind zwei Lösungen.}$$

$$\lambda_2 = -2i \quad \hookrightarrow \quad \varphi_3(x) = e^{-2ix}, \quad \varphi_4(x) = x \cdot e^{-2ix} \quad \text{sind zwei Lösungen.}$$

$\Rightarrow e^{2ix}, e^{-2ix}, xe^{2ix}, xe^{-2ix}$ ist ein komplexes Fundamentalsystem.

Übergang zum reellen Fundamentalsystem durch spezielle Linearkombinationen:

$$\frac{1}{2} (\varphi_1(x) + \varphi_3(x)) = \frac{1}{2} (e^{2ix} + e^{-2ix}) = \cos(2x)$$

$$\frac{1}{2i} (\varphi_1(x) - \varphi_3(x)) = \frac{1}{2i} (e^{2ix} - e^{-2ix}) = \sin(2x)$$

$$\frac{1}{2} (\varphi_2(x) + \varphi_4(x)) = x \frac{1}{2} (e^{2ix} + e^{-2ix}) = x \cdot \cos(2x)$$

$$\frac{1}{2i} (\varphi_2(x) - \varphi_4(x)) = x \frac{1}{2i} (e^{2ix} - e^{-2ix}) = x \cdot \sin(2x).$$

$$\Rightarrow \cos(2x), \sin(2x), x \cos(2x), x \sin(2x)$$

bildet ein reelles Fundamentalsystem. Die allgemeine Lösung lautet somit:

$$y(x) = c_1 \cos(2x) + c_2 \sin(2x) + c_3 x \cos(2x) + c_4 x \sin(2x). \quad \square$$

Beispiel 12.34. Gegeben ist die Differenzialgleichung

$$y'''(x) - y(x) = 0.$$

Gesucht ist ein reelles Fundamentalsystem. Mit dem Ansatz $y(x) = e^{\lambda x}$ in die Differenzialgleichung eingesetzt, erhält man das charakteristische Polynom

$$P(\lambda) = \lambda^3 - 1 \stackrel{!}{=} 0.$$

Die Nullstellen des charakteristischen Polynoms sind $\lambda_1 = 1$ und $\lambda_{2/3} = -\frac{1}{2} \pm \frac{1}{2} \sqrt{3}i$, so dass die Funktionen

$$e^x, \quad e^{\left(-\frac{1}{2} + \frac{1}{2} \sqrt{3}i\right)x}, \quad e^{\left(-\frac{1}{2} - \frac{1}{2} \sqrt{3}i\right)x}$$

ein komplexes Fundamentalsystem bilden. Mit den Linearkombinationen

$$\begin{aligned} \frac{1}{2} \left(e^{\left(-\frac{1}{2} + \frac{1}{2} \sqrt{3}i\right)x} + e^{\left(-\frac{1}{2} - \frac{1}{2} \sqrt{3}i\right)x} \right) &= e^{-\frac{1}{2}x} \frac{1}{2} \left(e^{\frac{1}{2} \sqrt{3}ix} + e^{-\frac{1}{2} \sqrt{3}ix} \right) \\ &= e^{-\frac{1}{2}x} \cos\left(\frac{1}{2} \sqrt{3}x\right) \end{aligned}$$

und

$$\begin{aligned} \frac{1}{2i} \left(e^{\left(-\frac{1}{2} + \frac{1}{2} \sqrt{3}i\right)x} - e^{\left(-\frac{1}{2} - \frac{1}{2} \sqrt{3}i\right)x} \right) &= e^{-\frac{1}{2}x} \frac{1}{2i} \left(e^{\frac{1}{2} \sqrt{3}ix} - e^{-\frac{1}{2} \sqrt{3}ix} \right) \\ &= e^{-\frac{1}{2}x} \sin\left(\frac{1}{2} \sqrt{3}x\right) \end{aligned}$$

bekommt man ein reelles Fundamentalsystem

$$e^x, \quad e^{-\frac{1}{2}x} \cos\left(\frac{1}{2} \sqrt{3}x\right), \quad e^{-\frac{1}{2}x} \sin\left(\frac{1}{2} \sqrt{3}x\right). \quad \square$$

Anwendungsbeispiel 12.35 (Freie gedämpfte Schwingung, [mit MAPLE-Worksheet](#)).

Wir kommen auf das Federpendel aus Beispiel 12.25 zurück. Für die Auslenkung $x(t)$ der Masse m aus der Ruhelage gilt unter Berücksichtigung von Reibung die Differenzialgleichung

$$m \ddot{x}(t) = -\beta \dot{x}(t) - D x(t) \quad \text{mit} \quad x(0) = x_0 \quad \text{und} \quad \dot{x}(0) = 0.$$

Mit den Parametern $\omega_0^2 = \frac{D}{m}$ und $\mu = \frac{1}{2} \frac{\beta}{m}$ ist

$$\ddot{x}(t) + 2\mu \dot{x}(t) + \omega_0^2 x(t) = 0, \quad x(0) = x_0, \quad \dot{x}(0) = 0.$$

Der Ansatz

$$x(t) = e^{\lambda t}$$

führt zum charakteristischen Polynom

$$P(\lambda) = \lambda^2 + 2\mu\lambda + \omega_0^2 = 0$$

mit den Nullstellen

$$\lambda_{1/2} = -\mu \pm \sqrt{\mu^2 - \omega_0^2}.$$

Das Vorzeichen der Diskriminante

$$\Delta := \mu^2 - \omega_0^2$$

entscheidet über die Art der Schwingung. Bei schwacher Dämpfung ist das mechanische System zu echten Schwingungen fähig (*Schwingungsfall*). Dieser Fall tritt ein, wenn $\mu < \omega_0$. Bei starker Dämpfung $\mu > \omega_0$ bewegt sich das System nicht-periodisch (= *aperiodisch*) auf die Gleichgewichtslage zu (*Kriechfall*). Für $\Delta = 0$, d.h. $\mu = \omega_0$, folgt der aperiodische Grenzfall. Im Folgenden werden wir jeden dieser drei Fälle getrennt behandeln:

1. Fall: $\Delta < 0$, d.h. $\mu < \omega_0$: *Gedämpfte Schwingung*
2. Fall: $\Delta = 0$, d.h. $\mu = \omega_0$: *Aperiodischer Grenzfall*
3. Fall: $\Delta > 0$, d.h. $\mu > \omega_0$: *Kriechfall*

⊗ 1. Gedämpfte Schwingung (schwache Dämpfung).

Bei schwacher Dämpfung ($\mu < \omega_0$) sind die Nullstellen des charakteristischen Polynoms komplex-konjugierte Zahlen

$$\lambda_{1/2} = -\mu \pm \sqrt{\mu^2 - \omega_0^2} = -\mu \pm i \sqrt{\omega_0^2 - \mu^2} = -\mu \pm i \omega$$

mit $\omega := \sqrt{\omega_0^2 - \mu^2} > 0$. Damit ist ein komplexes Fundamentalsystem

$$\varphi_1(t) = e^{\lambda_1 t} = e^{(-\mu + i\omega)t} = e^{-\mu t} e^{i\omega t}$$

$$\varphi_2(t) = e^{\lambda_2 t} = e^{(-\mu - i\omega)t} = e^{-\mu t} e^{-i\omega t}.$$

Mit dem reellen Fundamentalsystem

$$x_1(t) = \frac{1}{2} (\varphi_1(t) + \varphi_2(t)) = e^{-\mu t} \cos(\omega t)$$

$$x_2(t) = \frac{1}{2i} (\varphi_1(t) - \varphi_2(t)) = e^{-\mu t} \sin(\omega t)$$

bestimmt sich die **allgemeine Lösung** durch

$$x(t) = e^{-\mu t} (c_1 \cos(\omega t) + c_2 \sin(\omega t)).$$

Die Konstanten c_1, c_2 bestimmen sich aus den Anfangsbedingungen $x(0)$ und $\dot{x}(0)$. Um die Anfangswerte für $\dot{x}(0)$ einsetzen zu können, benötigen wir die Ableitung von $x(t)$:

$$\begin{aligned} \dot{x}(t) &= -\mu e^{-\mu t} (c_1 \cos(\omega t) + c_2 \sin(\omega t)) \\ &\quad + e^{-\mu t} (-c_1 \omega \sin(\omega t) + c_2 \omega \cos(\omega t)). \end{aligned}$$

Setzen wir die Anfangsbedingungen ein, folgen zwei Bestimmungsgleichungen für die Konstanten c_1 und c_2 :

$$x(0) = x_0: \quad x(0) = c_1 = x_0,$$

$$\dot{x}(0) = 0: \quad \dot{x}(0) = c_1(-\mu) + c_2\omega = 0 \Rightarrow c_2 = x_0 \left(\frac{\mu}{\omega}\right).$$

$$\Rightarrow \quad x(t) = x_0 e^{-\mu t} \left(\cos(\omega t) + \frac{\beta}{2m} \frac{1}{\omega} \sin(\omega t) \right).$$

Interpretation: Die Lösung setzt sich aus einer zeitlich abnehmenden Amplitude $x_0 e^{-\mu t}$ und einem periodischen Anteil $(\cos(\omega t) + \frac{\beta}{2m} \frac{1}{\omega} \sin(\omega t))$ zusammen. Der periodische Anteil der Funktion kann in der Form $(\frac{\omega_0}{\omega} \sin(\omega t + \varphi))$ mit $\tan \varphi = \frac{\omega}{\mu}$ dargestellt werden, so dass

$$x(t) = x_0 e^{-\mu t} \frac{\omega_0}{\omega} \sin(\omega t + \varphi).$$

Es liegt eine *gedämpfte Schwingung* vor. Das Federpendel schwingt mit der gegenüber der ungedämpften Schwingung **verkleinerten** Kreisfrequenz

$$\omega = \sqrt{\omega_0^2 - \mu^2} < \omega_0.$$

Abb. 12.20 zeigt den typischen Verlauf einer gedämpften Schwingung.

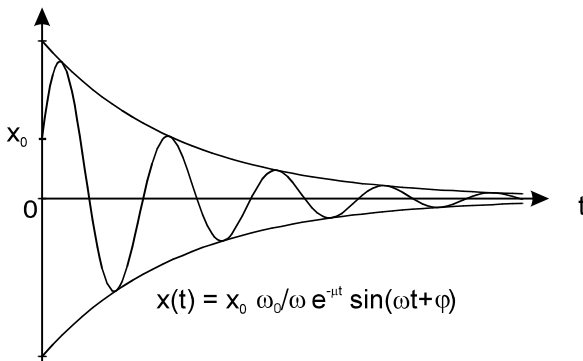


Abb. 12.20. Zeitlicher Verlauf einer gedämpften Schwingung

➤ 2. Aperiodischer Grenzfall.

$\Delta = 0$, d.h. $\mu = \omega_0$, beschreibt den aperiodischen Grenzfall, der die periodischen Bewegungen von den nicht-periodischen trennt. Für $\mu = \omega_0$ ist

$$\lambda_{1/2} = -\mu$$

eine **doppelte** Nullstelle des charakteristischen Polynoms $P(\lambda)$, und $\varphi_1(t) = e^{-\mu t}$ und $\varphi_2(t) = t \cdot e^{-\mu t}$ bilden ein reelles Fundamentalsystem. Die **allge-**

meine Lösung lautet dann

$$x(t) = c_1 e^{-\mu t} + c_2 t e^{-\mu t}.$$

Durch Einsetzen der Anfangsbedingungen bestimmen sich c_1, c_2 :

$$x(0) = x_0: \quad c_1 = x_0,$$

$$\dot{x}(0) = 0: \quad -\mu c_1 + c_2 = 0 \Rightarrow c_2 = \mu x_0.$$

$$\Rightarrow \quad x(t) = x_0 e^{-\mu t} (1 + \mu t).$$

Der Massepunkt bewegt sich nach seiner Auslenkung um x_0 auf die Gleichgewichtslage nicht-periodisch zu.

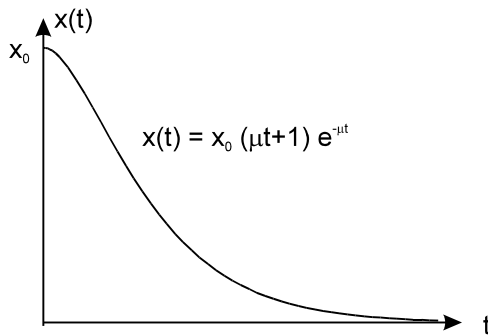


Abb. 12.21. Aperiodischer Grenzfall für $x(0) = x_0, x'(0) = 0$

③ 3. Kriechfall.

Bei *starker Dämpfung* ($\mu > \omega_0$) sind die Nullstellen des charakteristischen Polynoms zwei negative, reelle Zahlen

$$\lambda_{1/2} = -\mu \pm \sqrt{\mu^2 - \omega_0^2} < 0.$$

Setzen wir $k = \sqrt{\mu^2 - \omega_0^2}$, bilden $\varphi_1(t) = e^{(-\mu+k)t}$ und $\varphi_2(t) = e^{(-\mu-k)t}$ ein reelles Fundamentalsystem und die **allgemeine Lösung** lautet

$$x(t) = c_1 e^{(-\mu+k)t} + c_2 e^{(-\mu-k)t}.$$

Die Masse ist infolge zu starker Reibung zu keiner echten Schwingung fähig und bewegt sich im Lauf der Zeit nicht-periodisch auf die Gleichgewichtslage zu. Man bezeichnet diesen Fall in der Mechanik als *Kriechfall* oder als aperiodische Schwingung. Der genaue Verlauf hängt, wie im Fall 2, von den Anfangsbedingungen ab.

Für $x(0) = x_0$ und $\dot{x}(0) = 0$ bestimmen sich die Konstanten c_1 und c_2 zu:

$$x(0) = x_0: \quad x_0 = c_1 + c_2,$$

$$\dot{x}(0) = 0: \quad 0 = (-\mu + k) c_1 + (-\mu - k) c_2.$$

Die Lösung dieses linearen Gleichungssystems für c_1 und c_2 lautet $c_1 = x_0 \frac{k + \mu}{2k}$ und $c_2 = x_0 \frac{k - \mu}{2k}$

$$\Rightarrow x(t) = \frac{x_0}{2k} e^{-\mu t} ((k + \mu) e^{kt} + (k - \mu) e^{-kt})$$

(Exponentielles Abklingen ohne Schwingungsanteil).



Hinweis: Auf der CD-Rom befindet sich ein [MAPLE-Worksheet](#), in dem eine Animation aufzeigt wie sich das Schwingungsverhalten in Abhängigkeit von μ ändert. Die Animation beginnt bei schwacher Dämpfung und nähert sich dann dem aperiodischen Grenzfall an. Man erkennt, dass beim aperiodischen Grenzfall das System am schnellsten zur Ruhe kommt, wie dies z.B. bei Stoßdämpfern oder Messinstrumenten gefordert wird. Wird die Dämpfung noch kleiner, erhält man den Schwingungsfall; die Periode der Schwingung ändert sich dann in Abhängigkeit von μ . $\square$

► 12.3.4 Inhomogene DG n -ter Ord. mit konstanten Koeffizienten

In diesem Abschnitt betrachten wir das inhomogene Problem: Gegeben ist eine lineare DG n -ter Ordnung mit Störfunktion $f(x)$ (= Inhomogenität):

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = f(x). \quad (**)$$

Gesucht ist eine *partikuläre Lösung* $y_p(x)$.

Über das zugehörige LDG-System 1. Ordnung und Variation der Konstanten besitzt man für beliebige Inhomogenitäten eine Lösungsformel: Nachdem die Nullstellen des charakteristischen Polynoms bestimmt sind, bildet man nach Satz 12.15 ein Lösungsfundamentalsystem der homogenen linearen Differenzialgleichung. Anschließend transformiert man die lineare Differenzialgleichung n -ter Ordnung in ein System 1. Ordnung (Satz 12.12) und führt die Variation der Konstanten (Satz 12.11) durch. Dieser Weg bietet sich an, wenn man z.B. auf MAPLE zurückgreift, um die umfangreichen Integrale aufzustellen und berechnen zu lassen.

⊗ **Inhomogenität = Exponentialfunktion**

Im Folgenden werden wir aber für oft auftretende Inhomogenitäten eine partikuläre Lösung durch einen speziellen Lösungsansatz vorgeben. Zunächst betrachten wir den Fall, dass die Störfunktion eine reine Exponentialfunktion ist:

$$f(x) = c e^{\mu x} \quad \mu \in \mathbb{C}.$$

Dann lautet die Differenzialgleichung

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = c e^{\mu x}.$$

Gesucht ist eine partikuläre Lösung der Form

$$y_p(x) = k e^{\mu x}$$

mit einer noch unbekannten Konstanten k . Setzen wir $y_p(x)$ zusammen mit seinen Ableitungen in die Differenzialgleichung ein, folgt

$$k \mu^n e^{\mu x} + a_{n-1} k \mu^{n-1} e^{\mu x} + \dots + a_1 k \mu e^{\mu x} + a_0 k e^{\mu x} = c e^{\mu x}$$

$$\Rightarrow k \underbrace{(\mu^n + a_{n-1} \mu^{n-1} + \dots + a_1 \mu + a_0)}_{P(\mu)} = c$$

$$\Rightarrow k P(\mu) = c,$$

wenn das charakteristische Polynom $P(\lambda)$ an der Stelle μ ausgewertet wird. Ist μ **keine** Nullstelle des charakteristischen Polynoms, $P(\mu) \neq 0$, folgt für die Konstante $k = \frac{c}{P(\mu)}$ und die Lösung lautet $y_p(x) = \frac{c}{P(\mu)} \cdot e^{\mu x}$:

Satz 12.16: Gegeben ist die inhomogene, lineare Differenzialgleichung n -ter Ordnung

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = c e^{\mu x}.$$

Ist μ **keine** Nullstelle des charakteristischen Polynoms $P(\lambda)$, dann ist

$$y_p(x) = \frac{c}{P(\mu)} e^{\mu x}$$

eine **partikuläre Lösung**.

Beispiele 12.36:

- ① Gesucht ist eine partikuläre Lösung
- $y_p(x)$
- der Differenzialgleichung

$$y^{(4)}(x) + 2y''(x) + y(x) \stackrel{!}{=} 25e^{2x}.$$

Das zur Differenzialgleichung gehörende charakteristische Polynom ist

$$P(\lambda) = \lambda^4 + 2\lambda^2 + 1$$

und $\mu = 2$ ist keine Nullstelle von $P(\lambda)$: $P(2) = 25 \neq 0$. Somit erhält man durch

$$y_p(x) = k e^{2x}$$

eine partikuläre Lösung. Setzt man $y_p(x)$ in die Differenzialgleichung ein, bestimmt sich die Konstante k aus:

$$k 16 e^{2x} + k 8 e^{2x} + k e^{2x} = 25 e^{2x}$$

$$\hookrightarrow k = \frac{25}{25} = 1 \Rightarrow \boxed{y_p(x) = e^{2x}}.$$

- ② Gesucht ist eine partikuläre Lösung
- $y_p(x)$
- der Differenzialgleichung

$$y^{(4)}(x) + 2y''(x) + y(x) = 25e^{i2x}.$$

Das charakteristische Polynom ist

$$P(\lambda) = \lambda^4 + 2\lambda^2 + 1.$$

Wegen $\mu = 2i$ und $P(2i) = 16i^4 + 8i^2 + 1 = 9 \neq 0$ ist μ keine Nullstelle von $P(\lambda)$ und somit ist eine partikuläre Lösung

$$y_p(x) = \frac{25}{P(2i)} e^{i2x} = \frac{25}{9} e^{i2x}.$$

□

Beispiel 12.37. Gegeben ist die Differenzialgleichung

$$\boxed{y'''(x) - 2y''(x) - 2y'(x) + 2y(x) = 2 \sin x.} \quad (*)$$

Um eine partikuläre Lösung $y_p(x)$ nach Satz 12.16 zu berechnen, setzen wir die Differenzialgleichung ins Komplexe fort:

$$\boxed{\tilde{y}'''(x) - 2\tilde{y}''(x) - 2\tilde{y}'(x) + 2\tilde{y}(x) = 2e^{ix}.} \quad (\tilde{*})$$

Ist $\tilde{y}_p(x)$ eine Lösung der komplexen Differenzialgleichung $(\tilde{*})$, dann ist der Imaginärteil

$$\boxed{y_p(x) := \operatorname{Im} \tilde{y}_p(x)}$$

eine Lösung der reellen Differenzialgleichung $(*)$, wie man folgendermaßen erkennt

$$\begin{aligned}
& \tilde{y}_p''' - 2\tilde{y}_p'' - 2\tilde{y}_p' + 2\tilde{y}_p = 2e^{ix} \\
& \hookrightarrow \operatorname{Im} (\tilde{y}_p''' - 2\tilde{y}_p'' - 2\tilde{y}_p' + 2\tilde{y}_p) = \operatorname{Im} (2e^{ix}) \\
& \hookrightarrow \operatorname{Im} (\tilde{y}_p''') - 2 \operatorname{Im} (\tilde{y}_p'') - 2 \operatorname{Im} (\tilde{y}_p') + 2 \operatorname{Im} (\tilde{y}_p) = 2 \operatorname{Im} (e^{ix}) \\
& \hookrightarrow [\operatorname{Im} (\tilde{y}_p)]''' - 2 [\operatorname{Im} (\tilde{y}_p)]'' - 2 [\operatorname{Im} (\tilde{y}_p)]' + 2 [\operatorname{Im} (\tilde{y}_p)] = 2 \sin x. \\
& \Rightarrow y_p''' - 2y_p'' - 2y_p' + 2y_p = 2 \sin x.
\end{aligned}$$

Um $(*)$ zu lösen, wählen wir den Ansatz

$$\tilde{y}_p(x) = k e^{ix}.$$

In die Differenzialgleichung eingesetzt

$$\begin{aligned}
& \hookrightarrow k i^3 e^{ix} - k 2 i^2 e^{ix} - k 2 i e^{ix} + k 2 e^{ix} = 2 e^{ix} \\
& \hookrightarrow k (4 - 3i) e^{ix} = 2 e^{ix} \hookrightarrow k = \frac{2}{4 - 3i}. \\
& \Rightarrow \tilde{y}_p(x) = \frac{2}{4 - 3i} e^{ix}.
\end{aligned}$$

Übergang ins Reelle: Damit ist eine partikuläre Lösung $\tilde{y}_p(x)$ von $(*)$ gefunden. Die gesuchte Lösung $y_p(x)$ von $(*)$ ist

$$y_p(x) = \operatorname{Im} (\tilde{y}_p(x)) = \operatorname{Im} \left(\frac{2}{4 - 3i} e^{ix} \right).$$

Es gibt zwei unterschiedliche Methoden, um den Imaginärteil zu berechnen. Beide führen auf eine unterschiedliche Darstellung der Lösung. Im ersten Fall zerlegen wir sowohl $\frac{2}{4-3i}$ als auch e^{ix} in Real- und Imaginärteil, bestimmen in der algebraischen Normalform das Produkt der beiden komplexen Größen und lesen vom Ergebnis den Imaginärteil ab. Im zweiten Fall stellen wir $\frac{2}{4-3i}$ in der Exponentialform dar und multiplizieren mit e^{ix} in der Exponentialform; die partikuläre Lösung ist wieder der Imaginärteil des Ergebnisses.

(i) Zerlegung von $\frac{2}{4-3i}$ in Real- und Imaginärteil

$$\begin{aligned}
\frac{2}{4 - 3i} &= \frac{2}{4 - 3i} \cdot \frac{4 + 3i}{4 + 3i} = \frac{8}{25} + \frac{6}{25} i. \\
\Rightarrow \tilde{y}_p(x) &= \frac{2}{4 - 3i} e^{ix} = \left(\frac{8}{25} + \frac{6}{25} i \right) (\cos x + i \sin x) \\
&= \left(\frac{8}{25} \cos x - \frac{6}{25} \sin x \right) + i \left(\frac{6}{25} \cos x + \frac{8}{25} \sin x \right).
\end{aligned}$$

Damit ist

$$y_p(x) = \operatorname{Im} (\tilde{y}_p(x)) = \frac{6}{25} \cos x + \frac{8}{25} \sin x.$$

- (ii) Die komplexe Zahl $c = \frac{2}{4-3i} = \frac{8}{25} + \frac{6}{25}i$ lässt sich darstellen in der Exponentialform $c = |c| e^{i\varphi}$ mit

$$|c| = \frac{1}{25} \sqrt{8^2 + 6^2} = \frac{10}{25} \quad \text{und} \quad \tan \varphi = \frac{3}{4} \hookrightarrow \varphi = 36.9^\circ \Rightarrow c = \frac{10}{25} e^{i 36.9^\circ}$$

$$\begin{aligned} \Rightarrow \tilde{y}_p(x) &= \frac{2}{4-3i} e^{ix} = \frac{10}{25} e^{i 36.9^\circ} \cdot e^{ix} = \frac{10}{25} e^{i(x+36.9^\circ)} \\ &= \frac{10}{25} \cos(x + 36.9^\circ) + i \frac{10}{25} \sin(x + 36.9^\circ) . \end{aligned}$$

Damit ist

$$y_p(x) = \operatorname{Im}(\tilde{y}_p(x)) = \frac{10}{25} \sin(x + 36.9^\circ) . \quad \square$$

⊗ Inhomogenität = Polynom mal Exponentialfunktion

Der Ansatz für die spezielle Lösung aus Satz 12.16 führt zum Ziel, wenn μ **keine** Nullstelle des charakteristischen Polynoms $P(\lambda)$ ist. Welcher Ansatz muss aber gewählt werden, wenn μ eine Nullstelle ist? Allgemeiner noch betrachten wir den Fall einer Inhomogenität $f(x) = h(x) e^{\mu x}$ mit einem Polynom $h(x)$:

Satz 12.17: Gegeben ist die *inhomogene* lineare Differenzialgleichung

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = h(x) e^{\mu x} .$$

(i) Ist μ eine k -fache ($k \geq 0$) Nullstelle des charakteristischen Polynoms

(ii) und $h(x)$ ein Polynom vom Grade m ,

dann liefert der Ansatz

$$y_p(x) = g(x) \cdot e^{\mu x}$$

eine spezielle Lösung, wenn $g(x)$ ein Polynom vom Grade $m + k$ ist.

Bemerkungen:

- (1) Satz 12.16 ist ein Spezialfall von Satz 12.17: Für $f(x) = c e^{\mu x}$ und μ keine Nullstelle des charakteristischen Polynoms ist $k = 0$ und $m = 0$ (c ist ein Polynom vom Grade 0). Daher liefert die Ansatzfunktion

$$y_p(x) = K e^{\mu x}$$

mit einem Polynom vom Grade $k + m = 0$ eine partikuläre Lösung.

- (2) Besteht die Störfunktion aus mehreren Störgliedern, erhält man einen Ansatz für eine partikuläre Lösung $y_p(x)$ als Summe der Ansätze für die einzelnen Störglieder.

Beispiele 12.38 (Musterbeispiele):

- ① Gesucht ist eine partikuläre Lösung der Differenzialgleichung

$$2y''(x) + y'(x) = xe^{-x} :$$

Das zugehörige charakteristische Polynom ist

$$P(\lambda) = 2\lambda^2 + \lambda.$$

Die Inhomogenität ist

$$f(x) = xe^{-x} \hookrightarrow \mu = -1.$$

$\mu = -1$ ist keine Nullstelle von $P(\lambda)$, da $P(-1) = 1 \neq 0 \hookrightarrow k = 0$; x ist ein Polynom vom Grade 1 $\hookrightarrow m = 1. \Rightarrow k + m = 1$ und die Ansatzfunktion für eine partikuläre Lösung ist ein Polynom vom Grad 1 mal e^{-x} :

$$y_p(x) = (a_0 + a_1 x) e^{-x}.$$

Wir setzen diesen Ansatz für $y_p(x)$ zusammen mit den Ableitungen

$$\begin{aligned} y_p'(x) &= a_1 e^{-x} - (a_0 + a_1 x) e^{-x} \\ y_p''(x) &= -2a_1 e^{-x} + (a_0 + a_1 x) e^{-x} \end{aligned}$$

in die Differenzialgleichung ein, um a_0 und a_1 zu bestimmen. Damit folgt

$$2y_p'' + y_p'(x) = [(a_0 - 3a_1) + a_1 x] e^{-x} \stackrel{!}{=} x e^{-x}$$

$$\Rightarrow (a_0 - 3a_1) + a_1 x \stackrel{!}{=} x$$

Um die Koeffizienten zu bestimmen, führen wir einen Koeffizientenvergleich nach absteigenden Potenzen von x durch

$$\left. \begin{array}{l} x^1 : a_1 = 1 \\ x^0 : a_0 - 3a_1 = 0 \Rightarrow a_0 = 3. \end{array} \right\} \Rightarrow y_p(x) = (3 + x) e^{-x}.$$

- ② Gesucht ist eine partikuläre Lösung der Differenzialgleichung

$$2y''(x) + y'(x) = x :$$

Das zugehörige charakteristische Polynom ist

$$P(\lambda) = 2\lambda^2 + \lambda.$$

Die Inhomogenität ist

$$f(x) = x e^{0x} \hookrightarrow \mu = 0.$$

$\mu = 0$ ist einfache Nullstelle von $P(\lambda) \hookrightarrow k = 1$; x ist ein Polynom vom Grade $1 \hookrightarrow m = 1. \Rightarrow m + k = 2$. Die Ansatzfunktion für eine partikuläre Lösung lautet

$$y_p(x) = (a_0 + a_1 x + a_2 x^2) e^{0x} = a_0 + a_1 x + a_2 x^2.$$

Um die Koeffizienten a_0 , a_1 und a_2 zu bestimmen, setzen wir $y_p(x)$ zusammen mit den Ableitungen

$$\begin{aligned} y_p'(x) &= a_1 + 2a_2 x \\ y_p''(x) &= 2a_2 \end{aligned}$$

in die Differenzialgleichung ein. Dies liefert

$$2y_p''(x) + y_p'(x) = 4a_2 + a_1 + 2a_2 x \stackrel{!}{=} x.$$

Koeffizientenvergleich:

$$\begin{aligned} x^1: \quad 2a_2 &= 1 &\Rightarrow a_2 &= \frac{1}{2} \\ x^0: \quad 4a_2 + a_1 &= 0 &\Rightarrow a_1 &= -2. \end{aligned}$$

Für a_0 besteht **keine** Bedingung; a_0 kann somit z.B. auf Null gesetzt werden: $a_0 = 0$.

$$\Rightarrow y_p(x) = -2x + \frac{1}{2}x^2. \quad \square$$

Beispiele 12.39:

① $y''(x) + y(x) = e^{ix}:$

Das zugehörige charakteristische Polynom ist $P(\lambda) = \lambda^2 + 1$. Die Inhomogenität ist $e^{ix} \hookrightarrow \mu = i$. $\mu = i$ ist einfache Nullstelle $\hookrightarrow k = 1$; $m = 0 \Rightarrow k + m = 1$.

Ansatz:

$$\begin{aligned} y_p(x) &= (a_0 + a_1 x) e^{ix} \\ y_p'(x) &= a_1 e^{ix} + i(a_0 + a_1 x) e^{ix} \\ y_p''(x) &= 2i a_1 e^{ix} - (a_0 + a_1 x) e^{ix} \end{aligned}$$

In Differenzialgleichung einsetzen:

$$\begin{aligned} y_p''(x) + y_p(x) &= 2a_1 i e^{ix} - (a_0 + a_1 x) e^{ix} + (a_0 + a_1 x) e^{ix} \\ &= 2a_1 i e^{ix} \stackrel{!}{=} e^{ix} \end{aligned}$$

$$\Rightarrow a_1 = \frac{1}{2i} \quad \text{und} \quad a_0 = 0 \quad (\text{beliebig}).$$

$$\Rightarrow y_p(x) = \frac{1}{2i} x e^{ix} = -\frac{1}{2} i x e^{ix}.$$

$$\textcircled{2} \quad y''(x) + y(x) = \cos x: \quad (*)$$

Um eine partikuläre Lösung zu erhalten, setzen wir die Differenzialgleichung ins Komplexe fort

$$\tilde{y}''(x) + \tilde{y}(x) = e^{ix} \quad (\tilde{*})$$

$\tilde{y}_p(x) = -\frac{1}{2} i x e^{ix}$ ist nach $\textcircled{1}$ eine partikuläre Lösung von $(\tilde{*})$. Eine partikuläre Lösung von $(*)$ ist demnach gegeben durch den Realteil von $\tilde{y}_p(x)$

$$y_p(x) = \operatorname{Re}(\tilde{y}_p(x)).$$

Wegen
$$-\frac{1}{2} i x e^{ix} = +\frac{1}{2} e^{i 270^\circ} x e^{ix} = \frac{1}{2} x e^{i(x+270^\circ)}$$

$$\Rightarrow y_p(x) = \operatorname{Re}\left(\frac{1}{2} x e^{i(x+270^\circ)}\right) = \frac{1}{2} x \cos(x+270^\circ) = \frac{1}{2} x \sin x.$$

$$\textcircled{3} \quad y''(x) + y(x) = \sin x:$$

Nach dem Vorgehen in $\textcircled{2}$ ist eine partikuläre Lösung

$$y_p(x) = \operatorname{Im}(\tilde{y}_p(x)) = \operatorname{Im}\left(\frac{1}{2} x e^{i(x+270^\circ)}\right) = \frac{1}{2} x \sin(x+270^\circ).$$

$$y_p(x) = -\frac{1}{2} x \cos x.$$

□

⊗ Spezialfälle von Satz 12.17:

Um eine partikuläre Lösung der inhomogenen Differenzialgleichung

$$y^{(n)}(x) + a_{n-1} y^{(n-1)}(x) + \dots + a_1 y'(x) + a_0 y(x) = f(x)$$

zu erhalten, kann in manchen Spezialfällen direkt ein reeller Ansatz gewählt werden. Tabelle 12.2 gibt für häufig auftretende Störfunktionen die entsprechende Ansatzfunktion an.

Tabelle 12.2: Ansatzfunktionen für partikuläre Lösungen.

Störfunktion	NS von $P(\lambda)$	Ansatz
$f(x) = \sum_{i=0}^n a_i x^i$	0 keine NS	$y_p(x) = \sum_{i=0}^n A_i x^i$
	0 k -fache NS	$y_p(x) = x^k \sum_{i=0}^n A_i x^i$
$f(x) = a e^{\mu x}$	μ keine NS	$y_p(x) = A e^{\mu x}$
	μ k -fache NS	$y_p(x) = A x^k e^{\mu x}$
$f(x) = a \sin(\beta x)$	$i\beta$ keine NS	$y_p(x) = A \sin(\beta x) + B \cos(\beta x)$ $= C \sin(\beta x + \varphi)$
$f(x) = a \cos(\beta x)$	$i\beta$ k -fache NS	$y_p(x) = A x^k \sin(\beta x) + B x^k \cos(\beta x)$ $= C x^k \sin(\beta x + \varphi)$

Beispiel 12.40. Gegeben ist die Differenzialgleichung

$$y''(x) + 2y'(x) + y(x) = f(x) .$$

In der nachfolgenden Liste geben wir nach Tabelle 12.2 für verschiedene Störfunktionen f einen Ansatz für eine partikuläre Lösung sowie die Lösung für die freien Parameter an. Das charakteristische Polynom ist

$$P(\lambda) = \lambda^2 + 2\lambda + 1 = (\lambda + 1)^2$$

und damit $\lambda = -1$ eine doppelte Nullstelle.

$f(x)$	Ansatzfunktion	Parameter
$x^2 - 2x + 1$	$y_p(x) = a_0 + a_1 x + a_2 x^2$ ($\mu = 0$ keine NS von $P(\lambda)$)	$a_0 = 11, a_1 = -6, a_2 = 1$
$2e^x$	$y_p(x) = A e^x$ ($\mu = 1$ keine NS von $P(\lambda)$)	$A = \frac{1}{2}$
$\cos x$	$\tilde{y}_p(x) = A e^{ix} \rightarrow$ in kompl. DG $y_p(x) = \operatorname{Re}(A e^{ix})$	$A = -\frac{1}{2}i$ $y_p(x) = \frac{1}{2} \cos(x + \frac{3\pi}{2})$
$\cos x$	$y_p(x) = A \sin x + B \cos x$ ($\mu = i$ keine NS von $P(\lambda)$)	$A = \frac{1}{2}, B = 0$
$\sin x$	$y_p(x) = A \sin x + B \cos x$ ($\mu = i$ keine NS von $P(\lambda)$)	$A = 0, B = -\frac{1}{2}$
e^{-x}	$y_p(x) = a_2 x^2 e^{-x}$ ($\mu = -1$ ist doppelte NS)	$a_2 = \frac{1}{2}$
$-x^2 e^x$	$(a_0 + a_1 x + a_2 x^2) e^x$	$a_0 = -\frac{3}{8}, a_1 = \frac{1}{2}, a_2 = -\frac{1}{4}$
$x e^{-x}$	$(a_0 + a_1 x) x^2 e^{-x}$	$a_0 = 0, a_1 = \frac{1}{6}$

Zusammenfassung: Lineare Differenzialgleichung n -ter Ordnung mit konstanten Koeffizienten.

Die allgemeine Lösung der inhomogenen DG n -ter Ordnung

$$y^{(n)}(x) + a_{n-1}y^{(n-1)}(x) + \dots + a_1y'(x) + a_0y(x) = f(x) \quad (*)$$

setzt sich zusammen aus der **allgemeinen homogenen Lösung** $y_h(x)$ und **einer partikulären Lösung** $y_p(x)$ der inhomogenen DG:

$$y(x) = y_h(x) + y_p(x) \ .$$

1. Bestimmung der allgemeinen homogenen Lösung

$$y^{(n)}(x) + a_{n-1}y^{(n-1)}(x) + \dots + a_1y'(x) + a_0y(x) = 0 :$$

(1) Ansatz $y(x) = e^{\lambda x}$ in Differenzialgleichung.

(2) Charakteristisches Polynom

$$P(\lambda) = \lambda^n + a_{n-1}\lambda^{n-1} + \dots + a_1\lambda + a_0.$$

(3) Nullstellen des charakteristischen Polynoms $\lambda_1, \dots, \lambda_m$

$$\lambda_i \text{ einfache NS} \rightarrow \varphi_i(x) = e^{\lambda_i x}.$$

$$\lambda_i \text{ } k\text{-fache NS} \rightarrow \varphi_i(x) = e^{\lambda_i x}, x e^{\lambda_i x}, \dots, x^{k-1} e^{\lambda_i x}.$$

(4) $\varphi_1(x), \varphi_2(x), \dots, \varphi_n(x)$ ist Fundamentalsystem.

(5) Allgemeine, homogene Lösung

$$y_h(x) = c_1\varphi_1(x) + c_2\varphi_2(x) + \dots + c_n\varphi_n(x) \ .$$

2. Bestimmung einer partikulären Lösung der inhomogenen DG: Gemäß Tab. 12.2 wählt man spezielle Ansätze für eine partikuläre Lösung oder man geht auf Satz 12.17 zurück. Setzt sich die Störfunktion aus mehreren Störgliedern $f_1(x), \dots, f_l(x)$ zusammen, wählt man für jedes Störglied einen partikulären Ansatz $y_{p_1}(x), \dots, y_{p_l}(x)$ und löst

$$y_{p_i}^{(n)}(x) + a_{n-1}y_{p_i}^{(n-1)}(x) + \dots + a_0y_{p_i}(x) = f_i(x) \ .$$

Eine partikuläre Lösung für die Störfunktion $f(x) = f_1(x) + \dots + f_l(x)$ ist dann

$$y_p(x) = y_{p_1}(x) + \dots + y_{p_l}(x) \ .$$

3. Die allgemeine Lösung der Differenzialgleichung (*) lautet

$$y(x) = c_1y_1(x) + \dots + c_ny_n(x) + y_p(x) \ .$$

4. Die Koeffizienten $c_1, \dots, c_n$ bestimmen sich aus den Anfangsbedingungen $y(0), y'(0), \dots, y^{(n-1)}(0)$.

Beispiel 12.41 (Musterbeispiel).

Gesucht ist die Lösung von

$$y''(x) - 6y'(x) + 9y(x) = 4e^{2x} + 9x - 15$$

mit den Anfangsbedingungen $y(0) = y_0$, $y'(0) = 0$.

1. Lösen der homogenen Differenzialgleichung

$$y''(x) - 6y'(x) + 9y(x) = 0 :$$

Das charakteristische Polynom ist

$$P(\lambda) = \lambda^2 - 6\lambda + 9.$$

Die Nullstellen von $P(\lambda)$ sind $\lambda_1 = \lambda_2 = 3$ (doppelt). Die allgemeine Lösung der homogenen Differenzialgleichung ist somit

$$y_h(x) = c_1 e^{3x} + c_2 x e^{3x}.$$

2. Berechnung einer speziellen Lösung:

Die Störfunktion $f(x) = 4e^{2x} + 9x - 15$ besteht aus zwei Funktionstypen. In zwei Schritten bestimmen wir eine partikuläre Lösung:

$$(i) \quad y''(x) - 6y'(x) + 9y(x) = 4e^{2x}. \quad (1)$$

$\mu = 2$ ist keine NS von $P(\lambda)$ und daher erhalten wir eine partikuläre Lösung der Form

$$y_{p_1}(x) = A e^{2x}.$$

In die Differenzialgleichung (1) eingesetzt folgt:

$$A(4 - 6 \cdot 2 + 9) e^{2x} \stackrel{!}{=} 4e^{2x} \Rightarrow A = 4.$$

$$\Rightarrow y_{p_1}(x) = 4e^{2x}.$$

$$(ii) \quad y''(x) - 6y'(x) + 9y(x) = 9x - 15. \quad (2)$$

0 ist keine NS von $P(\lambda)$ und daher wird als Ansatz

$$y_{p_2}(x) = a_0 + a_1 x$$

in die Differenzialgleichung (2) eingesetzt:

$$\begin{aligned} y''_{p_2}(x) - 6y'_{p_2}(x) + 9y_{p_2}(x) &= -6a_1 + 9(a_0 + a_1 x) \\ &= (-6a_1 + 9a_0) + 9a_1 x \stackrel{!}{=} 9x - 15. \end{aligned}$$

Der Koeffizientenvergleich liefert

$$\begin{aligned} x^1: \quad 9a_1 &= 9 & \Rightarrow a_1 &= 1, \\ x^0: \quad -6a_1 + 9a_0 &= -15 & \Rightarrow a_0 &= -1. \\ & \Rightarrow y_{p_2}(x) &= -1 + x. \end{aligned}$$

(iii) Die partikuläre Lösung für die Störfunktion $4e^{2x} + 9x - 15$ setzt sich zusammen aus y_{p_1} und y_{p_2} :

$$y_p(x) = 4e^{2x} + x - 1.$$

3. Die allgemeine Lösung der Differenzialgleichung lautet somit

$$y(x) = c_1 e^{3x} + c_2 x e^{3x} + 4e^{2x} + x - 1.$$

4. Bestimmung der Konstanten c_1, c_2 über die Anfangsbedingungen:

$$\begin{aligned} y(0) &= c_1 + 4 - 1 \stackrel{!}{=} y_0 \Rightarrow c_1 = y_0 - 3 \\ y'(0) &= 3c_1 + c_2 + 9 \stackrel{!}{=} 0 \Rightarrow c_2 = -9 - 3c_1 = -3y_0. \end{aligned}$$

$$\Rightarrow y(x) = (y_0 - 3)e^{3x} - 3y_0 x e^{3x} + 4e^{2x} + x - 1. \quad \square$$

Anwendungsbeispiel 12.42 (Erzwungene, gedämpfte Schwingung, mit MAPLE-Worksheet).

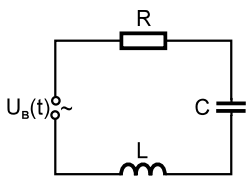


Abb. 12.22. RCL-Kreis

Ein elektrischer Schwingkreis besteht aus einem Ohmschen Widerstand R , einem Kondensator mit Kapazität C und einer Spule mit Induktivität L (siehe Abb. 12.22). Zur Zeit $t = 0$ wird der Stromkreis durch Anlegen einer äußeren Wechselspannung

$$U_B(t) = U_0 \sin(\omega t)$$

geschlossen. Nach dem Maschensatz gilt

$$L \frac{dI(t)}{dt} + RI(t) + \frac{1}{C} \int_0^t I(\tau) d\tau = U_0 \sin(\omega t).$$

Gehen wir zur komplexen Formulierung über, ist

$$L \dot{I}(t) + RI(t) + \frac{1}{C} \int_0^t I(\tau) d\tau = U_0 e^{i\omega t}$$

und nach Differenziation

$$\ddot{I}(t) + \frac{R}{L} \dot{I}(t) + \frac{1}{LC} I(t) = \frac{U_0}{L} i\omega e^{i\omega t}.$$

Setzen wir noch $\beta = \frac{R}{2L}$ (Dämpfung) und $\omega_0^2 = \frac{1}{LC}$, ist

$$\ddot{I}(t) + 2\beta \dot{I}(t) + \omega_0^2 I(t) = \frac{U_0}{L} i \omega e^{i\omega t}.$$

Die allgemeine, homogene Lösung wurde in Beispiel 12.35 diskutiert, so dass zur Lösung der inhomogenen Differenzialgleichung noch eine partikuläre Lösung gesucht ist. Das charakteristische Polynom zur Differenzialgleichung lautet

$$P(\lambda) = \lambda^2 + 2\beta\lambda + \omega_0^2.$$

$\hookrightarrow \mu = i\omega$ ist keine Nullstelle des charakteristischen Polynoms. Eine partikuläre Lösung liefert daher der Ansatz

$$\tilde{I}_p(t) = A e^{i\omega t}.$$

In die Differenzialgleichung eingesetzt, folgt

$$(i\omega)^2 A e^{i\omega t} + 2\beta(i\omega) A e^{i\omega t} + \omega_0^2 A e^{i\omega t} = \frac{U_0}{L} i \omega e^{i\omega t}$$

$$\Rightarrow A = i \frac{U_0 \omega}{L} \frac{1}{\omega_0^2 - \omega^2 + 2i\beta\omega}$$

und

$$\tilde{I}_p(t) = i \frac{U_0 \omega}{L} \frac{1}{\omega_0^2 - \omega^2 + 2i\beta\omega} e^{i\omega t}.$$

Übergang zu einer reellen partikulären Lösung: $I_p(t) = \text{Im}(\tilde{I}_p(t)).$

Wir stellen $\tilde{I}_p(t)$ in der exponentiellen Normalform dar, indem wir A in die Exponentialform überführen und die Multiplikation in dieser Normalform ausführen.

$$\begin{aligned} A &= i \frac{U_0 \omega}{L} \frac{1}{\omega_0^2 - \omega^2 + 2i\beta\omega} \cdot \frac{\omega_0^2 - \omega^2 - i2\beta\omega}{\omega_0^2 - \omega^2 - i2\beta\omega} \\ &= \frac{U_0 \omega}{L} \frac{1}{(\omega_0^2 - \omega^2)^2 + (2\beta\omega)^2} [2\beta\omega + i(\omega_0^2 - \omega^2)] \stackrel{!}{=} |A| e^{i\varphi} \end{aligned}$$

$$\text{mit } |A| = \frac{U_0 \omega}{L} \frac{1}{\sqrt{(\omega_0^2 - \omega^2)^2 + (2\beta\omega)^2}} \quad \text{und} \quad \tan \varphi = \frac{\text{Im } A}{\text{Re } A} = \frac{\omega_0^2 - \omega^2}{2\beta\omega}.$$

$$\Rightarrow \tilde{I}_p(t) = \underbrace{\frac{U_0 \omega}{L} \frac{1}{\sqrt{(\omega_0^2 - \omega^2)^2 + (2\beta\omega)^2}}}_{\text{komplexe Amplitude}} e^{i\varphi} e^{i\omega t}$$

$$\Rightarrow I_p(t) = \operatorname{Im} \left(\tilde{I}_p(t) \right) = \frac{U_0 \omega}{L} \frac{1}{\sqrt{(\omega_0^2 - \omega^2)^2 + (2\beta\omega)^2}} \sin(\omega t + \varphi).$$

Interpretation: Mit $\omega_0^2 = \frac{1}{LC}$ und $2\beta = \frac{R}{L}$ gilt für Amplitude und Phase des Stromes

$$\frac{U_0 \omega}{L} \frac{1}{\sqrt{(2\beta\omega)^2 + (\omega_0^2 - \omega^2)^2}} = \boxed{\frac{U_0}{\sqrt{R^2 + \left(\frac{1}{\omega C} - \omega L\right)^2}} =: I_0}$$

und

$$\tan \varphi(\omega) = \frac{\omega L - \frac{1}{\omega C}}{R}.$$

$$\Rightarrow I_p(t) = I_0 \sin(\omega t + \varphi) = \frac{U_0}{\sqrt{R^2 + \left(\frac{1}{\omega C} - \omega L\right)^2}} \sin(\omega t + \varphi). \quad (*)$$

Gleichung (*) ist das Ohmsche Gesetz der Wechselstromtechnik mit den Scheitelwerten U_0 und I_0 . Der reelle Scheinwiderstand beträgt

$$Z = \sqrt{R^2 + \left(\frac{1}{\omega C} - \omega L\right)^2} = Z(\omega).$$

Der Strom $I_p(t)$ ist um $\varphi(\omega)$ phasenverschoben zur Spannung. □

MAPLE-Worksheets zu Kapitel 12

Differenzialgleichungen erster Ordnung:

- Differenzialgleichungen erster Ordnung mit MAPLE
- Richtungsfelder mit MAPLE
- Beispiele mit MAPLE

Lineare Differenzialgleichungssysteme:

- Eigenwerte und Eigenvektoren mit MAPLE
- Beispiele mit MAPLE

Lineare Differenzialgleichungen n -ter Ordnung:

- Differenzialgleichungen n -ter Ordnung mit MAPLE
- Beispiele mit MAPLE

Numerisches Lösen von Differenzialgleichungen:

- Numerisches Lösen mit **dsolve**
- Numerisches Lösen mit **DEplot**
- Numerisches Lösen mit dem Euler-Verfahren
- Beispiele mit MAPLE

12.4 Aufgaben zu Differenzialgleichungen

Differenzialgleichungen 1. Ordnung

- 12.1 Wie lauten die allgemeinen Lösungen der folgenden linearen DG erster Ordnung mit konstanten Koeffizienten?

$$\begin{array}{lll} \text{a) } y' + 4y = 0 & \text{b) } 2y' + 4y = 0 & \text{c) } -3y' = 8y \\ \text{d) } ay' - by = 0 \quad (a \neq 0) & \text{e) } -3y' + 18y = 0 & \text{f) } L \frac{dI}{dt} + RI = 0 \end{array}$$

- 12.2 Lösen Sie folgende Anfangswertprobleme:

$$\begin{array}{l} \text{a) } 2v + \dot{v} = 0, \quad v(0) = 10 \frac{\text{m}}{\text{s}}. \quad \text{Wann ist } v(t) < 10 \frac{\text{cm}}{\text{s}}? \\ \text{b) } y' + \lambda y = 0, \quad y(0) = 1. \quad \text{Was ergibt sich für } \lambda, \text{ wenn } y(1) = \frac{1}{2}? \\ \text{c) } -\frac{dN}{dt} = \frac{1}{\tau} N, \quad N(0) = N_0. \quad \text{Wie groß ist } \tau, \text{ wenn } N(1,819 \cdot 10^{11}) = \frac{1}{2} N_0? \end{array}$$

- 12.3 Bestimmen Sie die allgemeine Lösung der folgenden linearen DG erster Ordnung:

$$\begin{array}{lll} \text{a) } y' + xy = 4x & \text{b) } y' + \frac{y}{1+x} = e^{2x} & \text{c) } xy' + y = x \cdot \sin x \\ \text{d) } y' \cos x - y \sin x = 1 & \text{e) } y' - 2 \cos xy = \cos x & \text{f) } xy' - y = x^2 + 4 \end{array}$$

- 12.4 Lösen Sie folgende inhomogenen DG:

$$\begin{array}{ll} \text{a) } y'(t) + \frac{1}{RC} y(t) = U_0 \sin(\omega t) & \text{(RC-Wechselstromkreis)} \\ \text{b) } y'(t) + \frac{R}{L} y(t) = U_0 e^{-2t} & \text{(RL-Wechselstromkreis)} \end{array}$$

- 12.5 Bestimmen Sie eine Lösung für die folgenden DG 1. Ordnung durch Trennung der Variablen:

$$\begin{array}{lll} \text{a) } x^2 y' = y^2 & \text{b) } y' (1 + x^2) = xy & \text{c) } y' = 1 - y^2 \\ \text{d) } y' = (1 - y)^2 & \text{e) } y' \sin y = -x & \text{f) } y' = e^y \cos x \end{array}$$

- 12.6 Lösen Sie die folgenden Anfangswertprobleme:

$$\begin{array}{l} \text{a) } y' + \cos x \cdot y = 0; \quad y\left(\frac{\pi}{2}\right) = 2\pi \\ \text{b) } x(x+1)y' = y; \quad y(1) = \frac{1}{2} \\ \text{c) } y^2 y' + x^2 = 1; \quad y(2) = 1 \\ \text{d) } x^2 y' = y^2 + xy; \quad y(1) = -1 \quad (\text{Substitution: } u = \frac{y}{x}) \\ \text{e) } y y' = 2e^{2x}; \quad y(0) = 2 \end{array}$$

- 12.7 Man löse durch Substitution ($u = \frac{y}{x}$)

$$\text{a) } xy' = y + 4x \qquad \text{b) } x^2 y' = \frac{1}{4} x^2 + y^2$$

- 12.8 Lösen Sie folgenden Differenzialgleichungen durch Trennung der Variablen:

$$\begin{array}{ll} \text{a) } \frac{1}{x^5} y'(x) = \frac{1}{(y(x))^2} & \text{b) } (x+1)(x-1)y'(x) = 2y(x) \\ \text{c) } xy' + \frac{1}{x} y' = y & \text{d) } y^2 - 2yy' + 1 = 0 \\ \text{e) } y' = \cos^2 y & \text{f) } 2y' = y^4 \cdot \sqrt{x} \end{array}$$

- 12.9 Welche Lösungen haben folgende Anfangswertprobleme:

$$\begin{array}{ll} \text{a) } y'(x) = \sin x \cdot y^2(x), \quad y(0) = 1 & \text{b) } y + y' = e^x, \quad y(0) = 1 \\ \text{c) } \sin x \cdot y' = \cos x \cdot y, \quad y\left(\frac{\pi}{2}\right) = \frac{\pi}{2} & \text{d) } y' \cdot (1 + x^2) = 2xy, \quad y(1) = 4 \end{array}$$

Anwendungen

- 12.11 Eine chemische Reaktion
- $A + B \rightarrow X$
- lässt sich durch

$$\frac{dx}{dt} = k (a - x) (b - x)$$

beschreiben, wenn die Anzahl der Moleküle vom Typ A bzw. B zu Beginn der Reaktion a bzw. b und $x(t)$ die Anzahl der Reaktionsmoleküle X zum Zeitpunkt t sind (k : Reaktionskonstante).

- a) Man löse die DG für $a \neq b$ und $x(0) = 0$ (mit MAPLE).
 b) Wann kommt die Reaktion zum Stillstand ($a > b$) ?

- 12.12 Die Sinkgeschwindigkeit
- $v(t)$
- eines Teilchens der Masse
- m
- in einer Flüssigkeit (
- k
- : Reibungsfaktor,
- g
- : Erdbeschleunigung) wird beschrieben durch

$$m \frac{dv}{dt} + k v = m g :$$

- a) Man bestimme (mit MAPLE) die Geschwindigkeit und Position zur Zeit $t > 0$ für die Anfangswerte $v(0) = v_0$ und $s(0) = 0$.
 b) Welche Geschwindigkeit $v_{\max}$ kann das Teilchen maximal erreichen?
- 12.13 Man löse Aufgabe 12.12, wenn das Medium einen Widerstand leistet, der gleich $k v^2$ ist und $v(0) = 0$.
- 12.14 Ein Körper besitze zur Zeit $t = 0$ die Temperatur T_0 und werde in der Folgezeit durch vorbeiströmende Luft der konstanten Temperatur T_L gemäß

$$\frac{dT}{dt} = -a (T - T_L) \quad (a > 0)$$

gekühlt. Bestimmen Sie den zeitlichen Verlauf der Körpertemperatur. Gegen welchen Endwert strebt diese?

- 12.15 Die radiale Geschwindigkeitsverteilung stationärer, laminarer Strömungen eines viskosen inkompressiblen Fluids (Viskosität
- η
-) längs eines Rohrstücks, in dem ein Druckabfall
- $\frac{\Delta p}{\Delta z}$
- wirkt, kann durch

$$-\frac{\Delta p}{\Delta z} + \eta \frac{1}{r} \frac{d}{dr} \left(r \frac{d}{dr} v_z(r) \right) = 0$$

beschrieben werden. Wie gross ist $v_z(r)$, wenn am Rand $v_z(R) = 0$ gilt? Hinweis: Integrieren Sie, nach geeigneten Umformungen, zunächst von 0 bis r und überlegen Sie, was sich für die Integrationskonstante bei $r = 0$ ergibt. Integrieren Sie dann von R bis r !

- 12.16 Ein Körper rollt eine schiefe Ebene (Winkel
- φ
-) hinunter und erfährt dabei Reibungskräfte proportional zu seiner Geschwindigkeit und einen Druckwiderstand proportional zum Quadrat seiner Geschwindigkeit. Es gilt:

$$m \cdot \dot{v} + R \cdot v + D \cdot v^2 = m \cdot g \cdot \sin \varphi .$$

Die Anfangsgeschwindigkeit sei $v(0) = 0$; die physikalischen Konstanten $m = 1$, $g = 10$, $\varphi = \frac{\pi}{3}$, $D = \frac{4}{5}$, $R = 3$.

- a) Was ergibt sich für $v(t)$ mit $D = 0$?
 b) Was ergibt sich für $v(t)$ mit $R = 0$?

Lineare Differenzialgleichungssysteme

12.17 Bestimmen Sie ein Lösungsfundamentalsystem des LDGS erster Ordnung

$$\begin{aligned} \text{a) } y'(t) &= A y(t) \quad \text{mit } A = \begin{pmatrix} 2 & 0 & -2 \\ 0 & 4 & 0 \\ -2 & 0 & 5 \end{pmatrix} \\ \text{b) } y'(t) &= B y(t) \quad \text{mit } B = \begin{pmatrix} -2 & -9 & 5 \\ -5 & -10 & 7 \\ -9 & -21 & 14 \end{pmatrix} \end{aligned}$$

Prüfen Sie, ob die Eigenvektoren eine Basis des $\mathbb{R}^3$ bilden.

12.18 Bestimmen Sie Lösungen des LDGS zweiter Ordnung

$$y''(t) = A y(t) \quad \text{mit } A = \begin{pmatrix} 1 & -2 \\ -2 & 4 \end{pmatrix}.$$

12.19 a) Die Bewegungsgleichungen eines geladenen Teilchens im Magnetfeld lauten

$$\dot{v}_x = -\frac{e}{m} B_z v_y, \quad \dot{v}_y = \frac{e}{m} B_z v_x$$

wenn $\vec{B} = B_z \vec{e}_z$. Man bestimme ein reelles Fundamentalsystem.

b) Man bestimme eine partikuläre Lösung, wenn neben dem Magnetfeld $\vec{B}$ noch ein elektrisches Feld $\vec{E} = E_0 \begin{pmatrix} 0 \\ t \\ 0 \end{pmatrix}$ wirkt:

$$\dot{v}_x = -\frac{e}{m} B_z v_y, \quad \dot{v}_y = \frac{e}{m} B_z v_x + E_0 \cdot t.$$

12.20 Gegeben sei die Matrix $A = \begin{pmatrix} -3 & 1 \\ 1 & -3 \end{pmatrix}$.

- Man bestimme zur Matrix A sämtliche Eigenwerte und Eigenvektoren.
- Man bestimme ein Fundamentalsystem von $\vec{y}'(t) = A \vec{y}(t)$.
- Man bestimme ein komplexes FS von $\vec{y}''(t) = A \vec{y}(t)$.
- Man bestimme ein reelles FS von $\vec{y}''(t) = A \vec{y}(t)$.
- Man stelle das zu $\vec{y}''(t)$ äquivalente LDGS 1. Ordnung auf.

12.21 Geben Sie je ein Fundamentalsystem für $\vec{y}' = A \vec{y}$ an:

$$\begin{aligned} \text{a) } A &= \begin{pmatrix} 3 & 4 \\ -5 & -5 \end{pmatrix} & \text{b) } A &= \begin{pmatrix} 3 & 1 & 1 \\ 1 & 5 & 1 \\ 1 & 1 & 3 \end{pmatrix} & \text{c) } A &= \begin{pmatrix} 3 & -1 & 1 \\ -1 & 3 & -1 \\ 1 & -1 & 3 \end{pmatrix} \end{aligned}$$

12.22 Lösen Sie das Anfangswertproblem:

$$\begin{aligned} y_1'(x) &= 3y_1(x) + 2y_2(x) - y_3(x), & y_1(0) &= 2 \\ y_2'(x) &= 2y_1(x) + 3y_2(x) - y_3(x), & y_2(0) &= 4 \\ y_3'(x) &= -y_1(x) - y_2(x) + 4y_3(x), & y_3(0) &= 0 \end{aligned}$$

12.23 "Knacken" Sie die Differenzialgleichung 2. Ordnung

$$y'' - 5y' + 6y = 0 \quad (*)$$

indem Sie die Hilfsfunktionen $y_1 = y$, $y_2 = y'$ einführen und $(*)$ als System schreiben und lösen! Wie lautet die Lösung für $y(0) = 1$, $y'(0) = 0$?

Differenzialgleichungen höherer Ordnung

- 12.25 Prüfen Sie nach, dass
- $\varphi_1 = 1 - \cos(2x)$
- und
- $\varphi_2 = 1 - \cos^2(x)$
- Lösungen von

$$y'' - (\tan x + \cot x) y' + 4y = 0.$$

sind. Bilden sie ein Fundamentalsystem?

- 12.26 Zeigen Sie, dass die beiden Funktionen
- $\sinh(kx)$
- und
- $\cosh(kx)$
- ein reelles Fundamentalsystem bilden von der Differenzialgleichung

$$y''(x) - k^2 y(x) = 0.$$

- 12.27 Lösen Sie die folgenden homogenen, linearen Differenzialgleichungen 2. Ordnung:

$$\begin{array}{ll} \text{a) } \ddot{u}(t) + 13\dot{u}(t) + 40u(t) = 0 & \text{b) } \ddot{v}(t) - 12\dot{v}(t) + 36v(t) = 0 \\ \text{c) } y''(x) + 6y'(x) + 34y(x) = 0 & \text{d) } z''(x) + 16z(x) = 0 \end{array}$$

- 12.28 Bestimmen Sie ein reelles Fundamentalsystem für

$$\begin{array}{l} \text{a) } y^{(4)}(x) - 10y''(x) + 9y(x) = 0 \\ \text{b) } u^{(3)}(t) - 2\ddot{u}(t) + \dot{u}(t) = 0 \\ \text{c) } y^{(6)}(x) - y(x) = 0 \end{array}$$

- 12.29 Gegeben ist die inhomogene, lineare Differenzialgleichung 2. Ordnung

$$y''(x) - 3y'(x) + 2y(x) = s(x)$$

mit der Inhomogenität $s(x)$. Ermitteln Sie partikuläre Lösungen für

$$\begin{array}{llll} \text{a) } s(x) = 6 & \text{b) } s(x) = x & \text{c) } s(x) = e^{2x} & \text{d) } s(x) = \cos x \\ \text{e) } s(x) = 4x + 10 \cos x & \text{f) } s(x) = x e^{2x} & \text{g) } s(x) = \cos x e^x \end{array}$$

- 12.30 Lösen Sie die Schwingungsprobleme

$$\begin{array}{l} \text{a) } \ddot{x}(t) + 16x(t) = 0, \quad x(0) = 3, \quad \dot{x}(0) = 4 \\ \text{b) } \ddot{x}(t) + 2\dot{x}(t) + 2x(t) = 0, \quad x(0) = 2, \quad \dot{x}(0) = 0 \\ \text{c) } \ddot{x}(t) + 13\dot{x}(t) + 40x(t) = 0, \quad x(0) = 3, \quad \dot{x}(0) = 0 \end{array}$$

- 12.31 Bestimmen Sie alle reellen Lösungen der folgenden DG (mit MAPLE)

$$\begin{array}{l} \text{a) } y^{(4)}(x) - 10y''(x) + 9y(x) = \sin(x) \\ \text{b) } y'''(x) - 7y'(x) - 6y(x) = 12e^x \\ \text{c) } y'''(x) - 2y''(x) + y'(x) - 2y(x) = \cos(x) \\ \text{d) } y'''(x) - 6y''(x) + 12y'(x) - 8y(x) = 6e^{2x} \end{array}$$

- 12.32 Lösen Sie die Differenzialgleichungen aus Aufgabe 12.6 numerisch mit dem Euler-Verfahren. Variieren Sie die Schrittweite und vergleichen Sie das numerische Ergebnis mit der exakten Lösung.

Kapitel 13
Laplace-Transformation

13

13	Laplace-Transformation	559
13.1	Die Laplace-Transformation	563
13.2	Inverse Laplace-Transformation	568
13.3	Zwei grundlegende Eigenschaften	569
13.3.1	Linearität	569
13.3.2	Laplace-Transformierte der Ableitung	571
13.4	Methoden der Rücktransformation	574
13.5	Anwendungen der Laplace-Transformation	577
13.6	Aufgaben zur Laplace-Transformation	581

Zusätzliche Abschnitte auf der CD-Rom

13.7	Transformationssätze	cd
13.7.1	Verschiebungssatz	cd
13.7.2	Dämpfungssatz	cd
13.7.3	Ähnlichkeitssatz	cd
13.7.4	Faltungssatz	cd
13.7.5	Grenzwertsätze	cd
13.8	Laplace-Transformation mit MAPLE	cd
13.8.1	Laplace-Transformation	cd
13.8.2	Inverse Laplace-Transformation	cd
13.8.3	Anwendungen der Laplace-Transformation mit MAPLE	cd

13 Laplace-Transformation

Eine elegante Methode zur Lösung von Differenzialgleichungen macht Gebrauch von der *Laplace-Transformation*. Das sog. Laplace-Integral eignet sich besonders zur Behandlung von Differenzialgleichungen und Differenzialgleichungssystemen mit Anfangsbedingungen. Die mathematische Formulierung der Laplace-Transformierten einer Zeitfunktion $f(t)$ lautet

$$\mathcal{L}\{f(t)\} = F(s) = \int_0^{\infty} f(t) e^{-st} dt.$$

Dabei wird der Zeitfunktion $f(t)$ eine Bildfunktion $F(s)$ zugeordnet, so dass man bei der Laplace-Transformation auch von einer *Funktionaltransformation* spricht.

Hinweis: Auf der CD-Rom befindet sich ein zusätzlicher Abschnitt über die [Eigenschaften](#) der Laplace-Transformation (Verschiebungssatz, Dämpfungssatz, Ähnlichkeitssatz, Faltungssatz), das Berechnen der Laplace-Transformation sowie der inversen Transformation mit MAPLE und zahlreiche weitere [Anwendungsbeispiele](#).

Bei der Lösung von Differenzialgleichungen zeigt sich, dass mit der Laplace-Transformation der gesuchten Funktion $y(t)$ in den Bildbereich $Y(s)$ die Differenzialgleichung in eine algebraische Gleichung umgeformt wird. Diese algebraische Gleichung für $Y(s)$ lässt sich i.A. einfacher lösen, als die Differenzialgleichung für $y(t)$. Durch Rücktransformation erhält man schließlich die gesuchte Funktion $y(t)$. Dieser allgemeine Lösungsweg ist schematisch im folgenden Diagramm aufgezeigt (vgl. Abb. 13.1).

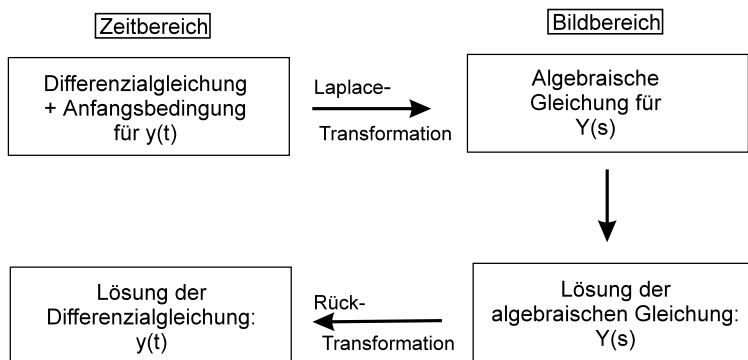


Abb. 13.1. Laplace-Transformation vom Zeitbereich in den Bildbereich

Damit zu gegebener Bildfunktion $Y(s)$ eine zugehörige Zeitfunktion $y(t)$ eindeutig bestimmt ist, muss die Laplace-Transformation und ihre *Rücktransformation* eindeutig sein, was für stetige Funktionen auch der Fall ist.

Von großem Vorteil für das Lösen von linearen Differenzialgleichungen mit der Laplace-Transformation ist, dass die Inhomogenitäten der Differenzialgleichung nicht stetig sein müssen und bei der Lösung automatisch die Anfangsbedingungen erfüllt werden. Anwendung findet die Laplace-Transformation u.a. in der Elektrotechnik, beim Lösen von linearen Differenzialgleichungssystemen mit Anfangsbedingungen.

Anwendungsbeispiel 13.1 (Elektrisches Netzwerk).

Gegeben ist das in Abb. 13.2 dargestellte elektrische Netzwerk. Gesucht sind die Einzelströme $I_1(t)$ und $I_2(t)$ in den beiden Zweigen, wenn die Anfangswerte $I_1(0) = I_2(0) = 0$.

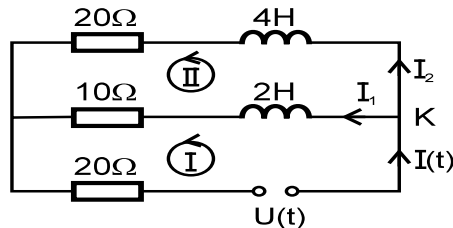


Abb. 13.2. Elektrisches Netzwerk

Aufstellen der Modellgleichungen: Wendet man die Kirchhoffschen Regeln auf das Netzwerk an, gilt nach dem Maschensatz

$$(I) \quad 20 I(t) + 2 \frac{d}{dt} I_1(t) + 10 I_1(t) = U(t)$$

$$(II) \quad -10 I_1(t) - 2 \frac{d}{dt} I_1(t) + 4 \frac{d}{dt} I_2(t) + 20 I_2(t) = 0$$

sowie nach dem Knotensatz

$$(K) \quad I(t) = I_1(t) + I_2(t).$$

Man erhält ein lineares Differenzialgleichungssystem 1. Ordnung

$$20 (I_1(t) + I_2(t)) + 2 I_1'(t) + 10 I_1(t) = U(t)$$

$$-10 I_1(t) + 20 I_2(t) - 2 I_1'(t) + 4 I_2'(t) = 0$$

mit den Anfangsbedingungen $I_1(0) = I_2(0) = 0$. Mit der Laplace-Transformation werden wir dieses System direkt unter Einbeziehung der Anfangsbedingungen lösen. $\square$

13.1 Die Laplace-Transformation

Die Laplace-Transformation ist eine Integraltransformation, die jeder *Zeitfunktion* $f(t)$, $t \geq 0$, eine *Bildfunktion* $F(s)$ gemäß

$$F(s) = \int_0^{\infty} f(t) e^{-st} dt$$

zuordnet. Da die Zeitintegration bei $t = 0$ beginnt, wird im Folgenden immer von $f(t) = 0$ für $t < 0$ ausgegangen.

Damit das uneigentliche Integral und damit die Bildfunktion $F(s)$ überhaupt definiert ist, muss das Integral für jedes s einen endlichen Wert annehmen. Eine hinreichende Bedingung hierfür ist, dass die Funktion $f(t)$ die beiden folgenden Eigenschaften besitzt:

Bedingung 1: $f : [0, \infty) \rightarrow \mathbb{R}$ ist eine **stückweise stetige** Funktion: Der Definitionsbereich der Funktion kann in endlich viele Teilintervalle unterteilt werden, in denen die Funktion stetig und beschränkt ist.

Bedingung 2: $f : [0, \infty) \rightarrow \mathbb{R}$ wächst nicht schneller als eine Exponentialfunktion $e^{\alpha t}$ mit geeignetem α : Es gibt ein $T > 0$ und Konstanten $\alpha \geq 0$, $M > 0$, so dass

$$|f(t)| \leq M e^{\alpha t} \quad \text{für } t \geq T.$$

Man nennt f dann von **höchstens exponentiellem Wachstum** der Ordnung α .

In Abb. 13.3 ist eine stückweise stetige Funktion mit höchstens exponentiellem Wachstum der Ordnung 1 gezeichnet. In jedem Teilintervall ist $f(t)$ stetig und für Zeiten $t > T$ ist $f(t) < e^t$.

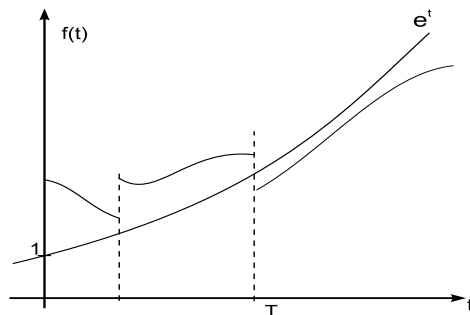


Abb. 13.3. Funktion von höchstens exponentiellem Wachstum

Beispiele 13.2:

- ① Funktionen von höchstens exponentiellem Wachstum:

$$f(t) = \text{const}, \quad t^n, \quad \cos(\omega t), \quad \sin(\omega t), \quad e^{\alpha t},$$

alle beschränkten Funktionen.

- ② Funktionen, die ein größeres Wachstum als das exponentielle besitzen:

$$e^{t^2}, \quad e^{\sin(t) \cdot t^3}.$$

Satz 13.1: Ist $f : [0, \infty) \rightarrow \mathbb{R}$ von höchstens exponentiellem Wachstum der Ordnung α , dann gilt

$$\lim_{t \rightarrow \infty} e^{-st} f(t) = 0 \quad \text{für } s > \alpha.$$

Begründung: Wenn $|f(t)| \leq M e^{\alpha t}$, dann ist für $s > \alpha$

$$|e^{-st} f(t)| \leq e^{-st} M e^{\alpha t} = M e^{(\alpha-s)t} \rightarrow 0 \quad \text{für } t \rightarrow \infty. \quad \square$$

Satz von Laplace: Sei $f : [0, \infty) \rightarrow \mathbb{R}$ eine stückweise stetige Funktion von höchstens exponentiellem Wachstum der Ordnung α (d.h. $|f(t)| \leq M e^{\alpha t}$ für $t > T$). Dann existiert

$$\mathcal{L}(f(t)) := F(s) := \int_0^\infty f(t) e^{-st} dt \quad \text{für } s > \alpha. \quad (*)$$

$\mathcal{L}(f(t))$ heißt **Laplace-Transformierte (Bildfunktion)** zur Zeitfunktion $f(t)$.

Bemerkungen:

- (1) I.A. ist $s = \delta + i\omega$ eine komplexe Variable und $F(s)$ eine komplexe Funktion. Im Folgenden werden wir aber bis auf die Angabe der inversen Laplace-Transformation s als reelle Variable und damit $F(s)$ als reellwertige Funktion betrachten.
- (2) Ein nach Formel (*) gebildetes Funktionenpaar $f(t)$ und $F(s)$ nennt man eine **Korrespondenz**. Man verwendet dafür auch die symbolische Schreibweise

$$f(t) \circ \bullet F(s).$$

- (3) Mit der Laplace-Transformation behandelt man zeitlich veränderliche Vorgänge, die zur Zeit $t = 0$ beginnen und die damit durch eine Funktion f mit $f(t) = 0$ für $t < 0$ beschrieben werden können.
- (4) Man kann allgemeiner die Laplace-Transformierte von Funktionen bilden, die statt Bedingung 1 die folgende allgemeinere Bedingung erfüllen:
In jedem endlichen Teilintervall von $[0, \infty)$ ist f stückweise stetig.
Diese Eigenschaft ist im Hinblick auf die Laplace-Transformierte von periodischen Funktionen von Bedeutung.

Beweis des Satzes von Laplace: Wir zeigen, dass das Integral

$$\int_0^{\infty} f(t) e^{-st} dt$$

für jedes $s > \alpha$ einen endlichen Wert annimmt: Da f stückweise stetig ist, lässt sich das Intervall $I = [0, \infty)$ in endlich viele Teilintervalle $I_1, \dots, I_n$ unterteilen, so dass f auf jedem dieser Intervalle $I_k = [t_{k-1}, t_k]$ ($k = 1, \dots, n$) stetig und beschränkt ist. Außerdem ist $f(t)$ von höchstens exponentiellem Wachstum der Ordnung α , d.h. es gibt ein T und Konstanten α, M , so dass

$$|f(t)| \leq M e^{\alpha t} \quad \text{für } t > T.$$

Wir nehmen an, dass $T > t_n$ und zerlegen den Definitionsbereich von f in

$$[0, \infty) = I_1 \cup I_2 \cup \dots \cup I_n \cup [t_n, T] \cup [T, \infty).$$

Dann ist

$$\begin{aligned} \int_0^{\infty} f(t) e^{-st} dt &= \int_0^{t_1} f(t) e^{-st} dt + \dots + \int_{t_{n-1}}^{t_n} f(t) e^{-st} dt \\ &\quad + \int_{t_n}^T f(t) e^{-st} dt + \int_T^{\infty} f(t) e^{-st} dt. \end{aligned}$$

Die ersten $n + 1$ Integrale sind endlich, da f darauf stetig und beschränkt ist. Das letzte Integral ist endlich, da f von höchstens exponentiellem Wachstum ist:

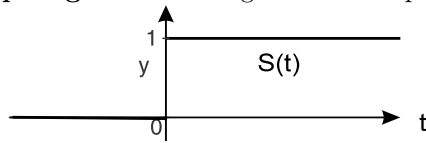
$$\begin{aligned} \left| \int_T^{\infty} f(t) e^{-st} dt \right| &\leq \int_T^{\infty} e^{-st} |f(t)| dt \leq M \int_T^{\infty} e^{-st} e^{\alpha t} dt \\ &= M \int_T^{\infty} e^{-(s-\alpha)t} dt = M \frac{1}{-(s-\alpha)} e^{-(s-\alpha)t} \Big|_T^{\infty} \\ &= \frac{M}{s-\alpha} e^{-(s-\alpha)T} - \frac{M}{s-\alpha} \lim_{t \rightarrow \infty} e^{-(s-\alpha)t} \\ &= \frac{M}{s-\alpha} e^{-(s-\alpha)T} \quad \text{für } s > \alpha. \end{aligned}$$

Damit sind alle Teilintegrale endlich und $F(s)$ für $s > \alpha$ definiert. □

Beispiele 13.3:

- ① Die Laplace-Transformierte der **Sprungfunktion**: Gegeben ist die Sprungfunktion (Heavisidefunktion)

$$S(t) := \begin{cases} 0 & \text{für } t < 0 \\ 1 & \text{für } t \geq 0 \end{cases}$$



Für $s > 0$ ist:

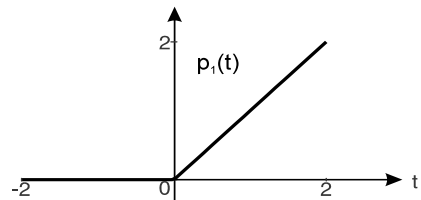
$$\mathcal{L}(S(t)) = \int_0^{\infty} 1 \cdot e^{-st} dt = \left[-\frac{1}{s} e^{-st} \right]_0^{\infty} = \frac{1}{s} \Rightarrow$$

$$S(t) \circ \bullet \frac{1}{s}.$$

- ② Laplace-Transformierte von **Potenzfunktionen**:

- (i) Die Laplace-Transformierte der unten gezeichneten linearen Funktion

$$p_1(t) := \begin{cases} 0 & \text{für } t < 0 \\ t & \text{für } t \geq 0 \end{cases}$$



erhält man mittels partieller Integration:

$$\mathcal{L}(p_1(t)) = \int_0^{\infty} t e^{-st} dt = t \frac{e^{-st}}{-s} \Big|_0^{\infty} + \frac{1}{s} \int_0^{\infty} e^{-st} dt = -\frac{1}{s^2} e^{-st} \Big|_0^{\infty} = \frac{1}{s^2}.$$

Die Korrespondenz lautet für $s > 0$

$$t S(t) \circ \bullet \frac{1}{s^2}.$$

- (ii) Die Laplace-Transformierte der Potenzfunktion

$$p_n(t) := \begin{cases} 0 & \text{für } t < 0 \\ t^n & \text{für } t \geq 0 \end{cases} = t^n S(t) \quad (n \in \mathbb{N})$$

lautet $\mathcal{L}(p_n(t)) = \frac{n!}{s^{n+1}}$. Man erhält diese Formel induktiv durch partielle Integration von

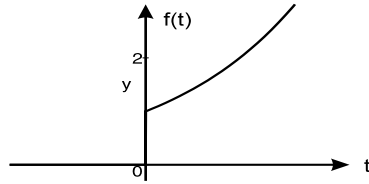
$$\begin{aligned} \mathcal{L}(p_{n+1}(t)) &= \int_0^{\infty} t^{n+1} e^{-st} dt = t^{n+1} \frac{e^{-st}}{-s} \Big|_0^{\infty} + \frac{n+1}{s} \int_0^{\infty} t^n e^{-st} dt \\ &= \frac{n+1}{s} \mathcal{L}(p_n(t)) = \frac{n+1}{s} \frac{n!}{s^{n+1}} = \frac{(n+1)!}{s^{n+2}}. \end{aligned}$$

Der Induktionsanfang ist durch ① bzw. ②(i) gegeben.

$$\Rightarrow t^n S(t) \circ \bullet \frac{n!}{s^{n+1}} \quad s > 0 \quad n \in \mathbb{N}_0 .$$

③ Die Laplace-Transformierte der **Exponentialfunktion**:

$$f(t) = \begin{cases} 0 & \text{für } t < 0 \\ e^{\alpha t} & \text{für } t \geq 0 \end{cases}$$



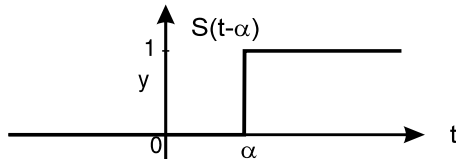
lautet $F(s) = \frac{1}{s-\alpha}$ für $s > \alpha$. Denn

$$\begin{aligned} \mathcal{L}(f(t)) &= \int_0^\infty e^{\alpha t} e^{-st} dt = \int_0^\infty e^{-(s-\alpha)t} dt \\ &= \frac{1}{-(s-\alpha)} e^{-(s-\alpha)t} \Big|_0^\infty = \frac{1}{s-\alpha} . \end{aligned}$$

$$\Rightarrow e^{\alpha t} S(t) \circ \bullet \frac{1}{s-\alpha} \quad \text{für } s > \alpha .$$

④ Die Laplace-Transformierte der **verschobenen Sprungfunktion** ($\alpha > 0$)

$$S(t-\alpha) = \begin{cases} 0 & \text{für } t < \alpha \\ 1 & \text{für } t \geq \alpha \end{cases}$$



lautet für $s > 0$

$$\mathcal{L}(S(t-\alpha)) = \int_0^\infty S(t-\alpha) e^{-st} dt = \int_\alpha^\infty 1 \cdot e^{-st} dt = \frac{e^{-st}}{-s} \Big|_\alpha^\infty = \frac{e^{-\alpha s}}{s}$$

$$\Rightarrow S(t-\alpha) \circ \bullet \frac{e^{-\alpha s}}{s} \quad \text{für } s > 0 .$$

□

13.2 Inverse Laplace-Transformation

Die Berechnung der Laplace-Transformierten $F(s)$ einer Zeitfunktion $f(t)$ wird als Laplace-Transformation bezeichnet. Die Rücktransformation aus dem Bildbereich in den Zeitbereich, d.h. die Bestimmung der Zeitfunktion aus einer gegebenen Bildfunktion, heißt **inverse Laplace-Transformation**.

Für die Rücktransformation vom Bildbereich in den Zeitbereich verwendet man folgende Symbole

$$\mathcal{L}^{-1}(F(s)) = f(t) \quad \text{oder} \quad F(s) \xrightarrow{\circlearrowleft} f(t).$$

Satz über die inverse Laplace-Transformation: Ist $F(s)$ die Laplace-Transformierte einer Funktion, dann gibt es genau eine stetige Funktion $f(t)$ von höchstens exponentiellem Wachstum, so dass $\mathcal{L}(f(t)) = F(s)$. $f(t)$ ist gegeben durch

$$f(t) = \mathcal{L}^{-1}(F(s)) = \frac{1}{2\pi i} \int_{\delta-i\infty}^{\delta+i\infty} F(s) e^{st} ds.$$

Die Integration wird im Komplexen durchgeführt. Sie lässt sich im Bedarfsfall mit den mathematischen Mitteln der Funktionentheorie auswerten, auf die wir aber nicht näher eingehen werden. In der Praxis wird die Rücktransformation nicht über die Berechnung des Integrals, sondern fast ausschließlich mit Hilfe von Tabellen oder mit **MAPLE** durchgeführt. Festzuhalten ist:

Bemerkungen:

- (1) Gilt für zwei Bildfunktionen $F(s) = G(s)$, so unterscheiden sich die zugehörigen Zeitfunktionen $f(t) = \mathcal{L}^{-1}(F(s))$ und $g(t) = \mathcal{L}^{-1}(G(s))$ höchstens an Unstetigkeitsstellen von f oder g .
- (2) Zwei **stetige** Zeitfunktionen f und g stimmen überein, wenn ihre Bildfunktionen $\mathcal{L}(f(t))$ und $\mathcal{L}(g(t))$ übereinstimmen.
- (3) Da die Rücktransformation $\mathcal{L}^{-1}$ im Wesentlichen eindeutig verläuft, besteht eine eindeutige Zuordnung zwischen Bild- und Zeitfunktion, so dass jede Korrespondenz

$$f(t) \xrightarrow{\circlearrowright} F(s)$$

von links nach rechts, aber auch von rechts nach links gelesen werden kann:

$$F(s) \xrightarrow{\circlearrowleft} f(t).$$

Beispiele 13.4:

① Aus der Korrespondenz

$$t^2 S(t) \circ \bullet \frac{2}{s^3} \quad \text{für } s > 0$$

folgt durch Rücktransformation (inverse Laplace-Transformation)

$$\frac{2}{s^3} \bullet \circ t^2 S(t) \quad \text{bzw.} \quad \mathcal{L}^{-1} \left(\frac{2}{s^3} \right) = t^2 S(t) .$$

② Aus der Korrespondenz

$$S(t - \alpha) \circ \bullet \frac{1}{s} e^{-\alpha s} \quad \text{für } s > 0$$

folgt durch Rücktransformation

$$\mathcal{L}^{-1} \left(\frac{1}{s} e^{-\alpha s} \right) = S(t - \alpha) \quad \text{bzw.} \quad \frac{1}{s} e^{-\alpha s} \bullet \circ S(t - \alpha) . \quad \square$$

13.3 Zwei grundlegende Eigenschaften**13.3**

Wir zeigen, dass die Laplace-Transformation linear ist: Die Transformierte einer Superposition von Funktionen ist gleich der Überlagerung der Transformaten. Die zweite wichtige Eigenschaft ist, dass die Transformierte der Ableitung einer Funktion gleich der Transformaten der Funktion multipliziert mit s minus $f(0)$ ist. Wir nehmen im Folgenden immer an, dass die betrachteten Funktionen den Voraussetzungen des Laplaceschen Satzes genügen, so dass die Laplace-Transformaten der Funktionen definiert sind.

► 13.3.1 Linearität

Wir betrachten die Superposition zweier Zeitfunktionen f_1 und f_2 und berechnen hierzu die Laplace-Transformaten:

$$\begin{aligned} \mathcal{L}(c_1 f_1(t) + c_2 f_2(t)) &= \int_0^\infty (c_1 f_1(t) + c_2 f_2(t)) e^{-st} dt \\ &= \int_0^\infty (c_1 f_1(t) e^{-st} + c_2 f_2(t) e^{-st}) dt \\ &= c_1 \int_0^\infty f_1(t) e^{-st} dt + c_2 \int_0^\infty f_2(t) e^{-st} dt \\ &= c_1 \mathcal{L}(f_1(t)) + c_2 \mathcal{L}(f_2(t)). \end{aligned}$$

Die obigen Umformungen beruhen auf der Eigenschaft des Integrals, dass das Integral über eine Summe von Funktionen gleich der Summe der Integrale und dass konstante Faktoren vor das Integral gezogen werden dürfen. Durch die Laplace-Transformation wird der Überlagerung (= Linearkombination) von

Originalfunktionen die gleiche Überlagerung von Bildfunktionen zugeordnet:

Satz: (Additionssatz)

Es gilt: $\mathcal{L}(c_1 f_1(t) + c_2 f_2(t)) = c_1 \mathcal{L}(f_1(t)) + c_2 \mathcal{L}(f_2(t))$.

Korrespondenz: $c_1 f_1(t) + c_2 f_2(t) \longleftrightarrow c_1 F_1(s) + c_2 F_2(s)$.

Beispiele 13.5:

- ① Zur Zeitfunktion $f(t) = 2t^3 - 5t^2 + 3$ soll die Laplace-Transformierte bestimmt werden. Es gilt nach dem Additionssatz

$$\begin{aligned} F(s) &= \mathcal{L}(2t^3 - 5t^2 + 3) \\ &= 2\mathcal{L}(t^3) - 5\mathcal{L}(t^2) + 3\mathcal{L}(1) \\ &= 2\frac{3!}{s^4} - 5\frac{2!}{s^3} + 3\frac{1}{s} = \frac{12 - 10s + 3s^3}{s^4}. \end{aligned}$$

- ② Man bestimme die Laplace-Transformierte von $f(t) = 4\sin(\omega t) + 5\cos(\omega t)$. Mit dem Additionssatz gilt

$$\begin{aligned} F(s) &= 4\mathcal{L}(\sin(\omega t)) + 5\mathcal{L}(\cos(\omega t)) \\ &= 4\frac{\omega}{s^2 + \omega^2} + 5\frac{s}{s^2 + \omega^2} = \frac{5s + 4\omega}{s^2 + \omega^2}. \end{aligned} \quad \square$$

Die Linearität der Laplace-Transformation überträgt sich durch die Korrespondenz auch auf die Rücktransformation:

Satz: (Additionssatz der inversen Laplace-Transformation)

$$\mathcal{L}^{-1}(c_1 F_1(s) + c_2 F_2(s)) = c_1 \mathcal{L}^{-1}(F_1(s)) + c_2 \mathcal{L}^{-1}(F_2(s)).$$

Beispiele 13.6:

- ① Man bestimme die Zeitfunktion zu

$$F(s) = \frac{3s + 8}{s^2 + 16}.$$

Mit der Zerlegung der Bildfunktion in die Partialbrüche

$$F(s) = 3\frac{s}{s^2 + 4^2} + 2\frac{4}{s^2 + 4^2}$$

erhält man unter Verwendung von Beispiel 13.5 ② die zugehörige Zeitfunktion

$$f(t) = 3 \cos(4t) + 2 \sin(4t) .$$

② Gesucht ist die zu

$$F(s) = \frac{5s^2 + 3s + 8}{s^3}$$

gehörende Zeitfunktion $f(t)$. Durch Zerlegung der Bildfunktion in Partialbrüche

$$F(s) = 5 \frac{1}{s} + 3 \frac{1}{s^2} + 4 \frac{1}{s^3}$$

folgt

$$f(t) = 5 + 3t + 4t^2 . \quad \square$$

➤ 13.3.2 Laplace-Transformierte der Ableitung

Für die Anwendung der Laplace-Transformation auf Differenzialgleichungen muss man die Laplace-Transformierte der Ableitung einer Funktion berechnen können:

Satz: (Laplace-Transformation der Ableitung)

Seien $f, f' : [0, \infty) \rightarrow \mathbb{R}$ stetig und von höchstens exponentiellem Wachstum. Dann gilt:

$$\mathcal{L}(f'(t)) = s \mathcal{L}(f(t)) - f(0).$$

Beweis: Mit partieller Integration bestimmt man die Laplace-Transformierte von f' :

$$\begin{aligned} \mathcal{L}(f'(t)) &= \int_0^\infty f'(t) e^{-st} dt = [f(t) e^{-st}]_0^\infty + s \int_0^\infty f(t) e^{-st} dt \\ &= \lim_{T \rightarrow \infty} f(T) e^{-sT} - f(0) + s \mathcal{L}(f(t)). \end{aligned}$$

Da $f(t)$ von höchstens exponentiellem Wachstum ist, gilt nach Satz 13.1

$$\lim_{T \rightarrow \infty} f(T) e^{-sT} = 0$$

und folglich ist

$$\mathcal{L}(f'(t)) = s \mathcal{L}(f(t)) - f(0) . \quad \square$$

Der Differenziation im Originalbereich entspricht die Multiplikation mit s und Subtraktion von $f(0)$ im Bildbereich. An die Stelle einer komplizierten Rechenoperation im Originalbereich tritt also eine einfache Multiplikation im Bildbereich. Wiederholtes Anwenden des Ableitungssatzes führt induktiv auf die Laplace-Transformierte der n -ten Ableitung:

Satz: (Laplace-Transformation der n -ten Ableitung)

Seien $f, f', \dots, f^{(n)} : [0, \infty) \rightarrow \mathbb{R}$ stetig und von höchstens exponentiellem Wachstum, dann gilt

$$\mathcal{L}(f^{(n)}(t)) = s^n \mathcal{L}(f(t)) - s^{n-1} f(0) - s^{n-2} f'(0) - \dots - f^{(n-1)}(0).$$

Die Bildfunktion der n -ten Ableitung von $f(t)$ ist gleich der Laplace-Transformierten von $f(t)$ multipliziert mit s^n minus einem Polynom $(n-1)$ -ten Grades in s , dessen Koeffizienten durch die Werte der Funktion f sowie deren Ableitungen an der Stelle $t = 0$ bestimmt sind.

Bemerkung: Von $f(t)$ wird vorausgesetzt, dass $f(t) = 0$ für $t < 0$. Damit ergibt sich der linksseitige Grenzwert an der Stelle $t = 0$ zu Null, $f(-0) = 0$, sowie $f'(-0) = \dots = f^{(n-1)}(-0) = 0$. Besitzt die Funktion $f(t)$ und deren Ableitungen bei $t = 0$ eine Sprungstelle, so müssen für die Anfangswerte $f(0), f'(0), \dots, f^{(n-1)}(0)$ in den Ableitungssätzen die rechtsseitigen Grenzwerte $f(+0), f'(+0), \dots, f^{(n-1)}(+0)$ verwendet werden.

⚠ Achtung: Diese Unterscheidung zwischen einem Wert an einer Stelle und dem Grenzwert bei Annäherung an diese Stelle ist wichtig. Mit diesen Grenzwerten sind Anfangswerte gemeint, von denen die Funktionen für $t > 0$ ausgehen und die dann bei $t = 0$ einen stetigen Anschluss der Funktion gewährleisten.

Beispiele 13.7 (Bestimmung der LT über den Ableitungssatz):

- ① Zur Funktion $f(t) = e^{at}$ soll die Laplace-Transformierte bestimmt werden. Die Anwendung des Ableitungssatzes auf

$$f'(t) = a e^{at}$$

führt mit $f(0) = 1$ auf

$$\mathcal{L}(a e^{at}) = s \mathcal{L}(e^{at}) - 1 \Rightarrow a \mathcal{L}(e^{at}) = s \mathcal{L}(e^{at}) - 1.$$

$$\Rightarrow \mathcal{L}(e^{at}) = \frac{1}{s-a} \quad \text{bzw.} \quad e^{at} \circ \bullet \frac{1}{s-a} \quad \text{für } s > a.$$

- ② Die trigonometrischen Funktionen $\sin(\omega t)$ und $\cos(\omega t)$ erfüllen die Differenzialgleichungen

$$f''(t) + \omega^2 f(t) = 0 \quad \text{mit} \quad \begin{cases} f(0) = 0, & f'(0) = \omega & \text{für } \sin(\omega t) \\ f(0) = 1, & f'(0) = 0 & \text{für } \cos(\omega t). \end{cases}$$

Ihre Laplace-Transformierten werden mit dem Ableitungssatz bestimmt. Aus der Differenzialgleichung folgt

$$\mathcal{L}(f''(t)) + \omega^2 \mathcal{L}(f(t)) = \mathcal{L}(0) = 0.$$

Für die Sinusfunktion ergibt sich unter Berücksichtigung der Anfangsbedingungen

$$s^2 \mathcal{L}(f(t)) - \omega = \omega^2 \mathcal{L}(f(t)) = 0$$

$$\Rightarrow \quad \mathcal{L}(\sin(\omega t)) = \frac{\omega}{s^2 + \omega^2} \quad \text{bzw.} \quad \sin(\omega t) \longleftrightarrow \frac{\omega}{s^2 + \omega^2}.$$

- ③ Analog findet man für die Kosinusfunktion

$$s^2 \mathcal{L}(f(t)) - s \cdot 1 + \omega^2 \mathcal{L}(f(t)) = 0$$

$$\Rightarrow \quad \mathcal{L}(\cos(\omega t)) = \frac{s}{s^2 + \omega^2} \quad \text{bzw.} \quad \cos(\omega t) \longleftrightarrow \frac{s}{s^2 + \omega^2}.$$

- ④ In Verallgemeinerung des vorhergehenden Beispiels wird die Laplace-Transformierte der Funktion $f(t) = \sin(\omega t + \varphi)$ gebildet. Diese Funktion genügt ebenfalls der Differenzialgleichung

$$f''(t) + \omega^2 f(t) = 0,$$

jedoch mit den Anfangsbedingungen $f(0) = \sin \varphi$ und $f'(0) = \omega \cos \varphi$. Mit der Laplace-Transformation folgt aus der Differenzialgleichung

$$\mathcal{L}(f''(t)) + \omega^2 \mathcal{L}(f(t)) = 0,$$

$$\hookrightarrow \quad s^2 \mathcal{L}(f(t)) - s \sin \varphi - \omega \cos \varphi + \omega^2 \mathcal{L}(f(t)) = 0.$$

$$\Rightarrow \quad \left\{ \begin{array}{l} \boxed{\mathcal{L}(\sin(\omega t + \varphi)) = \frac{s \sin \varphi + \omega \cos \varphi}{s^2 + \omega^2}} \\ \text{bzw.} \quad \boxed{\sin(\omega t + \varphi) \longleftrightarrow \frac{s \sin \varphi + \omega \cos \varphi}{s^2 + \omega^2}} \end{array} \right.$$

Mit $\varphi = \psi + \frac{\pi}{2}$ folgt

$$\left\{ \begin{array}{l} \mathcal{L}(\cos(\omega t + \psi)) = \frac{s \cos \psi - \omega \sin \psi}{s^2 + \omega^2} \\ \text{bzw.} \quad \cos(\omega t + \psi) \longleftrightarrow \frac{s \cos \psi - \omega \sin \psi}{s^2 + \omega^2}. \end{array} \right.$$

- ⑤ Die Hyperbelfunktionen $\sinh(\omega t)$ und $\cosh(\omega t)$ sind Lösungen der Differentialgleichung

$$f''(t) - \omega^2 f(t) = 0 \quad \text{mit} \quad \begin{cases} f(0) = 0, & f'(0) = \omega & \text{für } \sinh(\omega t) \\ f(0) = 1, & f'(0) = 0 & \text{für } \cosh(\omega t). \end{cases}$$

$$\Rightarrow \mathcal{L}(f''(t)) - \omega^2 \mathcal{L}(f(t)) = 0.$$

Mit den Anfangsbedingungen ergibt sich für $\sinh(\omega t)$

$$s^2 \mathcal{L}(f(t)) - \omega = \omega^2 \mathcal{L}(f(t)) = 0$$

$$\Rightarrow \mathcal{L}(\sinh(\omega t)) = \frac{\omega}{s^2 - \omega^2} \quad \text{bzw.} \quad \sinh(\omega t) \longleftrightarrow \frac{\omega}{s^2 - \omega^2}.$$

- ⑥ Analog erhält man

$$\mathcal{L}(\cosh(\omega t)) = \frac{s}{s^2 - \omega^2} \quad \text{bzw.} \quad \cosh(\omega t) \longleftrightarrow \frac{s}{s^2 - \omega^2}. \quad \square$$

13.4 Methoden der Rücktransformation

Die Rücktransformation, d.h. die Ermittlung der Funktion $f(t)$ im Zeitbereich zu einer gegebenen Bildfunktion $F(s)$, ist der schwierigste und aufwändigste Schritt. Die komplexe Umkehrformel wird in den seltensten Fällen benutzt, da sie detaillierte Kenntnisse der Funktionentheorie voraussetzt. Die im Folgenden angegebenen Verfahren führen jedoch in den meisten Fällen zum Erfolg.

⑤ Der Gebrauch von Tabellen

Die gebräuchlichste Methode zur Gewinnung der Originalfunktion zu gegebener Bildfunktion ist die Verwendung von Tabellen. Der Vorteil von Tabellen besteht darin, dass für viele vorkommende Fälle die Rechnung schon einmal durchgeführt wurde und in korrespondierenden Funktionenpaaren ihren Niederschlag gefunden hat.

In der folgenden Tabelle sind von vielen elementaren Funktionen die zugehörigen Laplace-Transformierten angegeben. Im Funktionsausdruck der Transformierten der Funktion t^p für $p > -1$ taucht die sog. Gamma-Funktion $\Gamma(p+1)$ auf. In einem separaten Abschnitt über die [Laplace-Transformation mit MAPLE](#) wird auf diese Funktion eingegangen.

Tabelle 13.1: Laplace-Transformierte elementarer Funktionen

Zeitfunktion $f(t)$	Bildfunktion $F(s)$	Zeitfunktion $f(t)$	Bildfunktion $F(s)$
$S(t)$	$\frac{1}{s} \quad s > 0$	$e^{at} t^n$	$\frac{n!}{(s-a)^{n+1}}$
t	$\frac{1}{s^2} \quad s > 0$	$\sin(\omega t + b)$	$\frac{(\sin b) s + \omega (\cos b)}{s^2 + \omega^2}$
t^n	$\frac{n!}{s^{n+1}} \quad s > 0$	$\cos(\omega t + b)$	$\frac{(\cos b) s - \omega (\sin b)}{s^2 + \omega^2}$
e^{at}	$\frac{1}{s-a} \quad s > a$	$e^{bt} \sinh(at)$	$\frac{a}{(s-b)^2 - a^2}$
$\sin(\omega t)$	$\frac{\omega}{s^2 + \omega^2}$	$e^{bt} \cosh(at)$	$\frac{s-b}{(s-b)^2 - a^2}$
$\cos(\omega t)$	$\frac{s}{s^2 + \omega^2}$	$t \sin(\omega t)$	$\frac{2\omega s}{(s^2 + \omega^2)^2}$
$t^p, p > -1$	$\frac{\Gamma(p+1)}{s^{p+1}}, s > 0$	$t \cos(\omega t)$	$\frac{s^2 - \omega^2}{(s^2 + \omega^2)^2}$
$e^{at} \sin(\omega t)$	$\frac{\omega}{(s-a)^2 + \omega^2}$		
$e^{at} \cos(\omega t)$	$\frac{s-a}{(s-a)^2 + \omega^2}$		
$\sinh(at)$	$\frac{a}{s^2 - a^2}$		
$\cosh(at)$	$\frac{s}{s^2 - a^2}$		

➤ **Die Methode der Partialbruchzerlegung, mit [MAPLE-Worksheet](#)**

Bei vielen Anwendungen der Laplace-Transformation sind die auftretenden Bildfunktionen gebrochenrationale Funktionen

$$F(s) = \frac{Z(s)}{N(s)}$$

der Variablen s . Zähler $Z(s)$ und Nenner $N(s)$ sind Polynome mit $\text{grad } Z(s) < \text{grad } N(s)$, und die Polynome $Z(s)$ und $N(s)$ besitzen reelle Koeffizienten. Je nach Art der auftretenden Polstellen von $F(s)$ ergeben sich für die Partialbruchzerlegung die folgenden Fälle:

- (1) **Bildfunktion mit nur einfachen reellen Polen:** In diesem Fall stellt sich die Bildfunktion $F(s)$ dar als

$$F(s) = \frac{a_1}{s-s_1} + \frac{a_2}{s-s_2} + \dots + \frac{a_n}{s-s_n}.$$

Mit der Korrespondenz

$$\frac{a_i}{s-s_i} \longleftrightarrow a_i e^{s_i t}$$

lässt sich jeder Summand zurück transformieren.

- (2) **Bildfunktion mit mehrfachen reellen Polen:** Ist s_0 eine reelle Polstelle der Vielfachheit k , so gilt für sie die Zerlegung

$$\frac{1}{(s-s_0)^k} = \frac{b_1}{s-s_0} + \frac{b_2}{(s-s_0)^2} + \dots + \frac{b_k}{(s-s_0)^k}.$$

Mit dem Dämpfungssatz (siehe Abschnitt auf der CD-Rom) und $t^i \longleftrightarrow \frac{i!}{s^{i+1}}$ gilt die Korrespondenz

$$\frac{b_i}{(s-s_0)^i} \longleftrightarrow b_i e^{s_0 t} \frac{t^{i-1}}{(i-1)!}.$$

Wieder lässt sich mit diesen Korrespondenzen die zugehörige Zeitfunktion ermitteln.

- (3) **Bildfunktion mit einfachen komplexen Polen:** Ist $s_0 = a + ib$ eine komplexe Nullstelle von $N(s)$, so ist mit s_0 auch die komplex konjugierte Zahl $s_0^* = a - ib$ eine Nullstelle (*Fundamentalsatz der Algebra*, §5.2.7). Es gilt dann mit

$$(s-s_0)(s-s_0^*) = s^2 - s(s_0 + s_0^*) + s_0 s_0^* = s^2 - 2as + a^2 + b^2$$

$$\Rightarrow F(s) = \frac{c_1 s + c_2}{s^2 - 2as + a^2 + b^2} + P(s) = F_0(s) + P(s),$$

wobei $P(s)$ die Summe der restlichen Partialbrüche darstellt. Die zu $F_0(s)$ gehörende Korrespondenz lautet mit dem Verschiebungssatz

$$f_0(t) = e^{at} \left[c_1 \cos(bt) + \frac{c_2 + ac_1}{b} \sin(bt) \right].$$

- (4) **Bildfunktion mit k-fachen komplexen Polen:** Ist $s_0 = a + ib$ eine k -fache komplexe Nullstelle von $N(s)$, dann ist auch s_0^* eine k -fache Nullstelle und der Ansatz unter (3) ist folgendermaßen zu modifizieren

$$\frac{1}{(s-s_0)^k} = \frac{c_1 s + d_1}{(s^2 - 2as + a^2 + b^2)^1} + \frac{c_2 s + d_2}{(s^2 - 2as + a^2 + b^2)^2} + \dots + \frac{c_k s + d_k}{(s^2 - 2as + a^2 + b^2)^k}.$$

Wieder muss der Verschiebungssatz angewendet werden, um jeden einzelnen Summanden analog (3) zurückzuführen.

13.5 Anwendungen der Laplace-Transformation

Lineare Differenzialgleichungen und Differenzialgleichungssysteme mit Anfangsbedingungen werden elegant mit der Laplace-Transformation gelöst. In diesem Abschnitt werden wir exemplarisch den RL-Kreis und ein elektrisches Netzwerk diskutieren.

Hinweis: In einem separaten Abschnitt auf der CD-Rom werden **zusätzliche Anwendungsbeispiele** behandelt und mit MAPLE gelöst.

Anwendungsbeispiel 13.8 (RL-Stromkreis, mit MAPLE-Worksheet).

Gegeben ist der in Abb. 13.4 dargestellte RL-Stromkreis. Zur Zeit $t = 0$ wird eine konstante Spannung U_0 angelegt. Gesucht ist der Verlauf des Stromes $I(t)$ für $I(0) = I_0$.

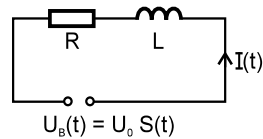


Abb. 13.4. RL-Kreis

Es gilt nach dem Maschensatz

$$L I'(t) + R I(t) = U_0 S(t).$$

1. Schritt: Man wende die Laplace-Transformation auf die Differenzialgleichung an. Wegen der Linearität gilt:

$$\mathcal{L}(L I'(t) + R I(t)) = \mathcal{L}(U_0 S(t))$$

$$\mathcal{L}(I'(t)) + \frac{R}{L} \mathcal{L}(I(t)) = \frac{1}{L} U_0 \mathcal{L}(S(t)).$$

2. Schritt: Man ersetze die Laplace-Transformierte der Ableitung unter Berücksichtigung der Anfangsbedingung und berechne die Laplace-Transformierte der Inhomogenität.

$$s \mathcal{L}(I(t)) - I_0 + \frac{R}{L} \mathcal{L}(I(t)) = \frac{U_0}{L} \frac{1}{s}.$$

3. Schritt: Man löse nach der Laplace-Transformierten $\mathcal{L}(I(t))$ auf.

$$(s + \frac{R}{L}) \mathcal{L}(I(t)) = \frac{U_0}{L} \frac{1}{s} + I_0$$

$$\mathcal{L}(I(t)) = \frac{U_0}{L} \frac{1}{s + \frac{R}{L}} \frac{1}{s} + \frac{I_0}{s + \frac{R}{L}}.$$

4. Schritt: Man suche die zugehörige Zeitfunktion. Mit einer Partialbruchzerlegung folgt

$$\mathcal{L}(I(t)) = \frac{U_0}{R} \left[\frac{1}{s} - \frac{1}{s + \frac{R}{L}} \right] + I_0 \frac{1}{s + \frac{R}{L}}.$$

Berücksichtigt man die Korrespondenzen

$$\begin{aligned}\mathcal{L}(1) &= \frac{1}{s} \\ \mathcal{L}(e^{-\frac{R}{L}t}) &= \frac{1}{s + \frac{R}{L}}\end{aligned}$$

folgt

$$I(t) = \frac{U_0}{R} \left(1 - e^{-\frac{R}{L}t} \right) + I_0 e^{-\frac{R}{L}t}.$$

In Abb. 13.5 sind für unterschiedliche Startwerte $I(0)$ der Verlauf des Stromes gezeichnet.

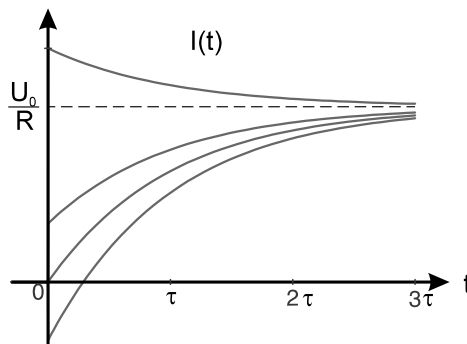


Abb. 13.5. Verlauf des Stromes für verschiedene Anfangswerte I_0 mit $\tau = R/L$

Unabhängig vom Startwert nähert sich der Strom $I(t)$ für lange Zeiten dem Wert $\frac{U_0}{R}$ an. Die Zeitkonstante ist $\tau = \frac{R}{L}$. Sie bestimmt, wie schnell sich der Endzustand einstellt.



Hinweis: Auf der CD-Rom befindet sich ein [MAPLE-Worksheet](#), welches dieselbe Schaltung mit einer Wechselspannung $U_0(t) = U_0 \sin(\omega t)$ behandelt. □

Anwendungsbeispiel 13.9 (Elektrisches Netzwerk mit konstanter Spannung, mit MAPLE-Worksheet).

Gegeben sei das im Einführungsbeispiel 13.1 diskutierte elektrische Netzwerk mit konstanter Spannung $U(t) = 120 \text{ V}$:

$$\begin{aligned} 20(I_1(t) + I_2(t)) + 2I_1'(t) + 10I_1(t) &= 120 \\ -10I_1(t) + 20I_2(t) - 2I_1'(t) + 4I_2'(t) &= 0 \end{aligned}$$

$$I_1(0) = I_2(0) = 0.$$

Mit den gleichen Schritten wie bei linearen Differenzialgleichungen verwendet man auch bei Systemen die Laplace-Transformation, um die Lösung direkt zu bestimmen:

1. Schritt: Anwenden der Laplace-Transformation auf das Differenzialgleichungssystem.

$$\begin{aligned} 20\mathcal{L}(I_1) + 20\mathcal{L}(I_2) + 2\mathcal{L}(I_1') + 10\mathcal{L}(I_1) &= 120\mathcal{L}(1) \\ -10\mathcal{L}(I_1) + 20\mathcal{L}(I_2) - 2\mathcal{L}(I_1') + 4\mathcal{L}(I_2') &= 0. \end{aligned}$$

2. Schritt: Ersetzen der Laplace-Transformierten der Ableitungen.

$$\begin{aligned} 30\mathcal{L}(I_1) + 20\mathcal{L}(I_2) + 2\mathcal{L}(I_1) \cdot s &= 120 \cdot \frac{1}{s} \\ -10\mathcal{L}(I_1) + 20\mathcal{L}(I_2) - 2\mathcal{L}(I_1) \cdot s + 4\mathcal{L}(I_2) \cdot s &= 0, \end{aligned}$$

da die Anfangsbedingungen $I_1(0) = I_2(0) = 0$ verschwinden. Man erhält das folgende lineare Gleichungssystem für $\mathcal{L}(I_1)$ und $\mathcal{L}(I_2)$:

$$\begin{aligned} (30 + 2s)\mathcal{L}(I_1) + 20\mathcal{L}(I_2) &= 120 \cdot \frac{1}{s} \\ (-10 - 2s)\mathcal{L}(I_1) + (20 + 4s)\mathcal{L}(I_2) &= 0. \end{aligned}$$

3. Schritt: Lösen des linearen Gleichungssystems z.B. mit dem Gauß-Algorithmus:

$$\mathcal{L}(I_1) = \frac{60}{s(s+20)}; \quad \mathcal{L}(I_2) = \frac{30}{s(s+20)}.$$

4. Schritt: Suchen der zugehörigen Zeitfunktion. Eine Partialbruchzerlegung für die beiden Laplace-Transformierten liefert

$$\mathcal{L}(I_1) = 3 \left(\frac{1}{s} - \frac{1}{s+20} \right) \quad \text{und} \quad \mathcal{L}(I_2) = \frac{3}{2} \left(\frac{1}{s} - \frac{1}{s+20} \right)$$

$$\Rightarrow I_1(t) = 3(1 - e^{-20t}) \quad \text{und} \quad I_2(t) = \frac{3}{2}(1 - e^{-20t})$$

$$\Rightarrow I(t) = I_1(t) + I_2(t) = \frac{9}{2} (1 - e^{-20t}).$$

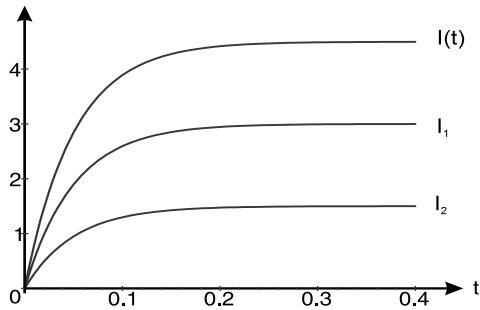


Abb. 13.6. Ströme im elektrischen Netzwerk bei konstanter Spannung

Im Abb. 13.6 sind die Teilströme $I_1(t)$ und $I_2(t)$ zusammen mit dem Gesamtstrom $I(t)$ als Funktion der Zeit dargestellt.



Hinweis: Auf der CD-Rom befindet sich ein [MAPLE-Worksheet](#), welches das elektrische Netzwerk simuliert, wenn die Eingangsspannung eine Wechselspannung oder eine Rechteckspannung ist. $\square$

MAPLE-Worksheets zu Kapitel 13



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 13 mit MAPLE zur Verfügung.

- Laplace-Transformation mit MAPLE
- Inverse Laplace-Transformation mit MAPLE
- Methoden der Rücktransformation mit MAPLE
- Eigenschaften der Laplace-Transformation mit MAPLE
- Beispiel RC-Kreis mit MAPLE
- Beispiel RL-Kreis mit MAPLE
- Beispiel elektrisches Netzwerk mit MAPLE
- Beispiel LDG-System mit MAPLE

13.6 Aufgaben zur Laplace-Transformation

13.1 Man berechne jeweils die Laplace-Transformierte von

a) $3e^{-4t}$ b) $2t^2$ c) $4\cos(5t)$ d) $\sin(\pi t)$ e) $\frac{-3}{\sqrt{t}}$

13.2 Berechnen Sie die Laplace-Transformierten von

a) $3t^4 - 2t^{\frac{3}{2}} + 6$ b) $5\sin(2t) - 3\cos(2t)$
 c) $3\sqrt[3]{t} - 4e^{2t}$ d) $\frac{1}{t^2}$

13.3 Bestimmen Sie die Laplace-Transformierten der Zeitfunktionen (mit MAPLE)

a) $f(t) = \begin{cases} A & 0 \leq t \leq t_0 \\ Ae^{-2(t-t_0)} & t > t_0 \end{cases}$ b) $f(t) = \begin{cases} 0 & \text{für } t < a \\ A & \text{für } a < t < b \\ 0 & \text{für } t > b \end{cases}$
 c) $f(t) = \begin{cases} t & \text{für } 0 \leq t \leq 3 \\ 3 & \text{für } t > 3 \end{cases}$ d) $f(t) = \begin{cases} \sin t & \text{für } t \leq \pi \\ 0 & \text{für } t > \pi \end{cases}$

13.4 Berechnen Sie

a) $\mathcal{L}^{-1}\left\{\frac{5}{s+2}\right\}$ b) $\mathcal{L}^{-1}\left\{\frac{4s-3}{s^2+4}\right\}$ c) $\mathcal{L}^{-1}\left(\frac{2s-5}{s^2}\right)$
 d) $\mathcal{L}^{-1}\left(\frac{1}{s^k}\right)_{k>0}$ e) $\mathcal{L}^{-1}\left\{\frac{4-5s}{s^{\frac{3}{2}}}\right\}$ f) $\mathcal{L}^{-1}\left\{\frac{1}{s^2+2s}\right\}$

13.5 Wenden Sie die Sätze der Laplace-Transformation an, um die inverse Laplace-Transformierte zu berechnen und überprüfen Sie die Ergebnisse mit MAPLE

a) $\mathcal{L}^{-1}\left\{\frac{2s+3}{s^2-2s+5}\right\}$ b) $\mathcal{L}^{-1}\left\{\frac{e^{-2s}}{s^2}\right\}$
 c) $\mathcal{L}^{-1}\left\{\frac{e^{-5s}}{s^4}\right\}$ d) $\mathcal{L}^{-1}\left(\frac{1-e^{-2s}}{s^3}\right)$

13.6 Man berechne durch Partialbruchzerlegung der Bildfunktion die zugehörige Zeitfunktion

a) $F(s) = \frac{2s^2-4}{(s-2)(s+1)(s-3)}$ b) $F(s) = \frac{3s+1}{(s-1)(s^2+1)}$
 c) $F(s) = \frac{5s^2-15s+7}{(s+1)(s-2)^2}$ d) $F(s) = \frac{3s^2-7s+6}{(s-1)^3}$

13.7 Lösen Sie die Differenzialgleichung

$$y''(t) + y(t) = S(t)$$

mit den Anfangsbedingungen $y(0) = 1$, $y'(0) = 0$ mit Hilfe der Laplace-Transformation.

13.8 Lösen Sie die Differenzialgleichung 4. Ordnung

$$y^{(4)}(t) + 2y''(t) + y(t) = \sin(t) S(t)$$

mit $y(0) = 1$, $y'(0) = -2$, $y''(0) = 3$, $y'''(0) = 0$ mit Hilfe der Laplace-Transformation.

Anwendungen der Laplace-Transformation

- 13.9 R, L, U_B sind mit einem Schalter S in Reihe geschaltet. Der Schalter ist zunächst geschlossen und wird zur Zeit $t = 0$ geöffnet: $I(t = 0) = I_0$. Lösen Sie die Differenzialgleichung

$$R I(t) + L \frac{d}{dt} I(t) = 0 \quad , \quad I(0) = I_0$$

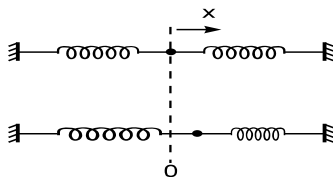
mit der Laplace-Transformation.

- 13.10 R, L, U_B sind mit einem Schalter S in Reihe geschaltet. Der Schalter ist zunächst geöffnet und wird zur Zeit $t = 0$ geschlossen ($I(t = 0) = 0$). Wie verhält sich $I(t)$, wenn $U_B(t) = U_0 \sin(\omega t)$? Lösen Sie die Differenzialgleichung

$$R I(t) + L \frac{d}{dt} I(t) = U_B(t)$$

mit der Laplace-Transformation.

- 13.11 Ein Teilchen bewegt sich auf der x -Achse und wird zum Ursprung 0 mit einer Kraft, die proportional zu der momentanen Entfernung von 0 ist, hingezogen. Wenn das Teilchen aus der Ruhe heraus bei $x = 5 \text{ cm}$ startet und erstmals nach 2 Sekunden die Stelle $x = 2.5 \text{ cm}$ erreicht, berechne man
- die Lage zu einer beliebigen Zeit t nach dem Start,
 - die Größe seiner Geschwindigkeit bei $x = 0$,
 - die momentane Beschleunigung.



- 13.12 Die Lage eines Teilchens, das sich entlang der x -Achse bewegt, ist durch die Gleichung

$$\frac{d^2}{dt^2} x(t) + 4 \frac{d}{dt} x(t) + 8 x(t) = 20 \cos(2t)$$

bestimmt. Wenn das Teilchen aus der Ruhe heraus bei $x = 0$ startet, berechne man x als Funktion von t .

- 13.13 Ein 2 kg schweres Gewicht hängt an einer Feder mit Federkonstanten $D = 200 \frac{\text{N}}{\text{m}}$ in Ruhe. Man berechne die Lage des Gewichtes zu einer beliebigen Zeit t , wenn seine Dämpfungskraft 40 mal der momentanen Geschwindigkeit ist.
- 13.14 Ein Gewicht an einer vertikalen Feder ist einer erzwungenen Schwingung ausgesetzt. Die Auslenkung aus der Ruhelage wird beschrieben durch

$$\frac{d}{dt^2} x(t) + 4 x(t) = 8 \sin(\omega t) \quad (\omega > 0) \quad .$$

Wenn $x(0) = 0$ und $\dot{x}(0) = 0$, berechne man x als Funktion von t und die Periode der von außen wirkenden Kraft, für welche Resonanz auftritt.

Kapitel 14
Fourier-Reihen

14

14	Fourier-Reihen	583
14.1	Einführung	585
14.2	Bestimmung der Fourier-Koeffizienten.....	587
14.3	Fourier-Reihen für 2π -periodische Funktionen	590
14.4	Fourier-Reihen für p -periodische Funktionen	597
14.5	Fourier-Reihen für komplexwertige Funktionen	605
14.6	Aufgaben zu Fourier-Reihen	610

Zusätzliche Abschnitte auf der CD-Rom

14.7	Zusammenstellung elementarer Fourier-Reihen	cd
14.8	MAPLE: Fourier-Reihen	cd
14.8.1	Berechnung der Fourier-Koeffizienten mit MAPLE.....	cd
14.8.2	Analyse T -periodischer Signale mit MAPLE.....	cd
14.8.3	MAPLE-Prozedur zur Berechnung der Fourier-Koeffizienten	cd
14.8.4	Berechnung der komplexen Fourier-Koeffizienten	cd
14.8.5	Zusammenstellung der MAPLE-Befehle	cd

14 Fourier-Reihen

Bei der Analyse periodischer Signale benötigt man die Darstellung des Signals in Form einer Fourier-Reihe

$$f(t) = a_0 + \sum_{n=1}^{\infty} a_n \cos(\omega_n t) + \sum_{n=1}^{\infty} b_n \sin(\omega_n t).$$

Denn durch eine solche Zerlegung des Signals in seine harmonischen Bestandteile geht hervor, welche Frequenzen mit welchen Amplituden im Signal enthalten sind.

Nach einer Einführung werden wir in 14.2 die Formeln für die Fourier-Reihe und die Fourier-Koeffizienten 2π -periodischer Funktionen aufstellen und in 14.3 auf Beispiele anwenden. In 14.4 übertragen wir die Formeln auf p -periodische Funktionen und in 14.5 gehen wir zur komplexen Formulierung über. Diese Formulierung stellt dann den Übergang zur Fourier-Transformation dar.

Hinweis: Auf der CD-Rom befinden sich zusätzlich weitere [Beispiele](#) mit [MAPLE](#), die Prozedur [fourier_reihen](#) zum numerischen Berechnen der Fourier-Koeffizienten und der Visualisierung der punktweisen Konvergenz der Reihe an die Funktion.

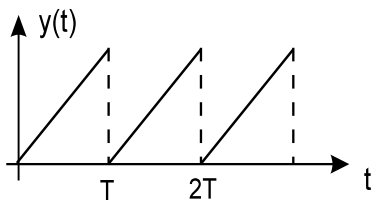
14.1 Einführung

14.1

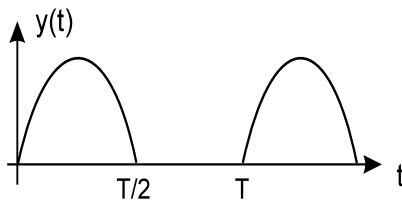
In der Physik lassen sich einfache, zeitlich periodische Vorgänge wie z.B. die Schwingung eines Federpendels oder Wechselspannungen durch die allgemeine Sinusfunktion der Form

$$y(t) = A \sin(\omega t + \varphi)$$

beschreiben. Man nennt diese Darstellung *harmonische Schwingung* mit Frequenz ω und Amplitude A . Sie treten vor allem bei der Beschreibung von schwingenden Saiten, Membranen, Pendel, elektromagnetischen Schwingungen, Schall- und Wellenausbreitung usw. auf. Häufig kommen aber auch Vorgänge vor, die zwar periodisch aber nicht mehr sinusförmig sind. Beispiele hierfür sind Kippschwingungen (Kippspannung, Kreidequietschen) oder der Sinusimpuls eines Gleichrichters.



Kippschwingung



Sinusimpuls eines Gleichrichters

Wenn man etwa an die Kippschwingung denkt, die zum Kreidequietschen führt, ist man an den dominanten Frequenzen und zugehörigen Amplituden interessiert. Wenn man bei einem Klavier die drei Töne c^1 , g^1 , e^2 gleichzeitig anschlägt und die Stärke der Anschläge so wählt, dass die in normierten Einheiten am Ohr erzeugten Überdrücke gleich 1.273, 0.424 und 0.255 sind, dann ist der Gesamtdruck $p(t)$ am Ohr gegeben durch deren Überlagerung:

$$p(t) = 1.273 \sin(2\pi\nu_1 t) + 0.424 \sin(2\pi\nu_3 t) + 0.255 \sin(2\pi\nu_5 t)$$

mit $\nu_1 = 128 \text{ Hz}$ (c^1), $\nu_3 = 3\nu_1 = 384 \text{ Hz}$ (g^1) und $\nu_5 = 5\nu_1 = 640 \text{ Hz}$ (e^2).

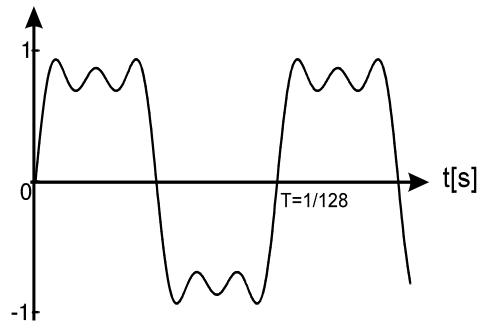


Abb. 14.1. Schalldruck am Ohr

Geübte Ohren können aufgrund des am Ohr erzeugten Überdrucks analysieren, welche Frequenzen (c^1 , g^1 , e^2) in dem Ton enthalten sind. Das Ohr unterzieht im Zusammenwirken mit dem Gehirn eine Analyse des periodischen Signals: Es zerlegt das Signal in Einzelfrequenzen (= *Signalanalyse*).

Von generellem Interesse bei der Signalanalyse ist die Zerlegung eines periodischen Zeitsignals in Grundschwingung und Oberschwingungen mit den zugehörigen Amplituden. Es stellt sich heraus, dass nahezu **jede** periodische Funktion $y(t)$ sich darstellen lässt als Überlagerung unendlich vieler harmonischer Schwingungen.

Den mathematischen Zusammenhang zwischen **periodischem** Signal und dessen Zerlegung in Grund- und Oberschwingungen mit zugehörigen Amplituden stellt die **Fourier-Reihe** dar:

$$y(t) = a_0 + \sum_{n=1}^{\infty} a_n \cos(n\omega_0 t) + \sum_{n=1}^{\infty} b_n \sin(n\omega_0 t) ,$$

wenn $\omega_0 = \frac{2\pi}{T}$ und T die Periodendauer der Funktion $y(t)$.

Die Entwicklung einer periodischen Funktion in eine Fourier-Reihe bezeichnet man als **Fourier-Analyse**. ω_0 ist die Grundschiwingung und $n \omega_0$ sind die Oberschwingungen. Die Koeffizienten $a_0, a_1, a_2, \dots; b_1, b_2, \dots$ heißen Fourier-Koeffizienten und geben die Amplituden der einzelnen Frequenzkomponenten an.

14.2 Bestimmung der Fourier-Koeffizienten

Zur Bestimmung der Amplituden in der Fourier-Zerlegung gehen wir von einer 2π -periodischen Funktion f aus (siehe z.B. Abb. 14.2)

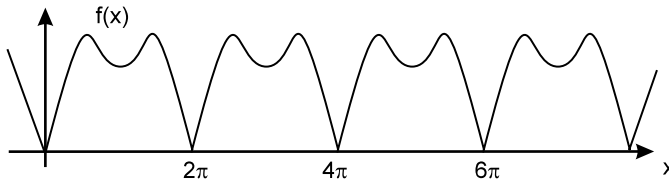


Abb. 14.2. 2π -periodische Funktion

und wählen den Ansatz:

$$f(x) = a_0 + \sum_{n=1}^{\infty} a_n \cos(nx) + \sum_{n=1}^{\infty} b_n \sin(nx). \quad (*)$$

Zur formalen Bestimmung der Koeffizienten $a_0; a_1, a_2, \dots; b_1, b_2, \dots$ benötigen wir die in Tabelle 14.1 zusammengestellten Integrale.

Tabelle 14.1: Zusammenstellung elementarer Sinus- und Kosinusintegrale

(1)	$\int_0^{2\pi} \sin(nx) dx = 0$	für $n = 1, 2, 3, \dots$
(2)	$\int_0^{2\pi} \cos(nx) dx = 0$	für $n = 1, 2, 3, \dots$
(3)	$\int_0^{2\pi} \cos(nx) \cos(mx) dx = \begin{cases} 0 & \text{für } m \neq n \\ \pi & \text{für } m = n \end{cases}$	$= \pi \delta(n - m)$
(4)	$\int_0^{2\pi} \sin(nx) \sin(mx) dx = \begin{cases} 0 & \text{für } m \neq n \\ \pi & \text{für } m = n \end{cases}$	$= \pi \delta(n - m)$
(5)	$\int_0^{2\pi} \sin(nx) \cos(mx) dx = 0$	für $n, m = 1, 2, 3, \dots$

In Tabelle 14.1 wird das *Kronecker-Symbol* $\delta(k)$ verwendet, das für alle ganzen Zahlen $k \in \mathbb{Z}$ definiert ist durch

$$\delta(k) := \begin{cases} 1 & \text{für } k = 0 \\ 0 & \text{für } k \in \mathbb{Z} \setminus \{0\}. \end{cases}$$

Bemerkung: Integral (1) und (2) rechnet man direkt nach. Mit der Formel $\cos \alpha \cos \beta = \frac{1}{2} (\cos(\alpha - \beta) + \cos(\alpha + \beta))$ gilt im Falle (3) für $m \neq n$

$$\begin{aligned} \int_0^{2\pi} \cos(nx) \cos(mx) dx &= \frac{1}{2} \left(\int_0^{2\pi} \cos((n-m)x) dx + \right. \\ &\quad \left. \int_0^{2\pi} \cos((n+m)x) dx \right) = 0; \end{aligned}$$

für $n = m$ ist

$$\int_0^{2\pi} \cos^2(nx) dx = \pi.$$

Formel (4) berechnet man analog zu (3) mit der Beziehung

$$\sin \alpha \sin \beta = \frac{1}{2} (\cos(\alpha - \beta) - \cos(\alpha + \beta)).$$

Zur Bestimmung von (5) verwendet man die Formel

$$\sin \alpha \cos \beta = \frac{1}{2} (\sin(\alpha - \beta) + \sin(\alpha + \beta)). \quad \square$$

⑤ Bestimmung von a_0

Wir integrieren Gleichung (*) gliedweise im Periodenintervall $[0, 2\pi]$:

$$\begin{aligned} \int_0^{2\pi} f(x) dx &= \underbrace{\int_0^{2\pi} a_0 dx}_{a_0 \cdot 2\pi} + \sum_{n=1}^{\infty} a_n \underbrace{\int_0^{2\pi} \cos(nx) dx}_{=0} \\ &\quad + \sum_{n=1}^{\infty} b_n \underbrace{\int_0^{2\pi} \sin(nx) dx}_{=0} \end{aligned}$$

Nach Tabelle 14.1 ergeben die Integrale $\int_0^{2\pi} \cos(nx) dx$ und $\int_0^{2\pi} \sin(nx) dx$ den Wert Null. In obiger Darstellung kommt nur das erste Integral $\int_0^{2\pi} a_0 dx$ mit dem von Null verschiedenen Wert $a_0 \cdot 2\pi$ vor. Lösen wir nach a_0 auf, folgt

$$\Rightarrow a_0 = \frac{1}{2\pi} \int_0^{2\pi} f(x) dx.$$

⊙ **Bestimmung von a_n**

Wir multiplizieren $f(x)$ in der Darstellung von Gleichung (*) zunächst mit $\cos(mx)$, ($m > 0$) und integrieren anschließend über das Periodenintervall $[0, 2\pi]$:

$$\begin{aligned} \int_0^{2\pi} f(x) \cos(mx) dx &= a_0 \int_0^{2\pi} \cos(mx) dx \\ &+ \sum_{n=1}^{\infty} a_n \int_0^{2\pi} \cos(nx) \cos(mx) dx \\ &+ \sum_{n=1}^{\infty} b_n \int_0^{2\pi} \sin(nx) \cos(mx) dx. \end{aligned}$$

Nach Tabelle 14.1 (5) verschwinden alle Summanden der zweiten Summe. Von der ersten Summe über a_n ist nur der Summand ungleich Null, bei dem der Laufindex n mit m übereinstimmt. Da auch $\int_0^{2\pi} \cos(mx) dx = 0$ ist, gilt

$$\int_0^{2\pi} f(x) \cos(mx) dx = \sum_{n=1}^{\infty} a_n \pi \delta(n-m) = \pi \cdot a_m.$$

$$\Rightarrow \quad a_m = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos(mx) dx \quad m = 1, 2, 3, \dots$$

⊙ **Bestimmung von b_n**

Analog dem Vorgehen zur Berechnung der Koeffizienten a_n multiplizieren wir (*) zunächst mit $\sin(mx)$ ($m > 0$) und integrieren anschließend über das Periodenintervall $[0, 2\pi]$:

$$\begin{aligned} \int_0^{2\pi} f(x) \sin(mx) dx &= a_0 \int_0^{2\pi} \sin(mx) dx \\ &+ \sum_{n=1}^{\infty} a_n \int_0^{2\pi} \cos(nx) \sin(mx) dx \\ &+ \sum_{n=1}^{\infty} b_n \int_0^{2\pi} \sin(nx) \sin(mx) dx. \end{aligned}$$

Alle Integrale, welche als Integranden $\cos(nx)$ enthalten, verschwinden nach Tabelle 14.1, ebenso $\int_0^{2\pi} \sin(mx) dx$. Die Integrale $\int_0^{2\pi} \sin(nx) \sin(mx) dx = \pi \delta(n-m)$ sind alle Null bis auf dasjenige mit $n = m$. Somit ist

$$\int_0^{2\pi} f(x) \sin(mx) dx = \sum_{n=1}^{\infty} b_n \pi \delta(n-m) = b_m \cdot \pi.$$

$$\Rightarrow \quad b_m = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin(mx) dx \quad m = 1, 2, 3, \dots$$

14.3 Fourier-Reihen für 2π -periodische Funktionen

Nach diesen Vorüberlegungen sind für eine 2π -periodische Funktion f die Fourier-Koeffizienten formal bestimmt. Die dadurch definierte Fourier-Reihe konvergiert für die meisten Funktionen und stimmt mit $f(x)$ überein. Hier muss jedoch gewarnt werden: Es gibt stetige, 2π -periodische Funktionen, deren Fourier-Reihe an unendlich vielen Stellen eines Periodenintervalls divergiert. Um sicherzustellen, dass die Fourier-Reihe einer 2π -periodischen Funktion f überall konvergiert und dass der Grenzwert mit $f(x)$ übereinstimmt, muss die Funktion f gewisse Forderungen erfüllen.

Im Folgenden geben wir ohne Beweis zwei Bedingungen an, die zusammen sowohl die Konvergenz als auch die Übereinstimmung der Fourier-Reihe mit der Funktion gewährleisten. Beide Bedingungen sind leicht überprüfbar und sind bei praktisch vorkommenden Funktionen so gut wie immer erfüllt.

Bedingung 1: Das Periodenintervall $[0, 2\pi]$ lässt sich durch endlich viele Teilpunkte $0 = x_1 < x_2 < \dots < x_N = 2\pi$ so zerlegen, dass in den offenen Teilintervallen (x_k, x_{k+1}) $1 \leq k \leq N-1$ die Funktion f differenzierbar und f' beschränkt ist. Man nennt solche Funktionen **stückweise stetig differenzierbare Funktionen**.

Bedingung 2: In den Teilpunkten x_k existiert der linksseitige und der rechtsseitige Grenzwert

$$f_l(x_k) = \lim_{\varepsilon \rightarrow 0} f(x_k - \varepsilon) \quad \text{bzw.} \quad f_r(x_k) = \lim_{\varepsilon \rightarrow 0} f(x_k + \varepsilon)$$

und für den Funktionswert gilt

$$f(x_k) = \frac{1}{2} (f_l(x_k) + f_r(x_k)) .$$

Man nennt diese Eigenschaft die **Mittelwerteigenschaft**.

Das Schaubild einer Funktion, die Bedingung (1) **und** (2) erfüllt, ist in Abb. 14.3 gezeigt.

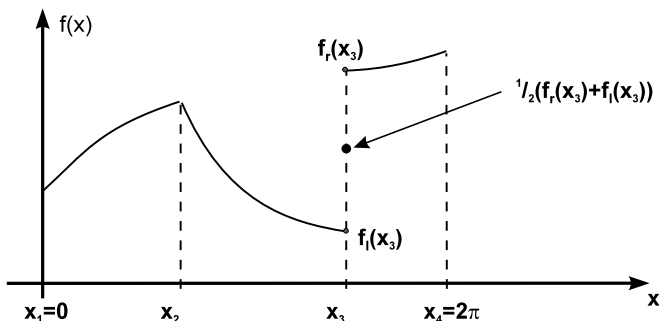


Abb. 14.3. Stückweise stetig differenzierbare Funktion

Stückweise stetig differenzierbare Funktionen dürfen also endlich viele Sprungstellen aufweisen. In Stetigkeitspunkten ist die Mittelwerteigenschaft immer erfüllt, in Unstetigkeitspunkten ist der Funktionswert der Mittelwert von links- und rechtsseitigem Grenzwert.

Satz von Fourier: Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ eine 2π -periodische Funktion, die stückweise stetig differenzierbar ist und für alle $x \in \mathbb{R}$ die Mittelwerteigenschaft erfüllt. Dann konvergiert die **Fourier-Reihe** und es gilt für alle $x \in \mathbb{R}$

$$f(x) = a_0 + \sum_{n=1}^{\infty} a_n \cos(nx) + \sum_{n=1}^{\infty} b_n \sin(nx)$$

mit den **Fourier-Koeffizienten**

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} f(x) dx$$

$$a_n = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos(nx) dx \quad n = 1, 2, 3, \dots$$

$$b_n = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin(nx) dx \quad n = 1, 2, 3, \dots$$

Bemerkung: Für eine 2π -periodische Funktion $f(x)$ gilt stets

$$\int_0^{2\pi} f(x) dx = \int_{\alpha}^{\alpha+2\pi} f(x) dx \quad \text{für beliebiges } \alpha \in \mathbb{R}.$$

Diese Formel besagt, dass zur Berechnung der Fourier-Koeffizienten ein beliebiges Periodenintervall der Länge 2π gewählt werden darf. Berücksichtigt man diese Eigenschaft, dann vereinfacht sich die Berechnung der Fourier-Koeffizienten für symmetrische Funktionen:

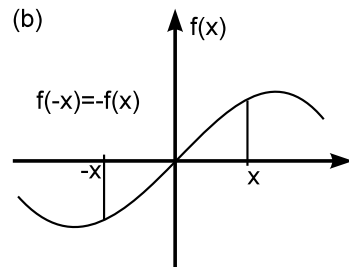
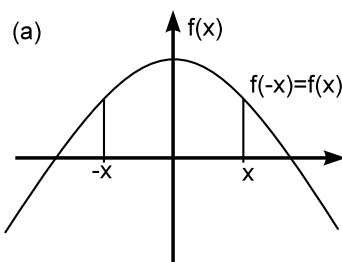


Abb. 14.4. Gerade (a) und ungerade (b) Funktionen

Symmetriebetrachtungen:

- (1) Für eine **gerade** 2π -periodische Funktion f ($f(-x) = f(x)$ für alle x) sind alle Fourier-Koeffizienten b_k ($k \in \mathbb{N}$) gleich Null.
- (2) Für eine **ungerade** 2π -periodische Funktion f ($f(-x) = -f(x)$ für alle x) sind alle Fourier-Koeffizienten a_k ($k \in \mathbb{N}_0$) gleich Null.

Begründung:

- (1) Ist $f(x)$ gerade, so ist auch $f(x) \cdot \cos(nx)$ eine gerade Funktion, $f(x) \cdot \sin(nx)$ dagegen ist ungerade. Wählt man als Integrationsintervall das Intervall $[-\pi, \pi]$, so erhält man für die ungerade Funktion $f(x) \cdot \sin(nx)$ (vgl. Abb. 14.4 (b))

$$b_n = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx = 0 \quad n \in \mathbb{N}$$

und für die Koeffizienten a_n :

$$a_0 = \frac{1}{\pi} \int_0^{\pi} f(x) dx, \quad a_n = \frac{2}{\pi} \int_0^{\pi} f(x) \cos(nx) dx, \quad n \in \mathbb{N}.$$

- (2) Ist $f(x)$ ungerade, so ist auch $f(x) \cdot \cos(nx)$ eine ungerade Funktion, während $f(x) \cdot \sin(nx)$ als Produkt zweier ungerader Funktionen gerade ist. Verwendet man wieder das Integrationsintervall $[-\pi, \pi]$ und berücksichtigt die Symmetrie, folgt (vgl. Abb. 14.4 (a))

$$a_0 = 0 \quad \text{und} \quad a_n = 0, \quad n \in \mathbb{N},$$

und für die Koeffizienten b_n die Formel

$$b_n = \frac{2}{\pi} \int_0^{\pi} f(x) \sin(nx) dx, \quad n \in \mathbb{N}.$$

□

Beispiel 14.1 (Mit MAPLE-Worksheet). Gegeben ist die in Abb. 14.5 gezeichnete Rechteckfunktion mit Periode 2π .

Diese Funktion wird im Periodenintervall $[0, 2\pi]$ beschrieben durch die Funktionsgleichung

$$f(x) = \begin{cases} 1 & 0 < x < \pi \\ 0 & x = 0, \pi, 2\pi \\ -1 & \pi < x < 2\pi \end{cases}.$$

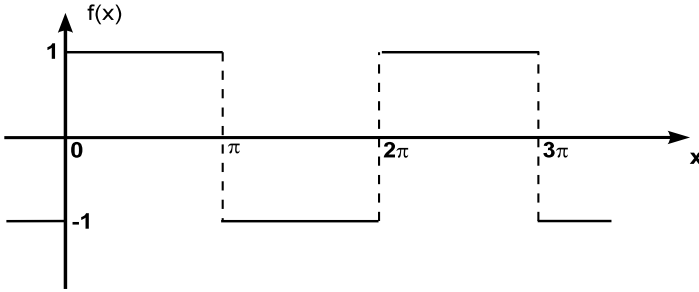


Abb. 14.5. Periodische Rechteckfunktion

f ist stückweise stetig differenzierbar und erfüllt in allen Punkten die Mittelwertsatz. Aufgrund der Punktsymmetrie bezüglich des Ursprungs gilt:

$$a_n = 0 \quad \text{für } n = 0, 1, 2, \dots$$

Es bleiben also nur die Fourier-Koeffizienten b_n zu berechnen. Nach Bemerkung (2) gilt

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_0^{2\pi} f(x) \sin(nx) \, dx \\ &= \frac{2}{\pi} \int_0^{\pi} \sin(nx) \, dx \\ &= \frac{2}{\pi} \left[-\frac{1}{n} \cos(nx) \right]_0^{\pi} \\ &= \frac{2}{\pi n} \{ -\cos(n\pi) + \cos(0) \} = \frac{2}{\pi n} \{ -(-1)^n + 1 \}, \end{aligned}$$

da $\cos(n\pi) = (-1)^n$ und $\cos(0) = 1$.

Für gerade n ist $(-1)^n = +1 \hookrightarrow b_n = 0$; für ungerade n ist $(-1)^n = -1 \hookrightarrow b_n = \frac{2}{\pi n} \cdot 2$.

$$\Rightarrow b_n = \begin{cases} 0 & \text{für } n = 0, 2, 4, \dots \\ \frac{4}{\pi n} & \text{für } n = 1, 3, 5, \dots \end{cases}$$

Die Fourier-Reihe der Funktion f lautet

$$\begin{aligned} f(x) &= \frac{4}{\pi} \left(\sin(x) + \frac{1}{3} \sin(3x) + \frac{1}{5} \sin(5x) + \frac{1}{7} \sin(7x) + \dots \right) \\ &= \sum_{\substack{n=1 \\ n \text{ ungerade}}}^{\infty} \frac{4}{n\pi} \sin(nx) = \sum_{n=0}^{\infty} \frac{4}{(2n+1)\pi} \sin((2n+1)x). \end{aligned}$$

In Abb. 14.6 (a) sind die Partialsummen dieser Reihe für $n = 3, 5, 7$ und in Abb. 14.6 (b) für $n = 40$ dargestellt.

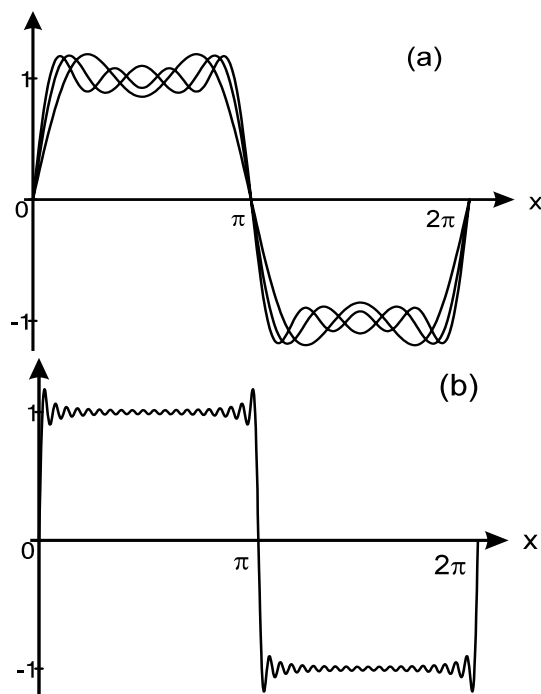


Abb. 14.6. Partialsummen der Fourier-Reihe (a) für $n = 3, 5, 7$ und (b) für $n = 40$

Diskussion: Man erkennt, dass viele Summenglieder notwendig sind, damit die Funktion f durch eine Partialsumme der Fourier-Reihe einigermaßen gut angenähert werden kann. Allerdings bauen sich selbst mit großem N immer noch Oszillationen vor der Sprungstelle auf. **Die Koeffizienten der Fourier-Reihe b_n verhalten sich proportional zu $\frac{1}{n}$.** $\square$



Animation: Auf der CD-Rom befindet sich eine [Animation](#), in der die punktweise Konvergenz der Fourier-Reihe an die Funktion visualisiert wird. Hierbei werden die aus dem Beispiel berechneten Fourier-Koeffizienten für die Darstellung der Fourier-Reihe verwendet.

Beispiel 14.2 (Mit MAPLE-Worksheet). Gegeben ist die in Abb. 14.7 gezeichnete 2π -periodische Funktion, die im Periodenintervall $[0, 2\pi]$ beschrieben wird durch die Funktionsgleichung

$$f(x) = \begin{cases} x & \text{für } 0 \leq x < \pi \\ 0 & \text{für } x = \pi, 2\pi \\ x - 2\pi & \text{für } \pi < x < 2\pi \end{cases}.$$

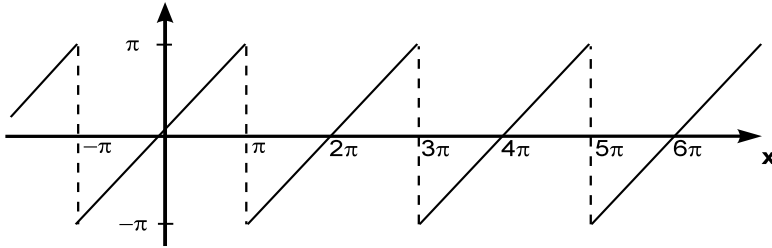


Abb. 14.7. Periodische Dreiecksfunktion

Gesucht ist deren Fourier-Reihe. f ist stückweise stetig differenzierbar und erfüllt in allen Punkten die Mittelwerteigenschaft. Aufgrund der Punktsymmetrie zum Ursprung ist

$$a_n = 0 \quad \text{für } n = 0, 1, 2, \dots$$

Nach Bemerkung (2) berechnet man die Fourier-Koeffizienten b_n durch die Formel

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_0^{2\pi} f(x) \sin(nx) \, dx \\ &= \frac{2}{\pi} \int_0^{\pi} f(x) \sin(nx) \, dx = \frac{2}{\pi} \int_0^{\pi} x \cdot \sin(nx) \, dx. \end{aligned}$$

Partielle Integration liefert

$$\begin{aligned} b_n &= \frac{2}{\pi} \left\{ \left[x \frac{-\cos(nx)}{n} \right]_0^{\pi} - \int_0^{\pi} \frac{-\cos(nx)}{n} \, dx \right\} \\ &= \frac{2}{\pi n} [-\pi \cos(n\pi) - 0] = -\frac{2}{n} (-1)^n = \frac{2}{n} (-1)^{n+1}, \end{aligned}$$

da $\cos(n\pi) = (-1)^n$. Folglich ist die Fourier-Reihe von f

$$\begin{aligned} f(x) &= 2 \left(\sin(x) - \frac{1}{2} \sin(2x) + \frac{1}{3} \sin(3x) - \frac{1}{4} \sin(4x) \pm \dots \right) \\ &= 2 \sum_{n=1}^{\infty} (-1)^{n+1} \frac{1}{n} \sin(nx). \end{aligned}$$

Wie in Beispiel 14.1 verhalten sich auch in diesem Beispiel die Fourier-Koeffizienten betragsmäßig $\sim \frac{1}{n}$. □

Beispiel 14.3 (Mit [MAPLE-Worksheet](#)). Gegeben ist die Funktion

$$f(x) = \frac{1}{\pi} (x - \pi)^2$$

im Intervall $0 \leq x \leq 2\pi$, die 2π -periodisch auf $\mathbb{R}$ fortgesetzt wird:

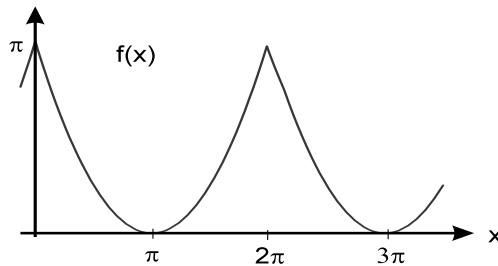


Abb. 14.8.

f ist stückweise stetig differenzierbar und erfüllt die Mittelwerteigenschaft. Aufgrund der Achsensymmetrie bezüglich der y -Achse gilt

$$b_n = 0 \quad \text{für } n = 1, 2, 3, \dots$$

Für den Koeffizient a_0 erhalten wir

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} \frac{1}{\pi} (x - \pi)^2 dx = \frac{1}{2\pi^2} \left[\frac{1}{3} (x - \pi)^3 \right]_0^{2\pi} = \frac{1}{3} \pi.$$

Die Koeffizienten a_n ($n \in \mathbb{N}$) berechnen sich aufgrund der Achsensymmetrie durch die Formel

$$a_n = \frac{2}{\pi} \int_0^\pi \frac{1}{\pi} (x - \pi)^2 \cos(nx) dx.$$

Zweimalige partielle Integration liefert das Ergebnis

$$\begin{aligned} a_n &= \frac{2}{\pi^2} \left\{ \left[(x - \pi)^2 \frac{1}{n} \sin(nx) \right]_0^\pi - 2 \int_0^\pi (x - \pi) \frac{1}{n} \sin(nx) dx \right\} \\ &= -\frac{4}{\pi^2 n} \int_0^\pi (x - \pi) \sin(nx) dx \\ &= -\frac{4}{\pi^2 n} \left\{ \left[(x - \pi) \frac{-1}{n} \cos(nx) \right]_0^\pi - \int_0^\pi \frac{-1}{n} \cos(nx) dx \right\} \\ &= \frac{4}{\pi n^2}. \end{aligned}$$

Somit ist

$$f(x) = \frac{\pi}{3} + \frac{4}{\pi} \sum_{n=1}^{\infty} \frac{1}{n^2} \cos(nx).$$

Die graphische Darstellung der Funktion f zusammen mit den ersten Partialsummen ist für $n = 5$ in Abb. 14.9 angegeben.

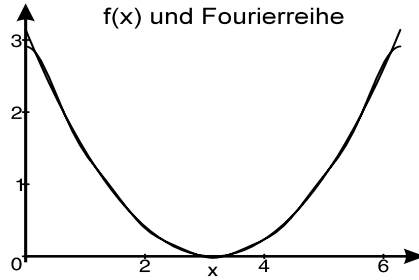


Abb. 14.9. Die Partialsumme der Funktion $f(x) = \frac{1}{\pi}(x - \pi)^2$ für $n = 5$

Diskussion: Im Gegensatz zu Beispiel 14.1 und 14.2 genügen vier Summenglieder, um die Funktion erkennbar durch die Fourier-Reihe anzunähern. Es bauen sich keine Oszillationen im Periodenintervall auf. **Die Koeffizienten der Fourier-Reihe verhalten sich proportional zu $\frac{1}{n^2}$.**

Nebenergebnis. Mit den Fourier-Reihen ist man in der Lage für einige in Abschnitt 9.1 diskutierte Reihen den Summenwert zu berechnen, indem in die Fourier-Reihe spezielle Werte eingesetzt werden. Man erhält so die folgenden beiden Nebenergebnisse:

$$x = 0: \quad f(0) = \pi = \frac{\pi}{3} + \frac{4}{\pi} \sum_{n=1}^{\infty} \frac{1}{n^2} \quad \Rightarrow \quad \sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{6}.$$

$$x = \pi: \quad f(\pi) = 0 = \frac{\pi}{3} + \frac{4}{\pi} \sum_{n=1}^{\infty} (-1)^n \frac{1}{n^2} \quad \Rightarrow \quad \sum_{n=1}^{\infty} (-1)^n \frac{1}{n^2} = \frac{\pi^2}{12}.$$

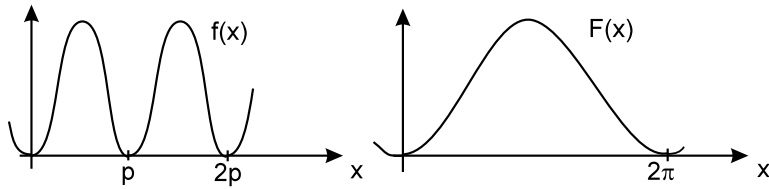
□

14.4 Fourier-Reihen für p -periodische Funktionen

14.4

Bisher wurde die Fourier-Reihendarstellung von 2π -periodischen Funktionen betrachtet. Um die entsprechende Darstellung mit den zugehörigen Fourier-Koeffizienten einer p -periodischen Funktion f zu bestimmen, gehen wir von f über zu der auf das Intervall $[0, 2\pi]$ gestauchten bzw. gestreckten Funktion

$$F(x) := f\left(\frac{p}{2\pi}x\right).$$

Abb. 14.10. p -periodische Funktion f und zugehörige 2π -periodische Funktion F

Ist f p -periodisch, dann ist F 2π -periodisch. Denn

$$F(x + 2\pi) = f\left(\frac{p}{2\pi}(x + 2\pi)\right) = f\left(\frac{p}{2\pi}x + p\right) = f\left(\frac{p}{2\pi}x\right) = F(x).$$

Für die 2π -periodische Funktion $F(x)$ stellt man formal die Fourier-Reihe auf

$$F(x) = a_0 + \sum_{n=1}^{\infty} a_n \cos(nx) + \sum_{n=1}^{\infty} b_n \sin(nx).$$

Die Fourier-Koeffizienten sind gegeben durch

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} F(x) dx; \quad a_n = \frac{1}{\pi} \int_0^{2\pi} F(x) \cos(nx) dx;$$

$$b_n = \frac{1}{\pi} \int_0^{2\pi} F(x) \sin(nx) dx.$$

Wir führen nun durch die Rücksubstitution

$$f(x) = F\left(\frac{2\pi}{p}x\right)$$

die Fourier-Reihe von f auf die Fourier-Reihe von F zurück. Die Formel für die Fourier-Reihe von f ergibt sich aus:

$$f(x) = F\left(\frac{2\pi}{p}x\right) = a_0 + \sum_{n=1}^{\infty} a_n \cos\left(n \frac{2\pi}{p}x\right) + \sum_{n=1}^{\infty} b_n \sin\left(n \frac{2\pi}{p}x\right).$$

Mit der Integralsubstitution $y = \frac{p}{2\pi}x$ erhält man die Fourier-Koeffizienten von f :

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} F(x) dx = \frac{1}{2\pi} \int_0^{2\pi} f\left(\frac{p}{2\pi}x\right) dx = \frac{1}{p} \int_0^p f(y) dy,$$

$$\begin{aligned} a_n &= \frac{1}{\pi} \int_0^{2\pi} F(x) \cos(nx) dx = \frac{1}{\pi} \int_0^{2\pi} f\left(\frac{p}{2\pi}x\right) \cos(nx) dx \\ &= \frac{2}{p} \int_0^p f(y) \cos\left(n \frac{2\pi}{p}y\right) dy, \end{aligned}$$

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_0^{2\pi} F(x) \sin(nx) dx = \frac{1}{\pi} \int_0^{2\pi} f\left(\frac{p}{2\pi}x\right) \sin(nx) dx \\ &= \frac{2}{p} \int_0^p f(y) \sin\left(n \frac{2\pi}{p}y\right) dy. \end{aligned}$$

Zusammenfassend gilt:

Satz von Fourier für p -periodische Funktionen: Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ eine p -periodische Funktion, die stückweise stetig differenzierbar ist und die für alle $x \in \mathbb{R}$ die Mittelwerteigenschaft erfüllt. Dann konvergiert die Fourier-Reihe

$$f(x) = a_0 + \sum_{n=1}^{\infty} a_n \cos\left(n \frac{2\pi}{p} x\right) + \sum_{n=1}^{\infty} b_n \sin\left(n \frac{2\pi}{p} x\right)$$

für alle $x \in \mathbb{R}$ und stimmt mit der Funktion f überein. Die Koeffizienten sind

$$a_0 = \frac{1}{p} \int_0^p f(x) dx$$

$$a_n = \frac{2}{p} \int_0^p f(x) \cos\left(n \frac{2\pi}{p} x\right) dx \quad n = 1, 2, 3, \dots$$

$$b_n = \frac{2}{p} \int_0^p f(x) \sin\left(n \frac{2\pi}{p} x\right) dx \quad n = 1, 2, 3, \dots$$

Bemerkung: Die Formeln für 2π -periodische Funktionen sind der Spezialfall $p = 2\pi$. Bemerkungen (1) und (2) übertragen sich sinngemäß auf p -periodische Funktionen.

Anwendung: Fourier-Zerlegung T -periodischer Signale

Sei $f(t)$ eine periodische Schwingung mit Periode T (= Schwingungsdauer). Dann gilt zu jedem Zeitpunkt die Fourier-Zerlegung

$$f(t) = a_0 + \sum_{n=1}^{\infty} a_n \cos(n \omega_0 t) + \sum_{n=1}^{\infty} b_n \sin(n \omega_0 t) \quad (*)$$

mit der Grundfrequenz $\omega_0 = \frac{2\pi}{T}$ und den angegebenen Fourier-Koeffizienten.

Die Entwicklung des Zeitsignals $f(t)$ in unendlich viele Sinus- und Kosinusfunktionen bedeutet aus physikalischer Sicht eine Zerlegung des Signals in seine harmonischen Bestandteile.

Ein periodisches Signal besitzt ein diskretes Spektrum: Es besteht aus der Grundschiwingung mit Frequenz ω_0 und den harmonischen Oberschwingungen mit Frequenzen $n\omega_0$. Die Fourier-Koeffizienten bestimmen dabei die Amplituden dieser harmonischen Teilschwingungen und damit den Beitrag der Oberschwingungen zum Signal.

Allerdings erhält man zu einer Frequenz $n\omega_0$ zunächst zwei Koeffizienten, nämlich a_n und b_n , da die Summanden in der Fourier-Reihe (*) die Überlagerung von jeweils zwei harmonischen Sinus- und Kosinusschwingungen gleicher Frequenz darstellen. Ist nach der Amplitude gefragt, mit welcher die Frequenz $n\omega_0$ im Signal vorkommt, muss man übergehen zu der Darstellung

$$a_n \cos(n\omega_0 t) + b_n \sin(n\omega_0 t) = A_n \cos(n\omega_0 t - \varphi_n)$$

mit

$$A_n = \sqrt{a_n^2 + b_n^2}$$

und

$$\tan \varphi_n = \frac{b_n}{a_n}.$$

A_n ist dann die gesuchte Amplitude und φ_n die Nullphase.

Begründung: Denn nach dem Additionstheorem für den Kosinus gilt

$$\begin{aligned} A_n \cos(n\omega_0 t - \varphi_n) &= A_n \cos(n\omega_0 t) \cos \varphi_n + A_n \sin(n\omega_0 t) \sin \varphi_n \\ &= a_n \cos(n\omega_0 t) + b_n \sin(n\omega_0 t) . \end{aligned}$$

Vergleicht man die Koeffizienten von $\cos(n\omega_0 t)$ und $\sin(n\omega_0 t)$, folgt

$$a_n = A_n \cos \varphi_n \quad \text{und} \quad b_n = A_n \sin \varphi_n .$$

Quadriert man beide Gleichungen und addiert sie dann, gilt

$$a_n^2 + b_n^2 = A_n^2 \cos^2 \varphi_n + A_n^2 \sin^2 \varphi_n = A_n^2 \quad \Rightarrow \quad A_n = \sqrt{a_n^2 + b_n^2} .$$

Dividiert man die zweite durch die erste Gleichung gilt

$$\frac{b_n}{a_n} = \tan \varphi_n .$$

□

In den Anwendungen benutzt man daher die Darstellung der Fourier-Reihe in der Form

$$f(t) = a_0 + \sum_{n=1}^{\infty} A_n \cos(n\omega_0 t - \varphi_n)$$

bzw. geht zur komplexen Formulierung (siehe 14.5) über.

A_n ist die **Gesamtamplitude**, mit der die Frequenz $\omega_n = n\omega_0$ im Signal vorkommt. φ_n ist die zugehörige **Phase**. Der Wert a_0 gibt den Koeffizienten $a_0 = \frac{1}{T} \int_0^T f(t) dt$ an. Er entspricht dem **Mittelwert** der Funktion während einer Schwingungsdauer. Man bezeichnet ihn als *Gleichspannungsanteil*. Zur graphischen Darstellung der Koeffizienten wählt man die folgenden Diagramme:

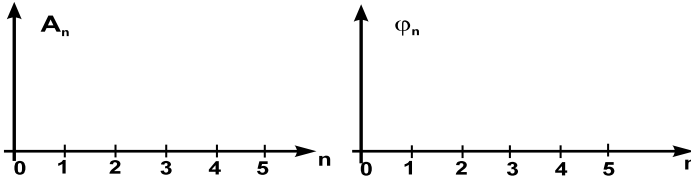


Abb. 14.11. Amplitudenspektrum A_n und Phasenspektrum φ_n

In diesen Schaubildern werden die Amplituden A_n und Phasen φ_n als Werte über den diskreten Frequenzen abgetragen (= *diskretes Spektrum* von f). Man bezeichnet die Darstellung der Amplituden als **Amplitudenspektrum** (links) und die der Phasen als **Phasenspektrum** (rechts).

Beispiel 14.4 (Amplitudenspektrum, mit MAPLE-Worksheet). Die Amplitudenspektren $A_n = \sqrt{a_n^2 + b_n^2}$ sind für die Beispiele 14.1, 14.2 und 14.3 in Abb. 14.12 gezeichnet.

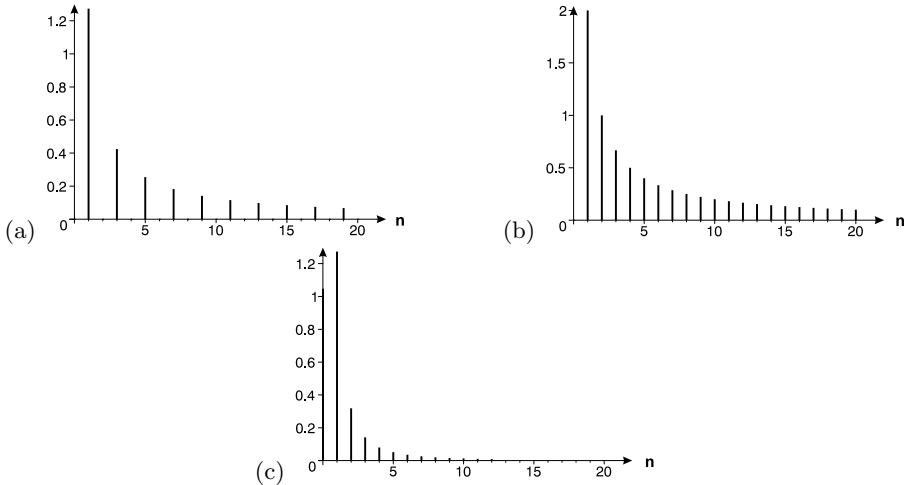


Abb. 14.12. Amplitudenspektrum zu (a) Beispiel 14.1, (b) 14.2 und (c) 14.3

Diese Darstellung beinhaltet als Information über das Signal sowohl die vorkommenden Frequenzen $n\omega_0$ als auch die zugehörigen Amplituden A_n . Man erkennt in den Amplitudenspektren, dass in den Beispielen 14.1 und 14.2 die Koeffizienten langsam abnehmen, während in Beispiel 14.3 die Koeffizienten sehr schnell zu Null gehen. Dies spiegelt die Tatsache wider, dass die Konvergenz in den beiden ersten Fällen $\sim \frac{1}{n}$, im dritten Fall jedoch $\sim \frac{1}{n^2}$. $\square$

Konvergenzbetrachtungen. Durch den Abbruch der Fourier-Reihe nach endlich vielen Gliedern, erhält man eine Näherungsfunktion für f in Form einer endlichen trigonometrischen Reihe. Ähnlich wie bei Potenzreihen gilt: Je mehr Glieder berücksichtigt werden, um so besser ist die Approximation. Man stellt zunächst für alle behandelten Beispiele fest:

$$a_n \rightarrow 0 \quad \text{und} \quad b_n \rightarrow 0 \quad \text{für} \quad n \rightarrow \infty .$$

Bei genauerer Betrachtung entdeckt man, dass die Näherungen für stetige Funktionen schneller konvergieren bzw. man wenige Glieder in der Fourier-Reihe benötigt, um die Funktion hinreichend gut zu approximieren. Unsere Beispiele zeigen, dass

$$a_n \sim \frac{1}{n^2} \quad \text{und} \quad b_n \sim \frac{1}{n^2} ,$$

wenn f keine Sprungstellen hat. Anschaulich kann man sagen: Die Fourier-Reihe einer p -periodischen Funktion konvergiert um so schneller, je "glatter" die Funktion f ist. Präziser gilt der Satz

Satz: Ist f eine p -periodische, $m+1$ -mal stetig differenzierbare Funktion, dann gilt für die Fourier-Koeffizienten von f

$$|a_n| \leq \frac{c}{n^{m+2}} \quad \text{und} \quad |b_n| \leq \frac{c}{n^{m+2}} .$$

Beispiele für $m = 0$ (stetige Funktionen) sind 14.3, 14.5 und Beispiele für $m = -1$ (Funktionen mit Sprungstelle) sind 14.1, 14.2.

Anwendungsbeispiel 14.5 (Fourier-Zerlegung eines Sinusimpuls-Einweggleichrichters, mit [MAPLE-Worksheet](#)).

Wir betrachten das in Abb. 14.13 gezeigte Signal $f(t)$ des Sinusimpulses eines Einweggleichrichters mit Periode T :

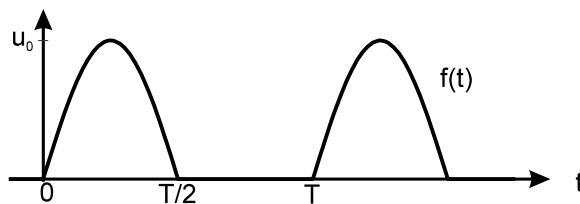


Abb. 14.13. Sinusimpuls eines Einweggleichrichters

Der Impuls wird im Periodenintervall $[0, T]$ durch die Funktionsgleichung

$$f(t) = \begin{cases} u_0 \sin(\omega_0 t) & 0 \leq t \leq \frac{T}{2} \\ 0 & \frac{T}{2} < t \leq T \end{cases}$$

mit $\omega_0 = \frac{2\pi}{T}$ beschrieben. Da das Signal im Bereich $\frac{T}{2} \leq t \leq T$ gleich Null ist, reduzieren sich die Integrationsgrenzen der Integrale auf $t = 0 \dots \frac{T}{2}$.

Berechnung des Koeffizienten a_0 :

$$\begin{aligned} a_0 &= \frac{1}{T} \int_0^T f(t) dt = \frac{1}{T} \int_0^{T/2} u_0 \sin(\omega_0 t) dt \\ &= -\frac{u_0}{T} \left[\frac{1}{\omega_0} \cos(\omega_0 t) \right]_0^{T/2} = \frac{u_0}{\pi}. \end{aligned}$$

Berechnung der Koeffizienten a_n :

$$a_n = \frac{2}{T} \int_0^T f(t) \cos(n\omega_0 t) dt = \frac{2u_0}{T} \int_0^{T/2} \sin(\omega_0 t) \cos(n\omega_0 t) dt.$$

Mit der trigonometrischen Formel

$$\sin(\alpha) \cos(\beta) = \frac{1}{2} (\sin(\alpha - \beta) + \sin(\alpha + \beta))$$

folgt für $n \neq 1$

$$\begin{aligned} a_n &= \frac{u_0}{T} \left\{ \int_0^{T/2} \sin((1-n)\omega_0 t) dt + \int_0^{T/2} \sin((1+n)\omega_0 t) dt \right\} \\ &= \frac{u_0}{T} \left\{ \left[-\frac{1}{(1-n)\omega_0} \cos((1-n)\omega_0 t) \right]_0^{T/2} \right. \\ &\quad \left. - \left[-\frac{1}{(1+n)\omega_0} \cos((1+n)\omega_0 t) \right]_0^{T/2} \right\} \\ &= \frac{u_0}{\pi} \frac{1}{n^2 - 1} ((-1)^n + 1). \end{aligned}$$

Anhand des Ergebnisses für a_n erkennt man, dass der Integralausdruck formal zwar berechnet wird, aber nur für $n \neq 1$ definiert ist. Der Koeffizient a_1 muss separat bestimmt werden:

$$\begin{aligned} a_1 &= \frac{u_0}{T} \left\{ \int_0^{T/2} 0 dt + \int_0^{T/2} \sin(2\omega_0 t) dt \right\} \\ &= \frac{u_0}{T} \left\{ \left[-\frac{1}{2\omega_0} \cos(2\omega_0 t) \right]_0^{T/2} \right\} = 0. \end{aligned}$$

Insgesamt erhält man also für die Koeffizienten a_n :

$$a_n = \begin{cases} 0 & \text{für } n \text{ ungerade} \\ -\frac{2u_0}{\pi(n^2-1)} & \text{für } n \text{ gerade, } n > 0. \end{cases}$$

Analog berechnen sich die Fourier-Koeffizienten

$$b_n = \frac{2}{T} \int_0^T f(t) \sin(n\omega_0 t) dt = \frac{2u_0}{T} \int_0^{T/2} \sin(\omega_0 t) \sin(n\omega_0 t) dt$$

mit der trigonometrischen Formel

$$\sin(\alpha) \sin(\beta) = \frac{1}{2} (\cos(\alpha - \beta) - \cos(\alpha + \beta)).$$

Es folgt für $n \neq 1$

$$b_n := -\frac{\sin(n\pi) u_0}{\pi(1+n)(-1+n)} = 0.$$

Auch hier muss der Koeffizient b_1 separat berechnet werden:

$$b_1 = \frac{1}{2} u_0.$$

Die Fourier-Reihe des Sinusimpulses hat demnach folgende Gestalt

$$\begin{aligned} f(t) &= \frac{u_0}{\pi} + \frac{u_0}{2} \sin(\omega_0 t) - \frac{2}{\pi} u_0 \left(\frac{1}{2^2-1} \cos(2\omega_0 t) + \frac{1}{4^2-1} \cos(4\omega_0 t) + \right. \\ &\quad \left. + \frac{1}{6^2-1} \cos(6\omega_0 t) + \dots \right) \\ &= \frac{u_0}{\pi} + \frac{u_0}{2} \sin(\omega_0 t) - \frac{2}{\pi} u_0 \sum_{\substack{n=2 \\ n \text{ gerade}}}^{\infty} \frac{1}{n^2-1} \cos(n\omega_0 t). \end{aligned}$$

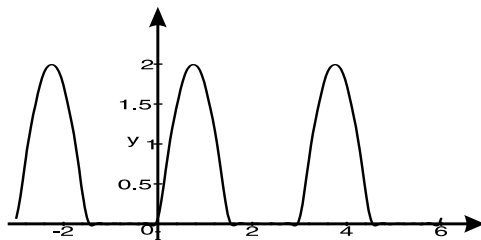


Abb. 14.14. Fourier-Reihe des Sinusimpulses für $n = 8$

Die graphische Darstellung der Fourier-Reihe ist für $n = 8$ in Abb. 14.14 gezeigt. Die Fourier-Reihe zeigt mit wenigen Summengliedern eine sehr gute Übereinstimmung mit der Funktion. $\square$

Hinweis: Auf der CD-Rom befindet sich die Fourier-Analyse einer [Kippspannung](#) sowie eine tabellarische [Zusammenfassung](#) der Beispiele.

14.5 Fourier-Reihen für komplexwertige Funktionen

Eine besonders einfache Gestalt nimmt die Fourier-Darstellung in komplexer Schreibweise an. Unter Benutzung der Eulerschen Formeln (vgl. Abschnitt 9.5.2)

$$\cos(x) = \frac{1}{2} (e^{ix} + e^{-ix}) \quad \text{und} \quad \sin(x) = \frac{1}{2i} (e^{ix} - e^{-ix})$$

lässt sich die reelle Fourier-Reihe einer p -periodischen Funktion f schreiben als:

$$f(x) = a_0 + \sum_{n=1}^{\infty} a_n \frac{1}{2} \left(e^{in \frac{2\pi}{p} x} + e^{-in \frac{2\pi}{p} x} \right) + \sum_{n=1}^{\infty} b_n \frac{1}{2i} \left(e^{in \frac{2\pi}{p} x} - e^{-in \frac{2\pi}{p} x} \right).$$

Nach Umordnung der Summanden in eine Summe über $e^{in \frac{2\pi}{p} x}$ und eine zweite Summe über $e^{-in \frac{2\pi}{p} x}$ ist

$$f(x) = a_0 + \sum_{n=1}^{\infty} \frac{1}{2} (a_n - i b_n) e^{in \frac{2\pi}{p} x} + \sum_{n=1}^{\infty} \frac{1}{2} (a_n + i b_n) e^{-in \frac{2\pi}{p} x}.$$

Betrachtet man die drei Summanden der Fourier-Reihe von f , so enthält die letzte Summe Terme mit Faktoren $e^{in \frac{2\pi}{p} x}$ für $n = -\infty, \dots, 1$, die mittlere Summe Terme $e^{in \frac{2\pi}{p} x}$ für $n = 1, \dots, \infty$ und der erste Summand den Term $e^{in \frac{2\pi}{p} x}$ für $n = 0$. Wir definieren

$$c_0 := a_0;$$

$$\begin{aligned} c_n &:= \frac{1}{2} (a_n - i b_n) \\ &= \frac{1}{2} \left(\frac{2}{p} \int_0^p f(x) \cos\left(n \frac{2\pi}{p} x\right) dx - i \frac{2}{p} \int_0^p f(x) \sin\left(n \frac{2\pi}{p} x\right) dx \right) \\ &= \frac{1}{p} \int_0^p f(x) \left(\cos\left(n \frac{2\pi}{p} x\right) - i \sin\left(n \frac{2\pi}{p} x\right) \right) dx \\ &= \frac{1}{p} \int_0^p f(x) e^{-in \frac{2\pi}{p} x} dx, \quad n \geq 1; \end{aligned}$$

und

$$\begin{aligned} c_{-n} &:= \frac{1}{2} (a_n + i b_n) \\ &= \frac{1}{2} \left(\frac{2}{p} \int_0^p f(x) \cos\left(n \frac{2\pi}{p} x\right) dx + i \frac{2}{p} \int_0^p f(x) \sin\left(n \frac{2\pi}{p} x\right) dx \right) \\ &= \frac{1}{p} \int_0^p f(x) \left(\cos\left(n \frac{2\pi}{p} x\right) + i \sin\left(n \frac{2\pi}{p} x\right) \right) dx \\ &= \frac{1}{p} \int_0^p f(x) e^{in \frac{2\pi}{p} x} dx, \quad n \geq 1. \end{aligned}$$

Dann stellt sich die Fourier-Reihe von f dar als **eine** Summe über $e^{i n \frac{2\pi}{p} x}$ für $n = -\infty, \dots, \infty$:

$$f(x) = \sum_{n=-\infty}^{\infty} c_n e^{i n \frac{2\pi}{p} x}$$

mit den einheitlichen Koeffizienten

$$c_n = \frac{1}{p} \int_0^p f(x) e^{-i n \frac{2\pi}{p} x} dx \quad \text{für } n \in \mathbb{Z}. \quad \square$$

Es gilt also folgende *komplexe* Formulierung des Satzes von Fourier

Satz von Fourier: (Komplexe Formulierung). Sei $f : \mathbb{R} \rightarrow \mathbb{C}$ eine komplexwertige Funktion mit reeller Periode p . f sei stückweise stetig differenzierbar und erfülle die Mittelwerteigenschaft. Dann konvergiert die *komplexe Fourier-Reihe* für alle $x \in \mathbb{R}$ gegen $f(x)$:

$$f(x) = \sum_{n=-\infty}^{\infty} c_n e^{i n \frac{2\pi}{p} x}.$$

Die *komplexen Fourier-Koeffizienten* sind gegeben durch

$$c_n = \frac{1}{p} \int_0^p f(x) e^{-i n \frac{2\pi}{p} x} dx \quad \text{für } n \in \mathbb{Z}.$$

⊗ **Bemerkungen: Eigenschaften der komplexen Formulierung**

- (1) Es gibt nur eine Summenformel für die Fourier-Reihe und die Koeffizienten c_n werden über eine einheitliche Formel bestimmt.
- (2) Da die Summation der komplexen Fourier-Reihe von $n = -\infty \dots \infty$ geht, kommen formal negative Frequenzen vor. Dies rührt nur von der komplexen Formulierung her: $e^{i \omega_n t}$ und $e^{-i \omega_n t}$ werden benötigt, um die reelle Schwingung $\sin(\omega_n t)$ bzw. $\cos(\omega_n t)$ mit reeller Frequenz $\omega_n > 0$ zu beschreiben, da $\cos(\omega_n t) = \frac{1}{2}(e^{i \omega_n t} + e^{-i \omega_n t})$ und $\sin(\omega_n t) = \frac{1}{2i}(e^{i \omega_n t} - e^{-i \omega_n t})$.

⊗ **und Vorteile dieser komplexen Formulierung**

Die Fourier-Koeffizienten c_n von reellen Signalen $f(x)$ besitzen die folgenden Eigenschaften:

- (3) $c_n^* = c_{-n}$: Die Koeffizienten zu negativen n sind das Komplex-konjugierte der entsprechenden Koeffizienten zu positiven n .
- (4) Folglich sind die Beträge von c_n und c_{-n} gleich, nämlich

$$|c_n| = \left| \frac{1}{2} (a_n - i b_n) \right| = \frac{1}{2} \sqrt{a_n^2 + b_n^2} = \frac{1}{2} A_n.$$

Der Betrag der komplexen Fourier-Koeffizienten stimmt bis auf den Faktor $\frac{1}{2}$ mit A_n überein. D.h. der **Betrag repräsentiert das Amplitudenspektrum** jeweils zur Hälfte von $-\infty$ bis -1 und von 1 bis ∞ .

- (5) Die Phase der komplexen Fourier-Koeffizienten ist bestimmt durch

$$\tan \varphi_n = \frac{\operatorname{Im} c_n}{\operatorname{Re} c_n} = \frac{-\frac{1}{2} b_n}{\frac{1}{2} a_n} = -\frac{b_n}{a_n}.$$

Bis auf das Vorzeichen ist dies das **Phasenspektrum**.

- (6) $c_0 = a_0$ stellt den **Gleichstromanteil** des Signals dar.

Der große Vorteil der komplexen Formulierung liegt also darin, dass eine einheitliche Formel für die Koeffizienten existiert und diese Koeffizienten das Amplitudenspektrum (bis auf den Faktor $\frac{1}{2}$) **und** das Phasenspektrum (bis auf das Vorzeichen) beinhalten.

Beispiel 14.6 (Mit MAPLE-Worksheet). Gesucht ist die komplexe Fourier-Reihe der unten gezeichneten Funktion f , die T -periodisch auf $\mathbb{R}$ fortgesetzt wird:

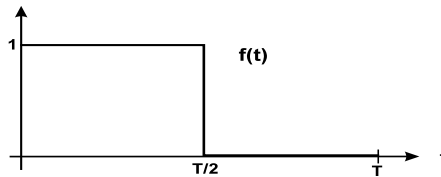


Abb. 14.15. Rechtecksignal

Setzen wir $\omega_0 = \frac{2\pi}{T}$, so sind die komplexen Fourier-Koeffizienten für $n \neq 0$ gegeben durch

$$\begin{aligned} c_n &= \frac{1}{T} \int_0^T f(t) e^{-i n \omega_0 t} dt = \frac{1}{T} \int_0^{\frac{T}{2}} e^{-i n \omega_0 t} dt \\ &= \frac{1}{T} \frac{1}{-i n \omega_0} \left[e^{-i n \omega_0 \frac{T}{2}} - 1 \right]. \end{aligned}$$

Mit $\omega_0 \cdot T = 2\pi$ und $\omega_0 \frac{T}{2} = \pi$ ist $e^{-i n \omega_0 \frac{T}{2}} = e^{-i n \pi} = (-1)^n$

$$\Rightarrow c_n = \frac{1}{-i n 2\pi} [(-1)^n - 1] = \begin{cases} \frac{1}{i n \pi} & n \text{ ungerade} \\ 0 & n \text{ gerade.} \end{cases}$$

Da der Gleichstromanteil des Signals $\frac{1}{2} \Rightarrow c_0 = \frac{1}{2}$.

$$\Rightarrow f(t) = \frac{1}{2} + \sum_{\substack{n=-\infty \\ n \text{ ungerade}}}^{\infty} \frac{1}{i n \pi} e^{i n \omega_0 t}.$$

Das Amplitudenspektrum dieser Funktion ist gegeben durch

$$a_0 = |c_0| = \frac{1}{2}$$

$$A_n = 2|c_n| = \begin{cases} \frac{2}{n\pi} & \text{für } n \text{ ungerade} \\ 0 & \text{für } n \text{ gerade} \end{cases}.$$

□

⊗ Berechnung der reellen Fourier-Koeffizienten aus den komplexen

Selbst wenn man die Rechnung im Komplexen durchgeführt hat, ist man gelegentlich an den reellen Fourier-Koeffizienten a_n und b_n interessiert. Aus den komplexen Fourier-Koeffizienten c_n , $n \in \mathbb{Z}$, lassen sich die reellen Fourier-Koeffizienten a_0 , a_n und b_n einfach zurückgewinnen, so dass sie nicht neu über die reellen Integralformeln berechnet werden müssen. Es gilt nach der Definition der c_n :

c_0	$=$	a_0	(1)
c_n	$=$	$\frac{1}{2} (a_n - i b_n)$	(2)
c_{-n}	$=$	$\frac{1}{2} (a_n + i b_n)$	(3)

Aus (1) folgt direkt

Addiert man (2) und (3) gilt

Subtrahiert man (3) von (2) gilt

$a_0 = c_0$	
$a_n = c_n + c_{-n}$	$n = 1, 2, 3, \dots$
$b_n = i (c_n - c_{-n})$	$n = 1, 2, 3, \dots$

□

MAPLE-Worksheets zu Kapitel 14



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 14 mit MAPLE zur Verfügung.

- [2π-periodischer Funktionen mit MAPLE](#)
- [T-periodische Signale mit MAPLE](#)
- [Numerische Bestimmung der Fourier-Koeffizienten mit MAPLE](#)
- [Komplexe Fourier-Koeffizienten mit MAPLE](#)
- [Darstellung periodischer Funktionen mit MAPLE](#)

Zusammenfassung: (Fourier-Reihen). Gegeben ist ein T -periodisches Signal $f(t)$, das stückweise stetig differenzierbar ist und die Mittelwert-eigenschaft erfüllt.

(1) Für alle $t \in \mathbb{R}$ gilt die reelle Fourier-Reihendarstellung

$$f(t) = a_0 + \sum_{n=1}^{\infty} a_n \cos\left(n \frac{2\pi}{T} t\right) + \sum_{n=1}^{\infty} b_n \sin\left(n \frac{2\pi}{T} t\right)$$

$$\text{mit } a_0 = \frac{1}{T} \int_0^T f(t) dt,$$

$$a_n = \frac{2}{T} \int_0^T f(t) \cos\left(n \frac{2\pi}{T} t\right) dt \quad n = 1, 2, 3, \dots,$$

$$b_n = \frac{2}{T} \int_0^T f(t) \sin\left(n \frac{2\pi}{T} t\right) dt \quad n = 1, 2, 3, \dots$$

(2) Für alle $t \in \mathbb{R}$ gilt

$$f(t) = a_0 + \sum_{n=1}^{\infty} A_n \cos\left(n \frac{2\pi}{T} t - \varphi_n\right)$$

$$\text{mit } A_n = \sqrt{a_n^2 + b_n^2} \text{ und } \tan \varphi_n = \frac{b_n}{a_n} \text{ für } n = 1, 2, 3, \dots$$

(3) Für alle $t \in \mathbb{R}$ gilt die komplexe Fourier-Reihendarstellung

$$f(t) = \sum_{n=-\infty}^{\infty} c_n e^{i n \frac{2\pi}{T} t}$$

$$\text{mit } c_n = \frac{1}{T} \int_0^T f(t) e^{-i n \frac{2\pi}{T} t} dt; \quad n \in \mathbb{Z}.$$

(4) Die Umrechnung der reellen Fourier-Koeffizienten zu den komplexen erfolgt mit den Formeln

$$\begin{aligned} c_0 &= a_0, \\ c_n &= \frac{1}{2} (a_n - i b_n) & n \in \mathbb{N}, \\ c_{-n} &= \frac{1}{2} (a_n + i b_n) & n \in \mathbb{N}. \end{aligned}$$

(5) Die Umrechnung der komplexen Fourier-Koeffizienten zu den reellen erfolgt mit den Formeln

$$\begin{aligned} a_0 &= c_0, \\ a_n &= c_n + c_{-n} & n \in \mathbb{N}, \\ b_n &= i (c_n - c_{-n}) & n \in \mathbb{N}. \end{aligned}$$

14.6 Aufgaben zu Fourier-Reihen

- 14.1 Bestimmen Sie für die unten skizzierten 2π -periodischen Funktionen im Periodenintervall einen formelmäßigen Ausdruck, suchen Sie nach eventuell vorhandenen Symmetrien und entwickeln Sie die Funktionen dann in eine reelle Fourier-Reihe:

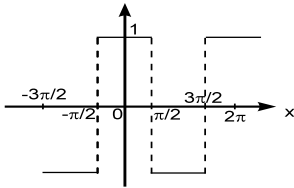


Abb. a

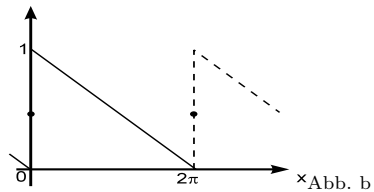


Abb. b

- 14.2 Skizzieren Sie die folgenden T -periodischen Funktionen und bestimmen Sie die Fourier-Reihe sowie das zugehörige Amplitudenspektrum:

$$\text{a) } f(t) = \begin{cases} e^t & -\frac{T}{2} \leq t \leq 0 \\ e^{-t} & 0 \leq t \leq \frac{T}{2} \end{cases} \quad \text{b) } f(t) = \begin{cases} \frac{2ht}{T} & 0 \leq t \leq \frac{T}{2} \\ \frac{h}{2} & \frac{T}{2} \leq t \leq T \end{cases}$$

- 14.3 Geben sei die komplexe Fourier-Entwicklung von

a) Aufgabe 14.1a) b) Aufgabe 14.1b) an.

- 14.4 Entwickeln Sie $f(t) = \sin^3 t$, $|t| \leq \pi$ in eine Fourier-Reihe (erst nachdenken!!!).

- 14.5 a) Entwickeln Sie $f(t) = t^2$, $0 \leq t \leq T$ in eine Fourier-Reihe.

b) Was ergibt sich für $T = 2\pi$?

c) Was erhält man bei b) für $\sum_{n=1}^{\infty} \frac{1}{n^2}$, wenn $t = 2\pi$ gesetzt wird?

- 14.6 In Abb. (c) ist der zeitliche Verlauf einer Kippspannung (Sägezahnimpuls) mit der Schwingungsdauer T angegeben. Man führe für diesen Impuls eine Fourier-Analyse durch. $u(t) = \frac{u_0}{T} t$ ($0 < t < T$).

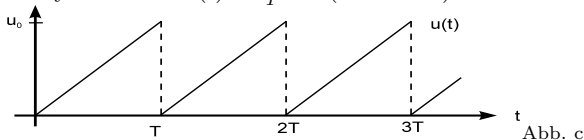


Abb. c

- 14.7 Zeichnen und entwickeln Sie die folgenden Funktionen in Fourier-Reihen unter Berücksichtigung möglicher Symmetrieeigenschaften:

$$\text{a) } f(x) = \begin{cases} 8 & 0 < x < 2 \\ -8 & 2 < x < 4 \end{cases} \text{ mit Periode } 4,$$

$$\text{b) } f(x) = \begin{cases} -x & -4 \leq x \leq 0 \\ x & 0 \leq x \leq 4 \end{cases} \text{ mit Periode } 8.$$

- 14.8 Zerlegen Sie (mit MAPLE) den trapezförmigen Impuls in seine harmonischen Komponenten. Welchen Wert nimmt die Fourier-Reihe in $t = T$ an?

$$f(t) = \begin{cases} 2/T & 0 \leq t \leq T/2 \\ 1 & T/2 \leq t \leq T \end{cases}$$

Kapitel 15
Fourier-Transformation

15

15	Fourier-Transformation	611
15.1	Fourier-Transformation und Beispiele	613
15.1.1	Übergang von der Fourier-Reihe zur Fourier-Transformation	614
15.1.2	Inverse Fourier-Transformation	619
15.2	Eigenschaften der Fourier-Transformation	623
15.2.1	Linearität	623
15.2.2	Symmetrieeigenschaft	623
15.2.3	Skalierungseigenschaft	625
15.2.4	Verschiebungseigenschaften	625
15.2.5	Modulationseigenschaft	627
15.2.6	Fourier-Transformation der Ableitung	629
15.2.7	Faltungstheorem	630
15.3	Fourier-Transformation der Deltafunktion	637
15.3.1	Deltafunktion und Darstellung der Deltafunktion	637
15.3.2	Fourier-Transformation der Deltafunktion	639
15.3.3	Darstellung der Fourier-Transformierten von $\delta(t)$	642
15.3.4	Korrespondenzen der Fourier-Transformation	644
15.4	Aufgaben zur Fourier-Transformation	645

Zusätzliche Abschnitte auf der CD-Rom

15.5	MAPLE: Fourier-Transformation	cd
15.6	Beschreibung von linearen Systemen	cd
15.6.1	LZK-Systeme	cd
15.6.2	Impulsantwort	cd
15.6.3	Die Systemfunktion (Übertragungsfunktion)	cd
15.6.4	Übertragungsfunktion elektrischer Netzwerke	cd
15.6.5	Zusammenhang zwischen der Sprung- und Deltafunktion	cd
15.7	Anwendungsbeispiele mit MAPLE	cd
15.7.1	Frequenzanalyse des Doppelpendelsystems mit MAPLE	cd
15.7.2	Frequenzanalyse eines Hochpasses mit MAPLE	cd

15 Fourier-Transformation

In diesem Kapitel wird mit der Fourier-Transformation untersucht, welche Frequenzen mit welchen Amplituden in einem *nichtperiodischen* Zeitsignal $f(t)$ enthalten sind. Man nennt dieses Vorgehen, wie bei den Fourier-Reihen, die *Frequenzanalyse* des Zeitsignals f . In 15.1 werden die Formeln zur Fourier-Transformation

$$F(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt$$

hergeleitet und an Beispielen verdeutlicht. Es werden weiterhin in 15.2 wichtige Eigenschaften der Fourier-Transformation vorgestellt und deren Bedeutung diskutiert. Zur Charakterisierung von linearen Systemen benötigt man eine Funktion, die alle Frequenzen mit gleicher Amplitude enthält. Dies führt auf den Begriff der *Deltafunktion*, die wir in Abschnitt 15.3 einführen und deren Eigenschaften wir diskutieren.

Hinweis: Im CD-Abschnitt 15.4 wird die *Fourier-Transformation*, die inverse Fourier-Transformation und deren Anwendung beim Lösen von Differenzialgleichungen mit MAPLE vorgestellt.

In einem separaten Abschnitt 15.5 behandeln wir die vollständige Beschreibung von *linearen Systemen* durch die Impulsantwort und stellen den Zusammenhang zur Übertragungsfunktion mit Hilfe der Fourier-Transformation her. In Abschnitt 15.6 wenden wir mit MAPLE die lineare Theorie auf das mechanische Problem der Doppelpendel und das elektronische Problem des *Hochpasses 5. Ordnung* an, indem wir eine Frequenzanalyse durchführen.

Im Hinblick auf die digitale Messdatenanalyse wird in einem eigenständigen CD-Kapitel auf die *diskrete Fourier-Transformation* (DFT), die inversen DFT und deren Eigenschaften eingegangen. In diesem Kapitel behandeln wir auch die *diskrete Fourier-Transformation* mit MAPLE, indem wir den FFT-Algorithmus verwenden. *Anwendungsbeispiele* zur Signal- und Systemanalyse mit Hilfe der DFT findet man ebenfalls in diesem Abschnitt.

15.1 Fourier-Transformation und Beispiele

15.1

Durch die Angabe der Fourier-Reihe ist es möglich, periodische Vorgänge zu analysieren, ungeachtet der Tatsache, wie groß die Periodendauer T ist. Eine T -periodische Funktion f lässt sich darstellen als Überlagerung unendlich vieler harmonischer Schwingungen mit Grundfrequenz $\omega_0 = \frac{2\pi}{T}$, Oberfrequenzen $\omega_n = n \frac{2\pi}{T}$ und den zugehörigen Amplituden. Periodische Funktionen besitzen damit ein diskretes Linienspektrum. Dieses Linienspektrum liefert eine eindeutige Zuordnung zwischen dem Zeitbereich der Funktion und dem Frequenzbereich.

Im Folgenden wird ein Verfahren (*Fourier-Transformation*) entwickelt, welches auch für nichtperiodische Funktionen alle Frequenzen mit zugehörigen Amplituden liefert, die in dem Signal enthalten sind.

15.1.1 Übergang von der Fourier-Reihe zur Fourier-Transformation

Um die Frequenzen in einem beliebigen Zeitsignal f zu bestimmen, interpretieren wir die Funktion f als periodische Funktion mit Periode $T \rightarrow \infty$. Um eine Darstellung für das Spektrum zu erhalten, gehen wir aber von einer $2p$ -periodischen Funktion $f(t)$ aus. Nach Kapitel 14.5 lautet die zugehörige komplexe Fourier-Reihe

$$f(t) = \sum_{n=-\infty}^{\infty} c_n e^{i n \frac{2\pi}{2p} t}$$

mit den komplexen Fourier-Koeffizienten

$$c_n = \frac{1}{2p} \int_0^{2p} f(t) e^{-i n \frac{2\pi}{2p} t} dt = \frac{1}{2p} \int_{-p}^p f(t) e^{-i n \frac{\pi}{p} t} dt.$$

Setzt man die Koeffizienten c_n in die Fourier-Reihe ein, folgt mit den Frequenzen $\omega_n = n \frac{\pi}{p}$ und dem Frequenzabstand $\Delta\omega = \omega_n - \omega_{n-1} = \frac{\pi}{p}$:

$$\begin{aligned} f(t) &= \sum_{n=-\infty}^{\infty} \left[\frac{1}{2p} \int_{-p}^p f(t) e^{-i \omega_n t} dt \right] e^{i \omega_n t} \\ &= \frac{1}{2\pi} \sum_{n=-\infty}^{\infty} \left[\int_{-\pi/\Delta\omega}^{\pi/\Delta\omega} f(t) e^{-i \omega_n t} dt \right] e^{i \omega_n t} \Delta\omega. \end{aligned}$$

Wir interpretieren nun eine beliebige, nicht notwendigerweise periodische Zeitfunktion $f(t)$ als periodische Funktion mit $p \rightarrow \infty$. Für $p \rightarrow \infty$ geht der Frequenzabstand $\Delta\omega$ gegen Null und die Summe geht in das Integral über. Die Frequenzspektren rücken näher zusammen und nach dem Grenzübergang erhält man eine kontinuierliche Funktion in ω :

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i \omega t} d\omega \quad (FI)$$

mit

$$F(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i \omega t} dt. \quad (FT)$$

Man bezeichnet die Darstellung von $f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega t} d\omega$ als **Fourier-Integral** und $F(\omega)$ die **Fourier-Transformierte** oder **Spektralfunktion** zur Zeitfunktion $f(t)$.



Animation: Auf der CD-Rom befindet sich eine [MAPLE-Animation](#), bei der am Beispiel der T -periodischen Fortsetzung der Rechteckfunktion aufgezeigt wird, wie für $T \rightarrow \infty$ die Spektrallinien zusammenrücken ($\Delta\omega \rightarrow 0$) und aus dem diskreten Spektrum eine kontinuierliche Frequenzfunktion hervorgeht.

Beim Übergang von der Fourier-Reihe zum Fourier-Integral wird formal der Grenzübergang $p \rightarrow \infty$ durchgeführt, ohne die Existenz des uneigentlichen Integrals zu begründen. Man kann jedoch allgemein zeigen, dass das uneigentliche Integral unter gewissen Voraussetzungen (die in der Praxis nahezu immer erfüllt sind) existiert:

Satz 15.1: (Fourier-Transformation). Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stückweise stetig differenzierbar und $\int_{-\infty}^{\infty} |f(t)| dt < \infty$, dann existiert für jedes $\omega \in \mathbb{R}$

$$F(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt \quad (\text{Fourier-Transformierte von } f).$$

Die **Fourier-Transformation** ordnet jeder Zeitfunktion $f(t)$ mit obigen Eigenschaften eine Frequenzfunktion $F(\omega) : \mathbb{R} \rightarrow \mathbb{C}$ zu. Man sagt, dass die Fourier-Transformation den Zeitbereich auf den **Spektralbereich (Frequenzbereich)** abbildet, indem sie der Zeitfunktion $f(t)$ die Spektralfunktion $F(\omega)$ zuweist. Um präzise anzugeben, zu welcher Zeitfunktion $F(\omega)$ gehört, verwendet man auch die Notation

$$\mathcal{F}(f)(\omega) = F(\omega) .$$

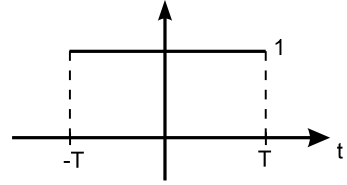
Diese Schreibweise drückt den transformatorischen Charakter der Fourier-Transformation aus: Der Funktion f wird eine neue Funktion $\mathcal{F}(f)$ zugeordnet. Teilweise wird auch die Korrespondenzschreibweise wie bei der Laplace-Transformation verwendet:

$$f(t) \circ \bullet F(\omega) .$$

Bemerkung: Eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ heißt stückweise stetig differenzierbar, wenn das Intervall in endlich viele Teilintervalle I_k zerlegt werden kann, so dass f im Innern der Intervalle I_k stetig differenzierbar ist und an den Grenzen der rechts- bzw. linksseitige Grenzwert von f existiert. f heißt auf $\mathbb{R}$ stückweise stetig differenzierbar, wenn f in jedem endlichen Intervall $[a, b] \subset \mathbb{R}$ stückweise stetig differenzierbar ist.

Beispiel 15.1. Für den unten dargestellten **Rechteckimpuls** soll die Fourier-Transformierte $F(\omega)$ berechnet werden.

$$f(t) := \begin{cases} 1 & \text{für } |t| < T \\ 0 & \text{für } |t| > T \end{cases} =: \text{rect}\left(\frac{t}{T}\right)$$



$$\begin{aligned} \mathcal{F}(f)(\omega) &= F(\omega) = \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt \\ &= \int_{-T}^T 1 e^{-i\omega t} dt = \left. \frac{e^{-i\omega t}}{-i\omega} \right|_{-T}^T \\ &= \frac{1}{-i\omega} (e^{-iT\omega} - e^{iT\omega}) = \frac{2}{\omega} \frac{1}{2i} (e^{iT\omega} - e^{-iT\omega}) \\ &= \frac{2}{\omega} \sin(\omega T) = 2T \frac{\sin(\omega T)}{\omega T}. \end{aligned}$$

⚠ Der Wert der Transformierten hängt nicht von der Wahl des Funktionswertes von $f(t)$ an den Stellen $t = -T$ bzw. $t = T$ ab!

Führt man die si-Funktion $\text{si}(x) := \frac{\sin x}{x}$ ein, schreibt man obiges Ergebnis als

$$\mathcal{F}\left(\text{rect}\left(\frac{t}{T}\right)\right)(\omega) = \frac{2}{\omega} \sin(\omega T) = 2T \text{si}(\omega T).$$

Der Verlauf der Fourier-Transformierten ist in Abb. 15.1 gezeigt.

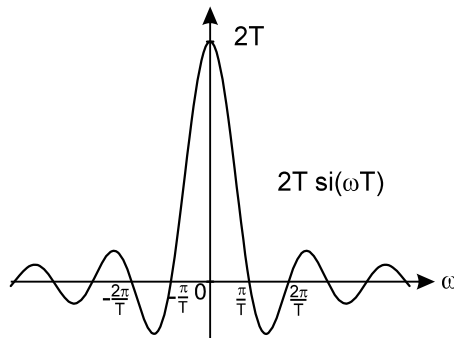


Abb. 15.1. Fourier-Transformierte der Rechteckfunktion

Das Spektrum ist eine Sinusfunktion in ω , dessen Amplitude mit $1/\omega$ abnimmt. Die Nullstellen von $F(\omega)$ sind bis auf $\omega = 0$ dieselben wie die der Sinusfunktion:

$$\omega_n T = n\pi \hookrightarrow \omega_n = n \frac{\pi}{T}.$$

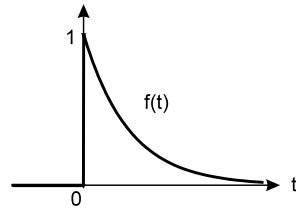
Für $\omega = 0$ muss die Regel von l'Hospital angewendet werden, um den Funktionswert zu berechnen:

$$F(\omega = 0) = \lim_{\omega \rightarrow 0} \frac{2 \sin(\omega T)}{\omega} = \lim_{\omega \rightarrow 0} \frac{2 \cos(\omega T) T}{1} = 2T.$$

Beobachtung: Je breiter das Rechteck $f(t)$ bzw. je größer T ist, desto schmaler wird das nullte Maximum von $F(\omega)$ bei $\omega = 0$, da die erste Nullstelle des Spektrums bei $\omega_1 = \frac{\pi}{T}$ liegt. Andererseits gilt: Je schmaler das Rechteck $f(t)$ d.h. je kleiner T ist, desto breiter wird das nullte Maximum. $\square$

Beispiel 15.2. Für die unten dargestellte **Exponentialfunktion** ist die zugehörige Spektralfunktion gesucht:

$$f(t) = e^{-\alpha t} \cdot S(t) = \begin{cases} 0 & \text{für } t < 0 \\ e^{-\alpha t} & \text{für } t > 0, \alpha > 0. \end{cases}$$



Dabei ist $S(t)$ die Sprungfunktion (Heavisidefunktion) $S(t) = \begin{cases} 0 & \text{für } t < 0 \\ 1 & \text{für } t > 0. \end{cases}$

Die zur Funktion f gehörende Spektralfunktion ist die Fourier-Transformierte $F(\omega)$:

$$\begin{aligned} F(\omega) &= \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt = \int_0^{\infty} e^{-\alpha t} e^{-i\omega t} dt = \int_0^{\infty} e^{-(\alpha+i\omega)t} dt \\ &= \left. \frac{e^{-(\alpha+i\omega)t}}{-(\alpha+i\omega)} \right|_{t=0}^{t=\infty} = \lim_{t \rightarrow \infty} \frac{e^{-(\alpha+i\omega)t}}{-(\alpha+i\omega)} + \frac{1}{\alpha+i\omega} \\ &= 0 + \frac{1}{\alpha+i\omega} = \frac{\alpha}{\alpha^2 + \omega^2} - i \frac{\omega}{\alpha^2 + \omega^2}. \end{aligned}$$

$$\Rightarrow F(\omega) = \mathcal{F}(e^{-\alpha t} \cdot S(t))(\omega) = \frac{1}{\alpha + i\omega}.$$

Anmerkung: Für die meisten Anwendungen spielt es keine Rolle, welchen Wert $S(t)$ an der Stelle $t = 0$ besitzt. Gebräuchlich sind $S(0) = 0$, $S(0) = 1$ aber auch $S(0) = \frac{1}{2}$. Mit der letzteren Festlegung gilt $S(t) = \frac{1}{2} + \frac{1}{2} \text{sign}(t)$, wenn $\text{sign}(t)$ die Vorzeichenfunktion darstellt. $\square$

⊗ **Darstellung der Fourier-Transformierten**

Man erkennt an diesem einfachen Beispiel, dass die Fourier-Transformierte $F(\omega)$ einer Zeitfunktion $f(t)$ i.A. komplexwertig ist. Den Graphen komplexwertiger Funktionen kann man nicht direkt zeichnen, sondern man muss entweder Real- und Imaginärteil getrennt darstellen oder man zerlegt die komplexe Funktion $F(\omega)$ in **Betrag** und **Phase**

$$F(\omega) = |F(\omega)| e^{i\varphi(\omega)} \quad \text{mit}$$

$$|F(\omega)| = \sqrt{F(\omega) \cdot F^*(\omega)} \quad (\text{Betrag})$$

$$\tan \varphi(\omega) = \frac{\operatorname{Im} F(\omega)}{\operatorname{Re} F(\omega)} \quad (\text{Phase})$$

(siehe Kapitel 5). Hierbei ist sowohl der Betrag als auch die Phase eine reelle Funktion von ω . Man spricht analog der Bezeichnung bei den Fourier-Reihen auch von *Amplituden-* und *Phasenspektrum*.

Beispiel 15.3 (Amplituden- und Phasenspektrum). Für Beispiel 15.2 erhält man für die Fourier-Transformierte

$$F(\omega) = \frac{1}{\alpha + i\omega}.$$

Wir zerlegen $F(\omega)$ in Betrag

$$|F(\omega)| = \sqrt{F(\omega) \cdot F^*(\omega)} = \sqrt{\frac{1}{\alpha + i\omega} \cdot \frac{1}{\alpha - i\omega}} = \frac{1}{\sqrt{\alpha^2 + \omega^2}}$$

und Phase

$$\tan \varphi(\omega) = \frac{\operatorname{Im} F(\omega)}{\operatorname{Re} F(\omega)} = -\frac{\omega}{\alpha}.$$

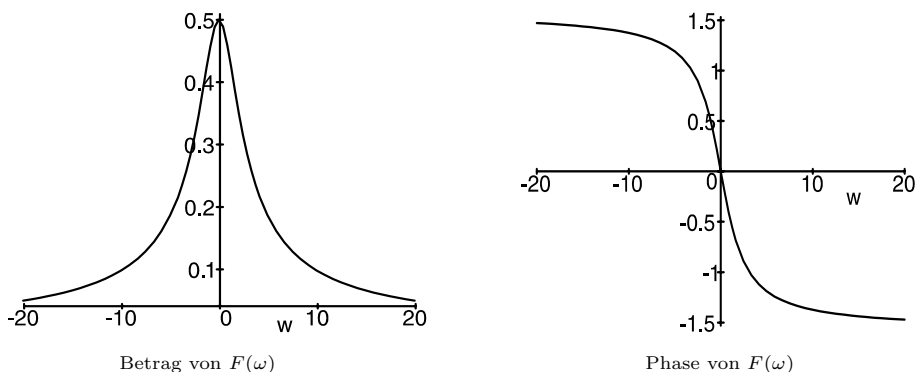


Abb. 15.2. Betrag und Phase der Fourier-Transformierten

In Abb. 15.2 sind sowohl das Amplituden- als auch Phasenspektrum für $\alpha = 1$ angegeben. □

► 15.1.2 Inverse Fourier-Transformation

Dass eine Zeitfunktion vollständig durch ihr Spektrum charakterisiert wird, zeigt sich durch die *inverse Fourier-Transformation*.

Satz 15.2: (Inverse Fourier-Transformation). Ist $F(\omega)$ die Fourier-Transformierte einer Funktion $f(t)$, so ist die Funktion $f(t)$ gegeben durch

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega t} d\omega \quad (\text{inverse Fourier-Transformation}),$$

falls $f(t)$ die Mittelwerteigenschaft erfüllt.

Die Mittelwerteigenschaft einer Funktion f besagt, dass der Funktionswert an jeder Stelle t durch den Mittelwert des linksseitigen und rechtsseitigen Funktionsgrenzwertes gegeben ist: $f(t) = \lim_{\varepsilon \rightarrow 0} \frac{1}{2} (f(t + \varepsilon) + f(t - \varepsilon))$ für alle $t \in \mathbb{D}_f$ (siehe Abschnitt 14.3). Für stetige Funktionen ist die Mittelwerteigenschaft immer erfüllt, für unstetige Funktionen muss an der Sprungstelle genau der Mittelwert des Sprunges als Funktionswert angenommen werden.

Um mit der Mittelwerteigenschaft konsistent zu sein, sollte die Sprungfunktion $S(t)$ für die Anwendungen bei der Fourier-Transformation an der Stelle $t_0 = 0$ durch $S(0) = \frac{1}{2}$ definiert werden!

Satz 15.2 impliziert, dass in der Fourier-Transformierten einer Funktion f dieselbe Information enthalten ist, wie in f selbst. Denn durch die Formel der inversen Fourier-Transformation kann das Zeitsignal $f(t)$ vollständig aus $F(\omega)$ rekonstruiert werden. Insbesondere kann die Funktion f alleine durch die zugehörige Spektralfunktion charakterisiert werden!

Bemerkungen:

- (1) Die Fourier-Transformierte ist auch für komplexwertige Funktionen $f(t) = f_1(t) + i f_2(t)$ definiert. $F(\omega)$ ist im Allgemeinen **immer** eine komplexwertige Funktion (siehe Beispiel 15.2).
- (2) Ist f eine reellwertige Funktion (ein reelles Signal), dann lässt sich die Fourier-Transformierte mit der Eulerschen Beziehung $e^{-i\omega t} = \cos(\omega t) - i \sin(\omega t)$ auch schreiben als

$$\begin{aligned} \mathcal{F}(f)(\omega) &= \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt = \int_{-\infty}^{\infty} f(t) (\cos(\omega t) - i \sin(\omega t)) dt \\ &= \int_{-\infty}^{\infty} f(t) \cos(\omega t) dt - i \int_{-\infty}^{\infty} f(t) \sin(\omega t) dt. \quad (*) \end{aligned}$$

Man spricht in diesem Zusammenhang oftmals von der *Kosinus-* und *Sinustransformierten*.

(3) Gerade und ungerade Funktionen

- (i) Für eine **gerade**, reelle Funktion f , d.h. $f(-t) = f(t)$, gilt

$$F(\omega) = 2 \int_0^{\infty} f(t) \cos(\omega t) dt.$$

Begründung: Mit $f(t)$ ist auch $f(t) \cdot \cos(\omega t)$ eine gerade Funktion und

$$\int_{-\infty}^{\infty} f(t) \cos(\omega t) dt = 2 \int_0^{\infty} f(t) \cos(\omega t) dt ,$$

da der Integrand symmetrisch zum Ursprung $t = 0$ integriert wird. Andererseits ist $f(t) \cdot \sin(\omega t)$ eine ungerade Funktion. Integriert man eine ungerade Funktion symmetrisch zum Ursprung, ist das bestimmte Integral Null:

$$\int_{-\infty}^{\infty} f(t) \sin(\omega t) dt = 0 .$$

Setzt man diese Integralergebnisse in die Formel (*) ein, folgt die Behauptung. $\square$

- (ii) Für eine **ungerade**, reelle Funktion f , d.h. $f(-t) = -f(t)$, gilt mit der zu (i) analogen Begründung

$$F(\omega) = -i 2 \int_0^{\infty} f(t) \sin(\omega t) dt.$$

Beispiel 15.4. Gegeben ist die gerade, reelle Funktion

$$f(t) = e^{-\alpha |t|} \quad \text{mit } \alpha > 0 .$$

Die Fourier-Transformierte von f berechnet sich nach Bemerkung 3 (i) aus

$$F(\omega) = 2 \int_0^{\infty} f(t) \cos(\omega t) dt = 2 \int_0^{\infty} e^{-\alpha t} \cos(\omega t) dt .$$

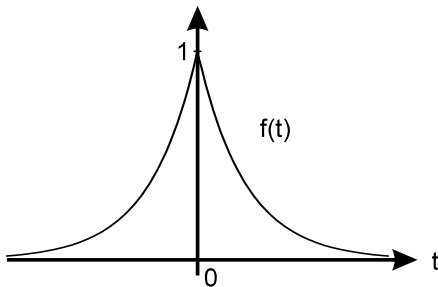
Die Beträge im Argument der Exponentialfunktion können weggelassen werden, da nur über positive t integriert wird. Zur Berechnung des Integrals er-

setzen wir $\cos(\omega t)$ durch $\operatorname{Re}(e^{i\omega t})$ und erhalten

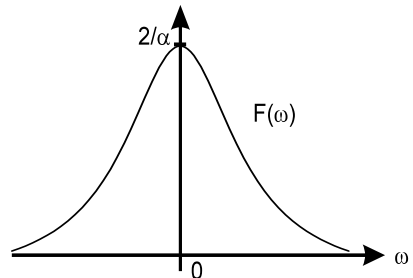
$$\begin{aligned} F(\omega) &= 2 \int_0^{\infty} e^{-\alpha t} \operatorname{Re}(e^{i\omega t}) dt = 2 \operatorname{Re} \int_0^{\infty} e^{(-\alpha + i\omega)t} dt \\ &= 2 \operatorname{Re} \left. \frac{e^{(-\alpha + i\omega)t}}{-\alpha + i\omega} \right|_{t=0}^{t=\infty} = 2 \operatorname{Re} \frac{-1}{-\alpha + i\omega} \\ &= 2 \operatorname{Re} \left\{ \frac{\alpha}{\alpha^2 + \omega^2} + i \frac{\omega}{\alpha^2 + \omega^2} \right\} = \frac{2\alpha}{\alpha^2 + \omega^2}. \end{aligned}$$

Das gleiche Ergebnis bekommt man auch durch zweimalige partielle Integration.

$$\Rightarrow F(\omega) = \mathcal{F}(e^{-\alpha|t|})(\omega) = \frac{2\alpha}{\alpha^2 + \omega^2}.$$



Zeitfunktion $f(t)$



Spektralfunktion $F(\omega)$

Abb. 15.3. Zeitfunktion und zugehörige Spektralfunktion

Beispiel 15.5. Gegeben ist die ungerade, reelle Funktion

$$f(t) = e^{-\alpha|t|} \operatorname{sign}(t)$$

mit $\alpha > 0$ und der Vorzeichenfunktion $\operatorname{sign}(t) = \begin{cases} 1 & \text{für } t > 0 \\ 0 & \text{für } t = 0 \\ -1 & \text{für } t < 0 \end{cases}$.

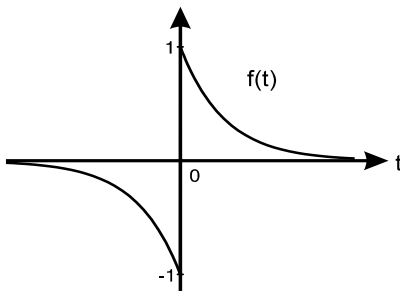
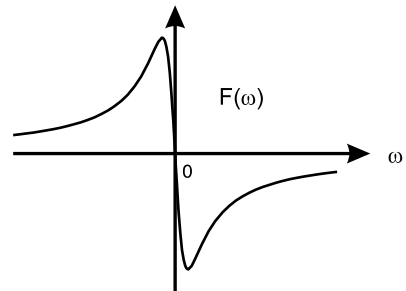
Die Fourier-Transformierte ist nach Bemerkung 3 (ii) gegeben durch

$$F(\omega) = -i 2 \int_0^{\infty} f(t) \sin(\omega t) dt = -i 2 \int_0^{\infty} e^{-\alpha t} \sin(\omega t) dt.$$

Ersetzt man zur einfachen Berechnung des Integrals $\sin(\omega t) = \operatorname{Im}(e^{i\omega t})$, so erhält man unter Berücksichtigung von Beispiel 15.4

$$F(\omega) = -i 2 \operatorname{Im} \int_0^{\infty} e^{-\alpha t} e^{i\omega t} dt = -i 2 \operatorname{Im} \left\{ \frac{\alpha}{\alpha^2 + \omega^2} + i \frac{\omega}{\alpha^2 + \omega^2} \right\}$$

$$\Rightarrow F(\omega) = -i \frac{2\omega}{\alpha^2 + \omega^2}.$$

Zeitfunktion $f(t)$ Fourier-Transformierte $F(\omega)$ **Abb. 15.4.** Zeitfunktion und zugehörige Fourier-Transformierte

Beispiel 15.6. Gegeben ist die ungerade, reelle Funktion

$$f(t) = \frac{1}{t}.$$

Obwohl die Funktion f für $t = 0$ nicht definiert ist (Polstelle), ist $\frac{\sin(\omega t)}{t}$ für alle $t \in \mathbb{R}$ stetig und beschränkt. Es gilt für die Fourier-Transformierte von $\frac{1}{t}$ nach Bemerkung 3 (ii) mit der Substitution ($x = \omega t \leftrightarrow dx = \omega dt$)

$$F(\omega) = -i 2 \int_0^{\infty} \frac{\sin(\omega t)}{t} dt = -i 2 \int_0^{\infty} \frac{\sin x}{x} dx \quad (\omega > 0).$$

Der Wert des bestimmten Integrals $\int_0^{\infty} \frac{\sin x}{x} dx = \frac{\pi}{2}$ wird im Beispiel 15.8 ③ berechnet. Unter Verwendung dieses Ergebnisses ist die Fourier-Transformierte

$$F(\omega) = \begin{cases} -i\pi & \text{für } \omega > 0 \\ i\pi & \text{für } \omega < 0 \end{cases} = -i\pi \operatorname{sign}(\omega),$$

wenn sign die Vorzeichenfunktion darstellt.

□

15.2 Eigenschaften der Fourier-Transformation

In diesem Abschnitt werden wichtige Eigenschaften der Fourier-Transformation vorgestellt und an Beispielen verdeutlicht. Im Folgenden wird immer davon ausgegangen, dass die Zeitfunktionen die Voraussetzungen der Fourier-Transformation erfüllen, so dass die Fourier-Transformierten der Funktionen definiert sind. $F(\omega) = \mathcal{F}(f)(\omega)$ sei stets die Fourier-Transformierte von f ; $F_1(\omega) = \mathcal{F}(f_1)(\omega)$ und $F_2(\omega) = \mathcal{F}(f_2)(\omega)$ die Fourier-Transformierten von f_1 bzw. f_2 .

► 15.2.1 Linearität

Gesucht ist das Spektrum einer Überlagerung $k_1 f_1(t) + k_2 f_2(t)$:

$$\begin{aligned}\mathcal{F}(k_1 f_1 + k_2 f_2)(\omega) &= \int_{-\infty}^{\infty} (k_1 f_1(t) + k_2 f_2(t)) e^{-i\omega t} dt \\ &= k_1 \int_{-\infty}^{\infty} f_1(t) e^{-i\omega t} dt + k_2 \int_{-\infty}^{\infty} f_2(t) e^{-i\omega t} dt \\ &= k_1 F_1(\omega) + k_2 F_2(\omega).\end{aligned}$$

Zusammenfassend gilt

$$(F_1) \quad \textbf{Linearität:} \quad \mathcal{F}(k_1 f_1 + k_2 f_2)(\omega) = k_1 F_1(\omega) + k_2 F_2(\omega)$$

Die Linearität besagt, dass das Spektrum einer Überlagerung von zwei Zeitfunktionen sich aus der entsprechenden Überlagerung der Einzelspektren zusammensetzt.

Beispiel 15.7. Gesucht ist das Spektrum der Funktion $4 \operatorname{rect}\left(\frac{t}{T}\right) + 3e^{-\alpha|t|}$:
Nach Beispiel 15.1 und 15.4 gilt mit der Eigenschaft (F_1)

$$\begin{aligned}\mathcal{F}\left(4 \operatorname{rect}\left(\frac{t}{T}\right) + 3e^{-\alpha|t|}\right)(\omega) &= 4 \mathcal{F}\left(\operatorname{rect}\left(\frac{t}{T}\right)\right)(\omega) + 3 \mathcal{F}\left(e^{-\alpha|t|}\right)(\omega) \\ &= 8 \frac{\sin(\omega T)}{\omega} + 6 \frac{\alpha}{\alpha^2 + \omega^2}.\end{aligned}\quad \square$$

► 15.2.2 Symmetrieeigenschaft

$$(F_2) \quad \textbf{Symmetrie:} \quad \mathcal{F}(\mathcal{F}(f))(t) = 2\pi f(-t)$$

Die Symmetrieeigenschaft besagt, dass die Fourier-Transformation zweimal auf die Funktion f angewendet, wieder die Funktion f als Ergebnis liefert, allerdings mit dem Faktor 2π und negativem Argument. Denn aufgrund des Fourier-Integrals (FI) gilt

$$\mathcal{F}(F)(t) = \int_{-\infty}^{\infty} F(\omega) e^{-i\omega t} d\omega = 2\pi \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega) e^{i\omega(-t)} d\omega = 2\pi f(-t).$$

Die Symmetrie wird ausgenutzt, um die Fourier-Transformation von Funktionen zu berechnen, die selbst Fourier-Transformierte sind:

Beispiele 15.8.

- ① Wegen $\mathcal{F}\left(\operatorname{rect}\left(\frac{t}{a}\right)\right)(\omega) = 2 \frac{\sin(\omega a)}{\omega}$ folgt für die achsensymmetrische Rechteck-Funktion rect :

$$\begin{aligned}\mathcal{F}\left(2 \frac{\sin(\omega a)}{\omega}\right)(t) &= \mathcal{F}\left(\mathcal{F}\left(\operatorname{rect}\left(\frac{t}{a}\right)\right)\right)(t) \\ &= 2\pi \operatorname{rect}\left(-\frac{t}{a}\right) = 2\pi \operatorname{rect}\left(\frac{t}{a}\right) .\end{aligned}$$

Damit erhält man nach dem Vertauschen der Variablen ω und t

$$\mathcal{F}\left(\frac{\sin(a t)}{t}\right)(\omega) = \pi \operatorname{rect}\left(\frac{\omega}{a}\right).$$

- ② Aus Beispiel 15.4 erhält man mit $\alpha = 1$

$$\mathcal{F}\left(e^{-|t|}\right)(\omega) = \frac{2}{1 + \omega^2} .$$

Mit (F_2) gilt dann

$$\mathcal{F}\left(\frac{2}{1 + \omega^2}\right)(t) = \mathcal{F}\left(\mathcal{F}\left(e^{-|t|}\right)\right)(t) = 2\pi e^{-|t|} .$$

Nach Vertauschung der Variablen ist

$$\mathcal{F}\left(\frac{1}{1 + t^2}\right)(\omega) = \pi e^{-|\omega|} .$$

- ③ Wir berechnen das bestimmte Integral $\int_0^\infty \frac{\sin x}{x} dx = \frac{\pi}{2}$ mit Hilfe der Fourier-Transformation. Nach Beispiel 15.8 ① ist

$$\mathcal{F}\left(\frac{\sin t}{t}\right)(\omega) = \pi \operatorname{rect}(\omega) .$$

Nach Bemerkung 3 (i) gilt aber auch für die gerade Funktion $\frac{\sin t}{t}$

$$\mathcal{F}\left(\frac{\sin t}{t}\right)(\omega) = 2 \int_0^\infty \frac{\sin t}{t} \cos(\omega t) dt .$$

Folglich gilt für $\omega = 0$:

$$2 \int_0^\infty \frac{\sin t}{t} dt = \pi \operatorname{rect}(0) = \pi ,$$

woraus der Wert für das bestimmte Integral folgt. □

➤ 15.2.3 Skalierungseigenschaft

Gesucht ist die Fourier-Transformation (das Spektrum) der Funktion $f(at)$, die aus $f(t)$ durch **Stauchung** ($a > 1$) bzw. **Streckung** ($0 < a < 1$) entsteht.

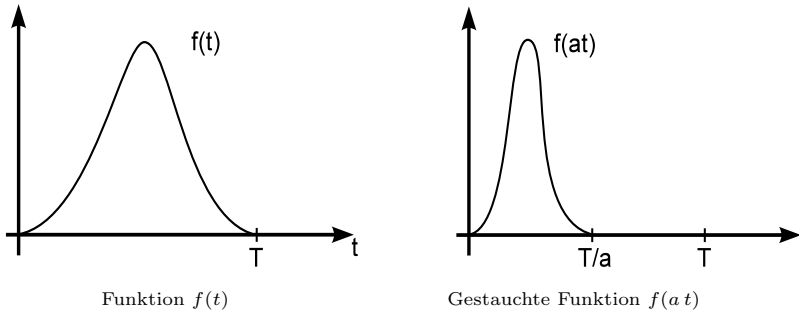


Abb. 15.5. Zur Skalierungseigenschaft

Die Fourier-Transformierte von f berechnet sich mit der Substitution $\xi = at$ für $a > 0$ durch

$$\mathcal{F}(f(at))(\omega) = \int_{-\infty}^{\infty} f(at) e^{-i\omega t} dt = \int_{-\infty}^{\infty} f(\xi) e^{-i\omega \frac{\xi}{a}} \frac{d\xi}{a} = \frac{1}{a} F\left(\frac{\omega}{a}\right).$$

Die Argumentation für $a < 0$ verläuft analog, daher gilt insgesamt

(F₃) Skalierung: $\mathcal{F}(f(at))(\omega) = \frac{1}{|a|} F\left(\frac{\omega}{a}\right), \quad a \in \mathbb{R}_{\neq 0}$

Die Skalierungseigenschaft besagt, dass in einem gestauchten Signal $f(at)$ ($a > 1$) die Frequenzen $F\left(\frac{\omega}{a}\right)$ vorkommen.

➤ 15.2.4 Verschiebungseigenschaften

➤ Zeitverschiebung

Der Verschiebungssatz macht eine Aussage über die Fourier-Transformierte einer zeitlich verschobenen Funktion $f(t - t_0)$.

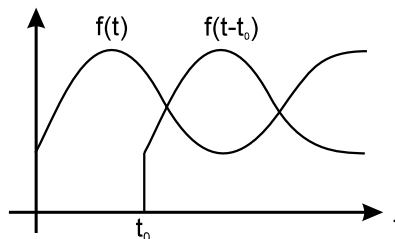


Abb. 15.6. Funktion $f(t)$ und um t_0 verschobene Funktion $f(t - t_0)$

Es gilt mit der Substitution $\xi = t - t_0$:

$$\mathcal{F}(f(t - t_0))(\omega) = \int_{-\infty}^{\infty} f(t - t_0) e^{-i\omega t} dt = \int_{-\infty}^{\infty} f(\xi) e^{-i\omega(\xi + t_0)} d\xi.$$

Spalten wir den Exponentialterm in zwei Faktoren auf und klammern $e^{-i\omega t_0}$ aus dem Integral aus (er ist unabhängig von der Integrationsvariablen), gilt weiter

$$\mathcal{F}(f(t - t_0))(\omega) = e^{-i\omega t_0} \int_{-\infty}^{\infty} f(\xi) e^{-i\omega \xi} d\xi = e^{-i\omega t_0} F(\omega).$$

Folglich gilt für die Fourier-Transformierte einer zeitverschobenen Funktion:

$$(F_4) \quad \text{Zeitverschiebung:} \quad \mathcal{F}(f(t - t_0))(\omega) = e^{-i\omega t_0} F(\omega)$$

Setzt man $F(\omega) = |F(\omega)| e^{i\varphi(\omega)}$, so ergibt sich für das Spektrum der zeitverschobenen Funktion $f(t - t_0)$

$$\mathcal{F}(f(t - t_0))(\omega) = e^{-i\omega t_0} |F(\omega)| e^{i\varphi(\omega)} = |F(\omega)| e^{-i(\varphi(\omega) - \omega t_0)}.$$

Das Spektrum von $f(t - t_0)$ besitzt dieselbe Amplitude wie $f(t)$ nur die Phase ist um ωt_0 verschoben. D.h. es kommen dieselben Frequenzen mit gleicher Amplitude aber phasenverschoben vor.

Beispiel 15.9. Gesucht ist das Spektrum des um $t_0 = T$ verschobenen Rechtecksignals $f(t) = \text{rect}\left(\frac{t}{T}\right)$:

$$\mathcal{F}(f(t - T)) = e^{-i\omega T} \mathcal{F}\left(\text{rect}\left(\frac{t}{T}\right)\right)(\omega) = e^{-i\omega T} 2 \frac{\sin(\omega T)}{\omega}. \quad \square$$

⊗ Frequenzverschiebung

Eine weitere Eigenschaft ist die der Frequenzverschiebung. Diese Eigenschaft trifft eine Aussage über das Spektrum der Funktion $e^{i\omega_0 t} f(t)$:

$$\begin{aligned} \mathcal{F}(e^{i\omega_0 t} f(t))(\omega) &= \int_{-\infty}^{\infty} e^{i\omega_0 t} f(t) e^{-i\omega t} dt \\ &= \int_{-\infty}^{\infty} f(t) e^{-i(\omega - \omega_0) t} dt = F(\omega - \omega_0). \end{aligned}$$

Folglich gilt für das Spektrum einer mit $e^{i\omega_0 t}$ multiplizierten Funktion:

$$(F_5) \quad \text{Frequenzverschiebung:} \quad \mathcal{F}(e^{i\omega_0 t} f(t))(\omega) = F(\omega - \omega_0)$$

Die Verschiebungseigenschaft besagt, dass das Spektrum der mit $e^{i\omega_0 t}$ multiplizierten Funktion dasselbe ist, wie das um ω_0 verschobene Spektrum der ursprünglichen Funktion. Als Beispiel und Anwendung der Frequenzverschiebung diskutieren wir die Modulationseigenschaft:

► 15.2.5 Modulationseigenschaft

Gesucht ist das Spektrum des *amplitudenmodulierten* Signals

$$f(t) \cos(\omega_0 t) .$$

Mit der Eulerschen Formel

$$\cos(\omega_0 t) = \frac{1}{2} (e^{i\omega_0 t} + e^{-i\omega_0 t})$$

berechnen wir unter Verwendung der Linearität (F_1) und der Frequenzverschiebungseigenschaft (F_5)

$$\begin{aligned} \mathcal{F}(\cos(\omega_0 t) f(t))(\omega) &= \mathcal{F}\left(\frac{1}{2} e^{i\omega_0 t} f(t) + \frac{1}{2} e^{-i\omega_0 t} f(t)\right)(\omega) \\ &= \frac{1}{2} \mathcal{F}(e^{i\omega_0 t} f(t))(\omega) + \frac{1}{2} \mathcal{F}(e^{-i\omega_0 t} f(t))(\omega) \\ &= \frac{1}{2} (F(\omega - \omega_0) + F(\omega + \omega_0)). \end{aligned}$$

$$(F_6) \text{ Modulation: } \mathcal{F}(f(t) \cos(\omega_0 t))(\omega) = \frac{1}{2} (F(\omega + \omega_0) + F(\omega - \omega_0))$$

Das mit $\cos(\omega_0 t)$ amplitudenmodulierte Signal besitzt als Spektrum das um $\pm\omega_0$ verschobene Spektrum der ursprünglichen Funktion f . Die Verschiebung des Spektrums von f erfolgt genau um die Modulationsfrequenz ω_0 !

Beispiele 15.10. Gesucht ist das Spektrum eines mit $\cos(\omega_0 t)$ modulierten Rechteckimpulses. Nach Beispiel 15.1 ist das Spektrum des Rechtecks

$$\mathcal{F}\left(\text{rect}\left(\frac{t}{T}\right)\right)(\omega) = 2 \frac{\sin(\omega T)}{\omega}.$$

Damit folgt mit der Modulationseigenschaft

$$\mathcal{F}\left(\cos(\omega_0 t) \cdot \text{rect}\left(\frac{t}{T}\right)\right)(\omega) = \frac{\sin(T(\omega - \omega_0))}{\omega - \omega_0} + \frac{\sin(T(\omega + \omega_0))}{\omega + \omega_0} .$$

Die Amplitudenmodulation des Rechtecksignals entspricht einer Verschiebung des Spektrums um die Modulationsfrequenz ω_0 nach rechts und links. Beide Teilspektren besitzen die halbe Amplitude.

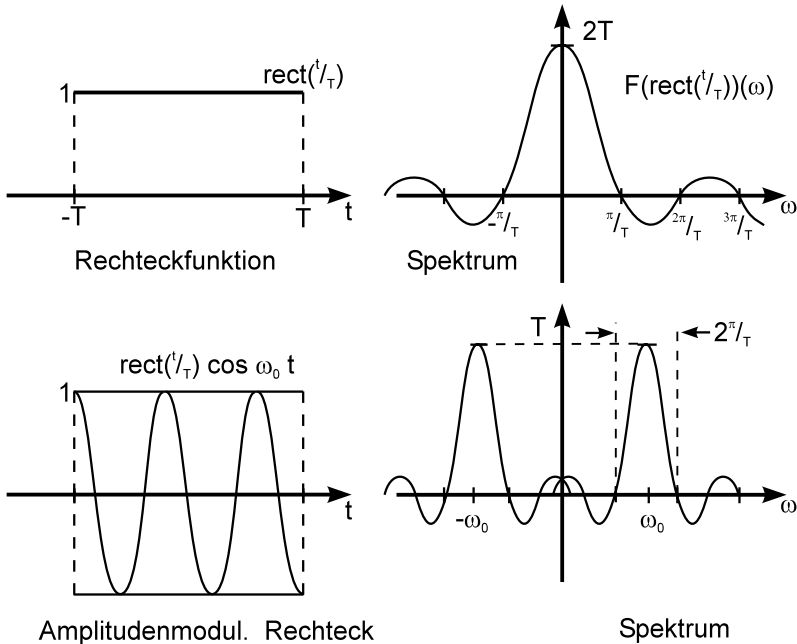


Abb. 15.7. Spektrum der amplitudenmodulierten Rechteckfunktion

Bemerkungen / Interpretation:

- (1) Bei der Übertragung von Nachrichten wird das Signal $f(t)$ oftmals amplitudenmoduliert, d.h. mit einer Trägerfrequenz ω_0 übertragen: $f(t) \cos(\omega_0 t)$. Ein Empfänger kann nun die Trägerfrequenz bestimmen, indem als Testsignal ein amplitudenmodulierter Rechteckimpuls übertragen wird. Das Spektrum dieses Signals ist dann das Spektrum des Rechteckimpulses um den Wert der Trägerfrequenz ω_0 verschoben.
- (2) Bei der experimentellen Analyse von periodischen Vorgängen erhält man in der Regel kein Linienspektrum, sondern eine Verbreiterung der Linie. Dieser Effekt lässt sich mit unserem Beispiel erklären:

Da man nicht für alle Zeiten $-\infty < t < \infty$ misst, sondern nur in einem endlichen Zeitintervall, entspricht dies der Analyse der Funktion $\text{rect}\left(\frac{t}{T}\right) \cdot \cos(\omega_0 t)$ statt der Analyse von $\cos(\omega_0 t)$. Obiges Beispiel 15.10 zeigt, dass man dann ein Spektrum der Form $2 \frac{\sin(\omega T)}{\omega}$ erhält, welches um ω_0 verschoben ist. Dieses Spektrum hat zwar bei ω_0 sein Maximum, aber eine endliche Breite $\frac{2\pi}{T}$. Nur im Falle $T \rightarrow \infty$ geht die Spektrenbreite gegen 0 und man erhält eine Linie bei ω_0 . $\square$

► 15.2.6 Fourier-Transformation der Ableitung

Für die Anwendung der Fourier-Transformation auf Differenzialgleichungen benötigt man die Fourier-Transformierte der Ableitung $\mathcal{F}(f')$. Es gibt wie im Falle der Laplace-Transformation ($\rightarrow$ 13.3.2) einen sehr einfachen Zusammenhang zwischen $\mathcal{F}(f')$ und $\mathcal{F}(f)$. Mit partieller Integration folgt für $\mathcal{F}(f')$

$$\begin{aligned}\mathcal{F}(f')(\omega) &= \int_{-\infty}^{\infty} f'(t) e^{-i\omega t} dt \\ &= f(t) e^{-i\omega t} \Big|_{-\infty}^{\infty} - \int_{-\infty}^{\infty} f(t) (-i\omega) e^{-i\omega t} dt.\end{aligned}$$

Wegen dem Abklingverhalten von f , $\lim_{t \rightarrow \pm\infty} f(t) = 0$, ist $f(t) e^{-i\omega t} \Big|_{-\infty}^{\infty} = 0$.

$$\Rightarrow \mathcal{F}(f')(\omega) = i\omega \int_{-\infty}^{\infty} f(t) e^{-i\omega t} dt = i\omega \mathcal{F}(f)(\omega).$$

Satz 15.3: (Fourier-Transformierte der Ableitung).

Für die Fourier-Transformierte der Ableitung $\mathcal{F}(f')$ gilt:

$$(F_7) \quad \textbf{Ableitung:} \quad \mathcal{F}(f')(\omega) = (i\omega) F(\omega)$$

Das Spektrum der differenzierten Funktion f' ist gleich dem mit $i\omega$ multiplizierten Spektrum der Funktion f . Wiederholtes Anwenden des Ableitungssatzes führt induktiv auf die Fourier-Transformierte der n -ten Ableitung:

Satz 15.4: (Fourier-Transformierte der n -ten Ableitung).

Für die Fourier-Transformierte der n -ten Ableitung $\mathcal{F}(f^{(n)})$ gilt:

$$(F_8) \quad \textbf{n-te Ableitung:} \quad \mathcal{F}(f^{(n)})(\omega) = (i\omega)^n F(\omega)$$

Beispiel 15.11. Gegeben ist die lineare Differenzialgleichung

$$y'(t) + \alpha y(t) = f(t) S(t)$$

mit dem konstanten Koeffizienten α und stetiger Funktion f . $S(t)$ ist die Sprungfunktion. Wenden wir auf diese Differenzialgleichung die Fourier-Transformation an und nutzen die Linearität (F_1) aus, folgt für die linke Seite

$$\mathcal{F}(y'(t) + \alpha y(t)) = \mathcal{F}(y'(t)) + \alpha \mathcal{F}(y(t)).$$

Wir ersetzen $\mathcal{F}(y'(t)) = i\omega \mathcal{F}(y(t))$ und erhalten insgesamt

$$i\omega \mathcal{F}(y(t)) + \alpha \mathcal{F}(y(t)) = \mathcal{F}(f(t) S(t))$$

$$\Rightarrow \mathcal{F}(y(t)) = \frac{1}{\alpha + i\omega} \mathcal{F}(f(t) S(t)) .$$

Dies ist die Fourier-Transformierte der gesuchten Lösung $y(t)$ der Differenzialgleichung. Sie ist gegeben als das Produkt von $\mathcal{F}(f(t) S(t))$ mit $\frac{1}{\alpha + i\omega}$. Nach Beispiel 15.2 ist

$$\frac{1}{\alpha + i\omega} = \mathcal{F}(e^{-\alpha t} S(t))(\omega) .$$

$$\Rightarrow \mathcal{F}(y(t)) = \mathcal{F}(f(t) S(t)) \cdot \mathcal{F}(e^{-\alpha t} S(t)) .$$

Es stellt sich also das Problem: Welche Zeitfunktion gehört zu einem Produkt von Frequenzfunktionen. Die Antwort liefert das Faltungstheorem:

► 15.2.7 Faltungstheorem

Gegeben ist das Spektrum der Funktion f als Produkt von zwei Frequenzfunktionen

$$\mathcal{F}(f) = \mathcal{F}(f_1) \cdot \mathcal{F}(f_2) .$$

Die gesuchte Zeitfunktion $f(t)$ ist dann eine Integralkombination der Zeitfunktionen $f_1(t)$ und $f_2(t)$

$$f(t) = \int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) d\tau ,$$

dem sog. **Faltungsintegral**. Die abkürzende Schreibweise für das Faltungsintegral ist

$$f(t) = (f_1 * f_2)(t) .$$

Faltungstheorem: Die Fourier-Transformierte des Faltungsintegrals

$$(f_1 * f_2)(t) := \int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) d\tau$$

ist gegeben durch das Produkt der Transformierten von f_1 und f_2 :

$$(F_9) \quad \textbf{Faltungstheorem:} \quad \mathcal{F}(f_1 * f_2) = \mathcal{F}(f_1) \cdot \mathcal{F}(f_2)$$

Begründung:

$$\begin{aligned}\mathcal{F}(f_1 * f_2) &= \int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) d\tau \right) e^{-i\omega t} dt \\ &= \int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) e^{-i\omega t} d\tau \right) dt .\end{aligned}$$

Nach Vertauschen der Integrationsreihenfolge und anschließender Substitution $\xi(t) = t - \tau \quad (\hookrightarrow d\xi = dt)$ folgt

$$\begin{aligned}\mathcal{F}(f_1 * f_2) &= \int_{-\infty}^{\infty} f_1(\tau) \left(\int_{-\infty}^{\infty} f_2(t - \tau) e^{-i\omega t} dt \right) d\tau \\ &= \int_{-\infty}^{\infty} f_1(\tau) \left(\int_{-\infty}^{\infty} f_2(\xi) e^{-i\omega(\xi + \tau)} d\xi \right) d\tau \\ &= \int_{-\infty}^{\infty} f_1(\tau) e^{-i\omega\tau} d\tau \cdot \int_{-\infty}^{\infty} f_2(\xi) e^{-i\omega\xi} d\xi \\ &= \mathcal{F}(f_1) \cdot \mathcal{F}(f_2) .\end{aligned} \quad \square$$

Beispiel 15.12. Gesucht ist nach Beispiel 15.11 die Zeitfunktion $y(t)$, die zum Spektrum von

$$\mathcal{F}(f(t) S(t)) \cdot \mathcal{F}(e^{-\alpha t} S(t))$$

gehört. Nach dem Faltungstheorem ist die Zeitfunktion $y(t)$ die Faltung der beiden Funktionen $f_1(t) = f(t) S(t)$ und $f_2(t) = e^{-\alpha t} S(t)$:

$$\begin{aligned}y(t) &= (f_1 * f_2)(t) = \int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) d\tau \\ &= \int_{-\infty}^{\infty} f(\tau) S(\tau) e^{-\alpha(t - \tau)} S(t - \tau) d\tau .\end{aligned}$$

Da $S(\tau) = 0$ für $\tau < 0$ ist die Integration erst ab der unteren Integrationsgrenze $\tau = 0$ durchzuführen. Für $\tau > 0$ ist $S(\tau) = 1$:

$$y(t) = \int_0^{\infty} f(\tau) e^{-\alpha t} e^{\alpha \tau} S(t - \tau) d\tau .$$

Wir spalten das Integral auf in zwei Teilintegrale

$$y(t) = \int_0^t f(\tau) e^{-\alpha t} e^{\alpha \tau} S(t - \tau) d\tau + \int_t^{\infty} f(\tau) e^{-\alpha t} e^{\alpha \tau} S(t - \tau) d\tau .$$

Das zweite Integral verschwindet, da hier $\tau > t$ und $S(t - \tau) = 0$ für $\tau > t$. Im ersten Integral ist $0 < \tau < t$ und für diesen Bereich $S(t - \tau) = 1$:

$$\Rightarrow \boxed{y(t) = e^{-\alpha t} \int_0^t e^{\alpha \tau} f(\tau) d\tau .} \quad \square$$

Folgerung: $y(t)$ ist nach Beispiel 15.11 die Lösung der Differenzialgleichung

$$y'(t) + \alpha y(t) = f(t) \quad \text{mit} \quad y(0) = 0.$$

Obige Formel hatten wir auch über die Variation der Konstanten für eine lineare Differenzialgleichung mit konstanten Koeffizienten erhalten ($\rightarrow$ Kap. 12.1.3).

Bemerkungen:

- (1) Das Faltungsintegral zweier Funktionen $f_1(t)$ und $f_2(t)$ ist kommutativ:

$$f_1 * f_2 = f_2 * f_1.$$

Denn mit der Substitution $\xi(\tau) = (t - \tau)$ ($\hookrightarrow d\xi = -d\tau$) ist

$$\begin{aligned} (f_1 * f_2)(t) &= \int_{-\infty}^{\infty} f_1(\tau) f_2(t - \tau) d\tau \\ &= \int_{-\infty}^{\infty} f_1(t - \xi) f_2(\xi) d\xi = (f_2 * f_1)(t). \end{aligned}$$

- (2) Bei Integralen der Form

$$\int_{-\infty}^{\infty} g(t - \tau) S(\tau) d\tau = \underbrace{\int_{-\infty}^0 g(t - \tau) \underbrace{S(\tau)}_{=0} d\tau}_{=0} + \int_0^{\infty} g(t - \tau) \underbrace{S(\tau)}_{=1} d\tau$$

tritt an der oberen (unteren) Grenze des 1. (2.) Teilintegrals bei $\tau = 0$ der unstetige Ausdruck $S(0)$ auf. Auf das Ergebnis der Integration hat dieser Funktionswert **keinen** Einfluss, da die Fläche unter einer Funktion sich nicht ändert, wenn die Funktion an endlich vielen Stellen abgeändert wird.

Beispiel 15.13 (Geometrische Interpretation). Das Faltungsintegral ist formal zwar einfach aufzustellen, aber zunächst recht unanschaulich. Im Folgenden geben wir eine geometrische Interpretation am Beispiel der Funktionen $f_1(t) = S(t)$ und $f_2(t) = S(t)$ an:

$$(f_1 * f_2)(t) = \int_{-\infty}^{\infty} S(\tau) S(t - \tau) d\tau.$$

Zur Bestimmung des Integrals betrachten wir die Bildfolge in Abb. 15.8 (a) - (d). In (a) ist die Funktion $S(\tau)$ graphisch dargestellt. Die Sprungfunktion ist Null für $\tau < 0$ und Eins für $\tau > 0$. Die Funktion $S(-\tau)$ geht aus der Funktion $S(\tau)$ durch Spiegelung (= **Faltung**) an der y -Achse hervor (b): $S(-\tau) = 0$ für $\tau > 0$ und $S(-\tau) = 1$ für $\tau < 0$. $S(t - \tau)$ entsteht aus $S(-\tau)$, indem der Graph von $S(-\tau)$ um t nach rechts verschoben wird (c). Anschließend ist

das Produkt von $S(\tau)$ und $S(t - \tau)$ in (d) dargestellt: $S(t - \tau) \cdot S(\tau) = 0$ für $\tau < 0$ und für $\tau > t$. Für das Integral $\int_{-\infty}^{\infty} S(\tau) S(t - \tau) d\tau$ mit der Integrationsvariablen τ bleibt nur der Bereich zwischen $0 \leq \tau \leq t$ ungleich Null und hat den Integralwert t

$$\Rightarrow (f_1 * f_2)(t) = \int_{-\infty}^{\infty} S(\tau) S(t - \tau) d\tau = t S(t) .$$

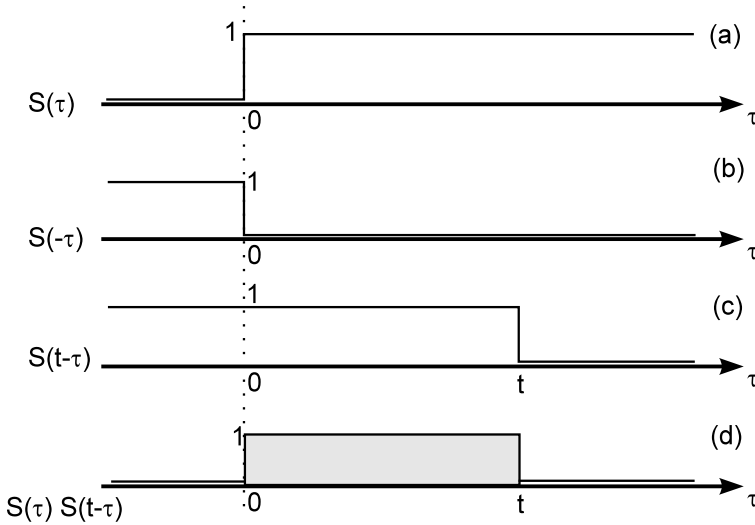


Abb. 15.8. Zur geometrischen Interpretation des Faltungsintegrals

Beispiel 15.14 (Faltungsintegral). Gesucht ist das Faltungsintegral $f * h$, wenn f und h die in Abb. 15.9 angegebenen Funktionen darstellen.

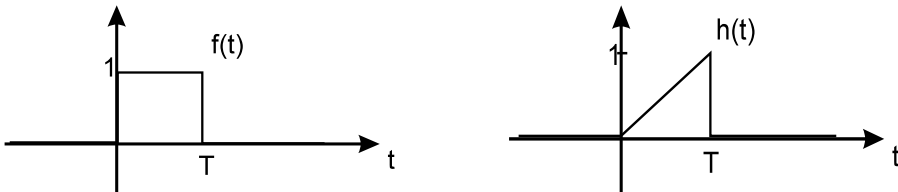


Abb. 15.9. Funktionen f und h

Wir bestimmen die Faltung $(f * h)(t) = \int_{-\infty}^{\infty} f(\tau) h(t - \tau) d\tau$ graphisch. Dazu gehen wir wie im obigen Beispiel 15.13 zunächst von der Funktion $h(\tau)$ durch Spiegelung zur Funktion $h(-\tau)$ über. Anschließend verschieben wir diese Funktion entlang der τ -Achse um den Wert T und multiplizieren dann mit der Rechteckfunktion. Diese vier Schritte sind schematisch in Abb. 15.10 gezeigt.

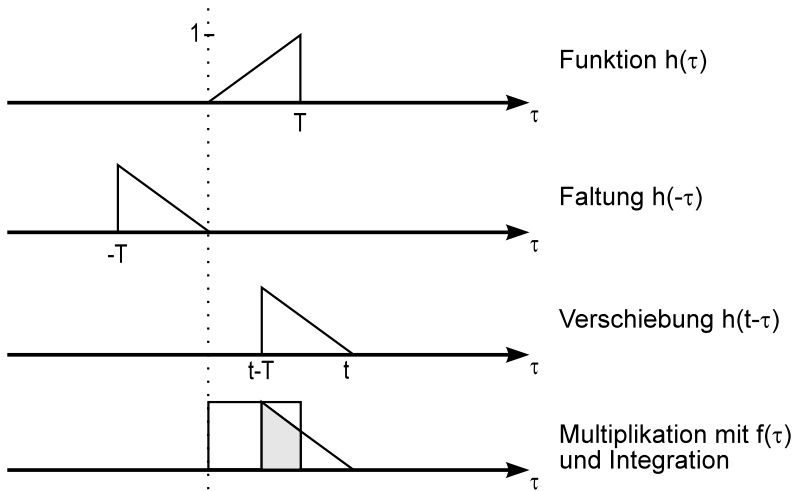


Abb. 15.10. Faltung der Rechteckfunktion mit der Dreiecksfunktion

Es treten bei der Bestimmung des Faltungsintegrals vier Fälle auf:

- (1) $t \leq 0$: Dann hat die Funktion $h(t - \tau)$ mit der Funktion $f(\tau)$ keinen Überlapp.
- (2) $0 \leq t \leq T$: Die Funktion $h(t - \tau)$ taucht mit der Spitze in den Graphen der Funktion $f(\tau)$ ein.
- (3) $T \leq t \leq 2T$: Die Funktion $h(t - \tau)$ tritt aus dem Graphen der Funktion $f(\tau)$ heraus.
- (4) $2T \leq t$: Die Funktion $h(t - \tau)$ hat mit der Funktion $f(\tau)$ keinen Überlapp.

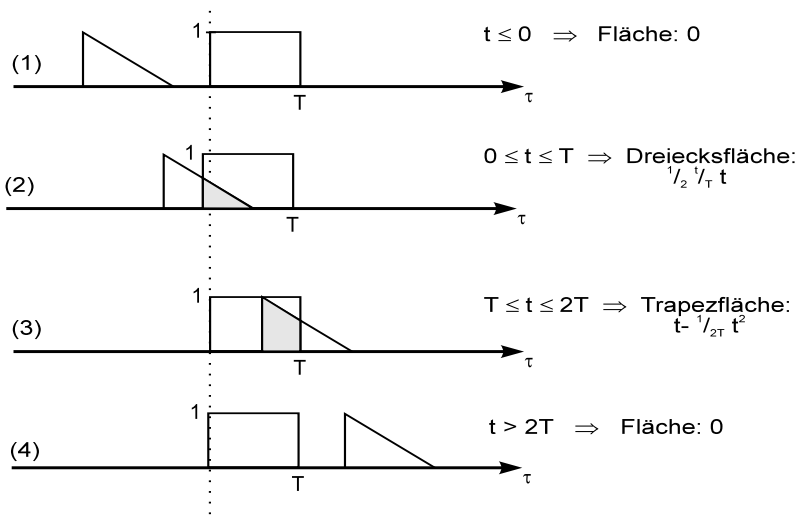


Abb. 15.11. Vier Fälle bei der Bestimmung des Faltungsintegrals

Das Ergebnis der Faltung lässt sich sowohl formelmäßig durch

$$\Rightarrow (f * h)(t) = \begin{cases} 0 & \text{für } t \leq 0 \\ \frac{1}{2T} t^2 & \text{für } 0 \leq t \leq T \\ t - \frac{1}{2T} t^2 & \text{für } T \leq t \leq 2T \\ 0 & \text{für } t > 2T \end{cases}$$

als auch graphisch darstellen:

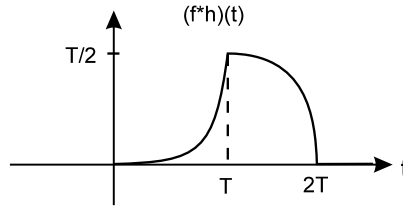


Abb. 15.12. Faltungsintegral $(f * h)(t)$

□

Beispiel 15.15 (Musterbeispiel: Lösen von Differenzialgleichungen mit der Fourier-Transformation).

Die Fourier-Transformation wird nicht nur zum Lösen von Differenzialgleichungen 1. Ordnung, sondern auch für lineare Differenzialgleichungen höherer Ordnung herangezogen: Gegeben ist die Differenzialgleichung 2. Ordnung

$$y''(t) - y(t) = \ln(t+1) S(t) .$$

1. Schritt: Durch Anwenden der Fourier-Transformation und Ausnutzung des Ableitungssatzes (F_8)

$$\mathcal{F}(y''(t))(\omega) = (i\omega)^2 \mathcal{F}(y(t))(\omega)$$

erhält man

$$(i\omega)^2 \mathcal{F}(y(t)) - \mathcal{F}(y(t)) = \mathcal{F}(\ln(t+1) S(t)) .$$

2. Schritt: Auflösen nach der Fourier-Transformierten: Die algebraische Gleichung für die Fourier-Transformierte $\mathcal{F}(y)$ wird nach $\mathcal{F}(y)$ aufgelöst:

$$\begin{aligned} \hookrightarrow \mathcal{F}(y(t)) &= \frac{1}{-1 - \omega^2} \mathcal{F}(\ln(t+1) S(t)) \\ &= -\frac{1}{2} \frac{2}{1 + \omega^2} \mathcal{F}(\ln(t+1) S(t)) \\ &= -\frac{1}{2} \mathcal{F}\left(e^{-1 \cdot |t|}\right) \cdot \mathcal{F}(\ln(t+1) S(t)) , \end{aligned}$$

da nach Beispiel 15.4: $\mathcal{F}(e^{-\alpha|t|}) = \frac{2\alpha}{\alpha^2 + \omega^2}$.

3. Schritt: Rücktransformation: Mit dem Faltungstheorem erhalten wir die Lösung der inhomogenen Differenzialgleichung 2. Ordnung

$$\begin{aligned} y(t) &= -\frac{1}{2} (e^{-|t|}) * (\ln(t+1) S(t)) \\ &= -\frac{1}{2} \int_{-\infty}^{\infty} e^{-|\tau|} \ln(t-\tau+1) S(t-\tau) d\tau \\ &= -\frac{1}{2} \int_{-\infty}^t e^{-|\tau|} \ln(t-\tau+1) d\tau. \end{aligned}$$



Hinweis: Weitere Beispiele zur Anwendung der Fourier-Transformation beim Lösen von Differenzialgleichungen findet man auf der CD-Rom in 15.4 und 15.6 sowie in den zugehörigen [MAPLE-Worksheets](#).

Zusammenfassung: Eigenschaften der Fourier-Transformation

$F(\omega)$ bezeichne die Fourier-Transformierte von f , $F_1(\omega)$ und $F_2(\omega)$ die Transformaten von f_1 und f_2 .

(F₁) **Linearität:** $\mathcal{F}(k_1 f_1 + k_2 f_2)(\omega) = k_1 F_1(\omega) + k_2 F_2(\omega)$

(F₂) **Symmetrie:** $\mathcal{F}(\mathcal{F}(f))(t) = 2\pi f(-t)$

(F₃) **Skalierung:** $\mathcal{F}(f(at))(\omega) = \frac{1}{|a|} F\left(\frac{\omega}{a}\right) \quad a \in \mathbb{R}_{\neq 0}$

(F₄) **Zeitverschiebung:** $\mathcal{F}(f(t-t_0))(\omega) = e^{-i t_0 \omega} F(\omega)$

(F₅) **Frequenzversch.:** $\mathcal{F}(e^{i\omega_0 t} f(t))(\omega) = F(\omega - \omega_0)$

(F₆) **Modulation:** $\mathcal{F}(f(t) \cos(\omega_0 t))(\omega) = \frac{1}{2} (F(\omega + \omega_0) + F(\omega - \omega_0))$

(F₇) **Ableitung:** $\mathcal{F}(f')(\omega) = i\omega F(\omega)$

(F₈) **n-te Ableitung:** $\mathcal{F}(f^{(n)})(\omega) = (i\omega)^n F(\omega)$

(F₉) **Faltungstheorem:** $\mathcal{F}(f_1 * f_2)(\omega) = F_1(\omega) \cdot F_2(\omega),$
mit $(f_1 * f_2)(t) = \int_{-\infty}^{\infty} f_1(\tau) f_2(t-\tau) d\tau$

15.3 Fourier-Transformation der Deltafunktion

Wenn man Systeme bezüglich ihrem Frequenzverhalten analysiert, benötigt man eine Funktion $\delta(t)$, welche alle Frequenzen mit derselben Amplitude enthält. Mit dieser Funktion als Eingangssignal regt man das System mit allen Frequenzen gleichermaßen an. (Mit einem beliebigen Eingangssignal würde man andernfalls das System für verschiedene Frequenzen unterschiedlich anregen.) Führt man anschließend eine Frequenzanalyse des Ausgangssignals durch, ist diese Information charakterisierend für das Frequenzverhalten des Systems.

Gesucht ist daher eine Funktion, die alle Frequenzen mit gleicher Amplitude enthält:

$$\mathcal{F}(\delta)(\omega) \equiv 1.$$

Diese Forderung führt auf die *Dirac-* oder *Deltafunktion*.

► 15.3.1 Deltafunktion und Darstellung der Deltafunktion

In der Systemtheorie und in vielen anderen Gebieten der Technik und Physik spielt die Impulsfunktion $\delta(t)$ eine sehr wichtige Rolle. Man bezeichnet diese Funktion auch oftmals nach ihrem Erfinder Dirac-Funktion oder auch aufgrund der Notation als Deltafunktion. In der Physik werden dieser Funktion angeblich die folgenden Eigenschaften zugewiesen:

Eigenschaften der Deltafunktion:

$$(1) \quad \delta(t) = 0 \quad \text{für } t \neq 0$$

$$(2) \quad \delta(0) = \infty$$

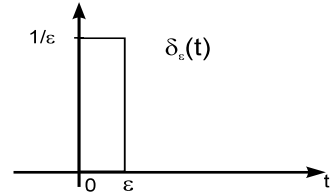
$$(3) \quad \int_{-\infty}^{\infty} \delta(t) dt = 1$$

$$(4) \quad \int_{-\infty}^{\infty} \delta(t) f(t) dt = f(0) \quad \text{für jede stetige Funktion } f : \mathbb{R} \rightarrow \mathbb{R}.$$

Freilich gibt es solche Funktionen im üblichen Sinne nicht und auf die Theorie der *Distributionen* (= Verallgemeinerten Funktionen) können wir uns hier nicht einlassen. Nur so viel: Die letzte Gleichung ist das Wesentliche, (3) ist der Spezialfall $f = 1$, Gleichung (1) und (2) haben wenig zu bedeuten!

Zur Erklärung gehen wir von dem folgenden Experiment aus. Auf einen frei beweglichen Körper wirke ein Kraftstoß $F(t) = \frac{m v_0}{\varepsilon}$ mit konstanter Stärke in der endlichen Zeitspanne ε . Definieren wir die Funktion

$$\delta_\varepsilon(t) = \begin{cases} 0 & \text{für } t < 0 \\ \frac{1}{\varepsilon} & \text{für } 0 < t < \varepsilon \\ 0 & \text{für } t > \varepsilon \end{cases}$$



lässt sich der Impulsstoß $F(t)$ schreiben als

$$F(t) = m v_0 \delta_\varepsilon(t - t_0) .$$

Der gesamte, übertragene Impuls ist dann

$$\Delta p = \int_{-\infty}^{\infty} F(t) dt = m v_0 \int_{-\infty}^{\infty} \delta_\varepsilon(t - t_0) dt = m v_0 \int_{t_0}^{t_0 + \varepsilon} \frac{1}{\varepsilon} dt = m v_0 .$$

Das Ergebnis ist unabhängig von der Zeitdauer ε ! Für $\varepsilon \rightarrow 0$ wird also derselbe Impuls übertragen als für ein endliches ε . Der Grenzwert der Funktionenfamilie $\delta_\varepsilon(t)$ für $\varepsilon \rightarrow 0$ ergibt die sogenannte **Deltafunktion (Diracfunktion; manchmal bezeichnet man sie auch nur mit Impulsfunktion)**.

$$\delta(t) := \lim_{\varepsilon \rightarrow 0} \delta_\varepsilon(t) .$$

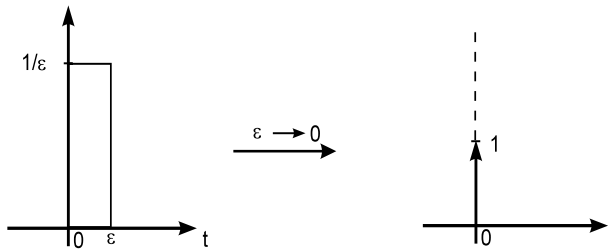


Abb. 15.13. Vom Rechteckimpuls zur Deltafunktion

Diese so definierte Funktion besitzt die Eigenschaften (1) - (4). Eigenschaft (1) - (3) sind offensichtlich erfüllt und Eigenschaft (4) prüft man folgendermaßen nach:

$$\int_{-\infty}^{\infty} \delta_\varepsilon(t) f(t) dt = \int_0^\varepsilon \frac{1}{\varepsilon} f(t) dt = f(\xi) \int_0^\varepsilon \frac{1}{\varepsilon} dt = f(\xi) ,$$

denn nach dem Mittelwertsatz der Integralrechnung darf f an einer geeigneten, aber unbekannten Zwischenstelle $\xi \in [0, \varepsilon]$ aus dem Integral gezogen werden (vgl. Abschnitt 8.2). Für $\varepsilon \rightarrow 0$ geht zum einen $\xi \rightarrow 0$, da $0 \leq \xi \leq \varepsilon$, und zum anderen $\delta_\varepsilon(t) \rightarrow \delta(t)$. Damit ist

$$\int_{-\infty}^{\infty} \delta(t) f(t) dt = \int_{-\infty}^{\infty} \lim_{\varepsilon \rightarrow 0} \delta_\varepsilon(t) f(t) dt = \lim_{\varepsilon \rightarrow 0} \int_{-\infty}^{\infty} \delta_\varepsilon(t) f(t) dt = f(0) .$$

Wichtig ist sich zu merken, dass die Deltafunktion ein *Funktional* ist, das durch die Integraleigenschaft

$$\int_{-\infty}^{\infty} f(t) \delta(t) dt = f(0)$$

charakterisiert wird, welche für jede stetige Funktion $f(t)$ gilt. Dies ist die *universelle Eigenschaft* der Deltafunktion. Allgemeiner gilt sogar die folgende Eigenschaft

$$\int_{-\infty}^{\infty} f(\tau) \delta(t - \tau) d\tau = f(t).$$

Denn

$$\int_{-\infty}^{\infty} f(\tau) \delta(t - \tau) d\tau = \int_{-\infty}^{\infty} f(t + \xi) \delta(\xi) d\xi = f(t + \xi)|_{\xi=0} = f(t).$$

Man nennt diese Beziehung die **Ausblendeigenschaft** der Deltafunktion, da von der Funktion f ein einzelner Wert nämlich der bei t "ausgeblendet" wird.

Bemerkung: Mit der Ausblendeigenschaft kann man auch zeigen, dass die Deltafunktion unabhängig von der gewählten Funktionenfamilie $\delta_\varepsilon(t)$ ist:

Denn sei $\delta_1(t)$ der Grenzwert einer Funktionenfamilie $\delta_{1\varepsilon}(t)$ und $\delta_2(t)$ der Grenzwert einer anderen Funktionenfamilie $\delta_{2\varepsilon}(t)$, dann gilt aufgrund der Ausblendeigenschaft von $\delta_1(t)$

$$\delta_2(t) = \int_{-\infty}^{\infty} \delta_1(t - \tau) \delta_2(\tau) d\tau = \int_{-\infty}^{\infty} \delta_1(\xi) \delta_2(t - \xi) d\xi = \delta_1(t) .$$

Die letzte Gleichheit gilt wegen der Ausblendeigenschaft von $\delta_2(t)$. Also ist $\delta_2(t) = \delta_1(t)$ für alle $t \in \mathbb{R}$. $\square$

► 15.3.2 Fourier-Transformation der Deltafunktion

Aufgrund der grundlegenden Eigenschaft der Deltafunktion, dass für jede stetige Funktion $\varphi(t)$

$$\int_{-\infty}^{\infty} \delta(t) \varphi(t) dt = \varphi(0) ,$$

gilt speziell für $\varphi(t) = e^{-i\omega t}$

$$\int_{-\infty}^{\infty} \delta(t) e^{-i\omega t} dt = e^{-i\omega t}|_{t=0} = 1 .$$

Die linke Seite der Gleichung ist die Fourier-Transformierte der Funktion $\delta(t)$; damit ist die Fourier-Transformierte der Deltafunktion die konstante Funktion

$$\mathcal{F}(\delta)(\omega) = 1.$$

Die Deltafunktion ist also genau die Funktion, die wir für die Systemanalyse benötigen: Sie enthält alle Frequenzen mit der Amplitude 1.

Bemerkungen:

- (1) Setzen wir dieses Ergebnis wiederum in die Umkehrformel der Fourier-Transformation (FI) ein, folgt

$$\delta(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \mathcal{F}(\delta)(\omega) e^{i\omega t} d\omega = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{i\omega t} d\omega. \quad (*)$$

Durch die Berechnung des uneigentlichen Integrals (*)

$$\begin{aligned} \delta(t) &= \lim_{\varepsilon \rightarrow 0} \frac{1}{2\pi} \int_{-1/\varepsilon}^{1/\varepsilon} e^{i\omega t} d\omega = \lim_{\varepsilon \rightarrow 0} \frac{1}{2\pi} \frac{1}{it} (e^{i\omega t}) \Big|_{\omega=-1/\varepsilon}^{\omega=1/\varepsilon} \\ &= \lim_{\varepsilon \rightarrow 0} \frac{1}{2\pi} \frac{1}{it} \left(e^{i\frac{1}{\varepsilon}t} - e^{-i\frac{1}{\varepsilon}t} \right) = \lim_{\varepsilon \rightarrow 0} \frac{\sin\left(\frac{1}{\varepsilon}t\right)}{\pi t} \end{aligned}$$

erhalten wir die Deltafunktion auch als Grenzwert der Funktionenfamilie $\frac{\sin(\frac{1}{\varepsilon}t)}{\pi t}$ für $\varepsilon \rightarrow 0$, welche wir schon in Beispiel 15.1 angegeben und graphisch diskutiert hatten.

- (2) Wenden wir auf Gleichung (*) nochmals die Fourier-Transformation an, erhält man mit der Symmetrieeigenschaft (F_2) die Fourier-Transformierte der konstanten Funktion:

$$\mathcal{F}(1)(\omega) = 2\pi \delta(\omega).$$

Beispiele 15.16.

- ① Wir berechnen die Fourier-Transformierte der Sprungfunktion

$$S(t) = \frac{1}{2} + \frac{1}{2} \operatorname{sign}(t).$$

Nach Beispiel 15.6 ist die Fourier-Transformierte von $f(t) = \frac{1}{t}$ gegeben durch

$$\mathcal{F}\left(\frac{1}{t}\right)(\omega) = -i\pi \operatorname{sign}(\omega).$$

Nach der Symmetrieeigenschaft (F_2) gilt daher

$$\mathcal{F}(\text{sign}(t))(\omega) = \frac{1}{-i\pi} \mathcal{F}\left(\mathcal{F}\left(\frac{1}{t}\right)\right)(\omega) = \frac{1}{-i\pi} 2\pi \left(\frac{1}{-\omega}\right) = \frac{2}{i\omega}$$

und wegen der Linearität (F_1)

$$\mathcal{F}\left(\frac{1}{2} + \frac{1}{2} \text{sign}(t)\right)(\omega) = \frac{1}{2} \mathcal{F}(1)(\omega) + \frac{1}{2} \mathcal{F}(\text{sign}(t))(\omega)$$

$$\Rightarrow \mathcal{F}(S(t))(\omega) = \pi \delta(\omega) + \frac{1}{i\omega}.$$

② Wegen den Verschiebungseigenschaften (F_4), (F_5) ist

$$\begin{aligned} \mathcal{F}(\delta(t - t_0))(\omega) &= e^{-i\omega t_0}, \\ \mathcal{F}(e^{i\omega_0 t})(\omega) &= 2\pi \delta(\omega - \omega_0). \end{aligned}$$

③ Wegen der Modulationseigenschaft (F_6) gilt mit $f(t) = 1$

$$\mathcal{F}(\cos(\omega_0 t))(\omega) = \pi \delta(\omega - \omega_0) + \pi \delta(\omega + \omega_0).$$

Dieses Ergebnis besagt, dass in $\cos(\omega_0 t)$ nur eine Frequenz ω_0 enthalten ist. Die Fourier-Transformation liefert als Spektrum von $\cos(\omega_0 t)$ nur eine Linie. Misst man $\cos(\omega_0 t)$ in einem endlichen Zeitintervall, so kommt es nach Beispiel 15.10 allerdings zur Verbreiterung dieser Linie!

④ Mit der Ausblendeigenschaft der Deltafunktion gilt

$$\delta(t - t_0) * f(t) = \int_{-\infty}^{\infty} \delta(\tau - t_0) f(t - \tau) d\tau = f(t - \tau)|_{\tau=t_0} = f(t - t_0).$$

Beispiel 15.17 (Fourier-Transformation periodischer Funktionen). Sei f eine T -periodische Funktion. Dann gilt nach dem Satz von Fourier für periodische Funktionen mit $\omega_0 = \frac{2\pi}{T}$ in der komplexen Schreibweise

$$f(t) = \sum_{n=-\infty}^{\infty} c_n e^{in\omega_0 t}.$$

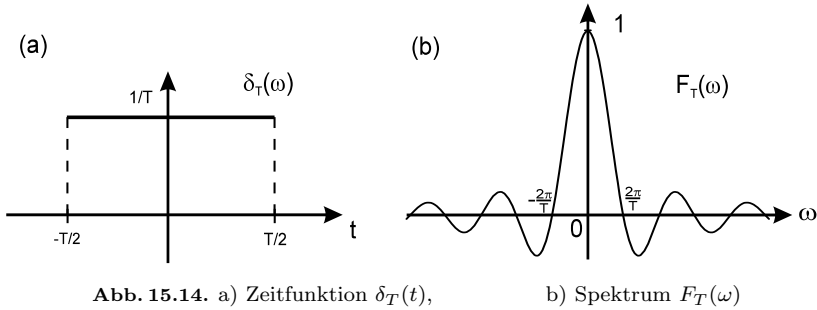
Aufgrund der Linearität der Fourier-Transformation gilt mit Beispiel 15.16 ②

$$\mathcal{F}(f(t))(\omega) = \sum_{n=-\infty}^{\infty} c_n \mathcal{F}(e^{in\omega_0 t})(\omega) = 2\pi \sum_{n=-\infty}^{\infty} c_n \delta(\omega - n\omega_0).$$

Die Fourier-Transformierte einer periodischen Funktion ist durch das Spektrum der Fourier-Reihe gegeben. Daher bezeichnet man als **Spektrum** eines beliebigen Signals $f(t)$ die **Fourier-Transformierte** $F(\omega)$. $\square$

➤ 15.3.3 Darstellung der Fourier-Transformierten von $\delta(t)$

Wir zeigen im Folgenden den Übergang des Spektrums der Funktionenfamilie $\delta_\varepsilon(t)$ für $\varepsilon \rightarrow 0$ zum Spektrum der Deltafunktion auf. Nach Beispiel 15.1 gilt die Beziehung



$$F_T(\omega) := \mathcal{F}\left(\frac{1}{T} \operatorname{rect}\left(\frac{2t}{T}\right)\right)(\omega) = \frac{\sin(\omega T/2)}{\omega T}.$$

Für $T \rightarrow 0$ geht die Zeitfunktion $\delta_T(t) = \frac{1}{T} \operatorname{rect}(\frac{2t}{T})$ gegen die Deltafunktion. Wir diskutieren das Spektrum $F_T(\omega)$ in Abhängigkeit des Parameters T : Die Maximalamplitude ist 1; unabhängig von T . Die erste Nullstelle des Spektrums liegt bei $\omega = \pm \frac{2\pi}{T}$; abhängig von T . Für $T \rightarrow 0$ strebt diese Nullstelle gegen $\pm\infty$:

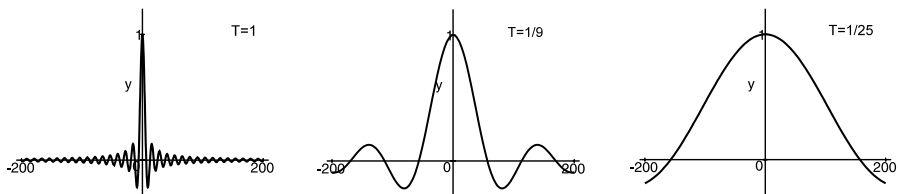
$$\delta_T(t) \xrightarrow{T \rightarrow 0} \delta(t)$$

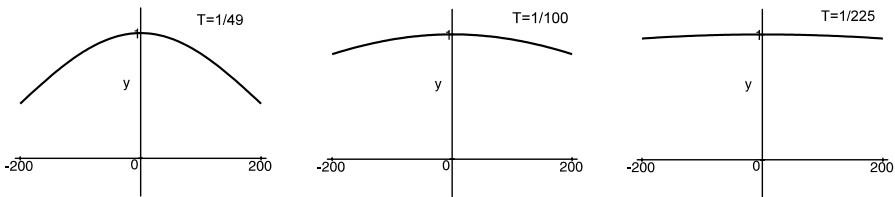
$$F_T(\omega) \xrightarrow{T \rightarrow 0} 1$$

Der Übergang $F_T(\omega) \rightarrow 1$ für $T \rightarrow 0$ lässt sich graphisch mit [MAPLE](#) sehr schön veranschaulichen. Wir definieren das Spektrum zu den Zeitfunktionen $F_T(t)$:

$$F_T := 2 \frac{\sin\left(\frac{1}{2} w T\right)}{w T}$$

Für kleiner werdendes T werden diese Funktionen in Form einer Animation über T graphisch dargestellt:





Animation: Man erkennt an den Einzelbildern, dass die Maximalamplitude stets bei 1 bleibt, die Nullstellen aber gegen $\pm\infty$ wandern, so dass als Grenzfunktion die konstante Funktion $F(\omega) = 1$ herauskommt. $\square$



Hinweis/Ausblick: Auf der CD-Rom sind die Anwendungen der Deltafunktion zur **Signal- und Systemtheorie** für lineare Systeme in Abschnitt 15.5 und 15.6 enthalten. Zentral ist dabei der Begriff der **Impulsantwort**, der die Reaktion eines Systems auf die Impulsfunktion (genauer auf die Deltafunktion) beschreibt. Die Fourier-Transformierte der Impulsantwort ist die **Systemfunktion** (= Übertragungsfunktion), welche das System bezüglich seinem Frequenzverhalten charakterisiert.

MAPLE-Worksheets zu Kapitel 15



Die folgenden elektronischen Arbeitsblätter stehen für Kapitel 15 mit MAPLE zur Verfügung.

- Von den Fourier-Reihen zur Fourier-Transformation
- Fourier-Transformation mit MAPLE
- Lösen von DG mit der Fourier-Transformation
- Deltafunktion mit MAPLE

Worksheets zu den Themen auf der CD-Rom:

- Systeme mit einem und zwei Energiespeicher mit MAPLE
- Systemanalyse des Doppelpendels über die Impulsantwort
- Systemanalyse eines Hochpasses über die Impulsantwort
- Diskrete Fourier-Transformation (DFT) mit MAPLE
- Diskrete Fourier-Transformation mit der FFT
- Signalanalyse diskreter Signale mit der FFT
- Systemanalyse eines Tiefpasses mit der FFT

15.3.4 Korrespondenzen der Fourier-Transformation

$f(t)$	$F(\omega) = \mathcal{F}(f)(\omega)$
$\delta(t)$	1
1	$2\pi \delta(\omega)$
$\cos(\omega_0 t)$	$\pi \delta(\omega - \omega_0) + \pi \delta(\omega + \omega_0)$
$\sin(\omega_0 t)$	$\frac{\pi}{i} \delta(\omega - \omega_0) - \frac{\pi}{i} \delta(\omega + \omega_0)$
$\text{sign}(t)$	$\frac{2}{i\omega}$
$S(t)$	$\pi \delta(\omega) + \frac{1}{i\omega}$
$S(t) \cos(\omega_0 t)$	$\frac{\pi}{2} \delta(\omega - \omega_0) + \frac{\pi}{2} \delta(\omega + \omega_0) + \frac{i\omega}{\omega_0^2 - \omega^2}$
$S(t) \sin(\omega_0 t)$	$\frac{\pi}{2i} \delta(\omega - \omega_0) - \frac{\pi}{2i} \delta(\omega + \omega_0) + \frac{\omega_0}{\omega_0^2 - \omega^2}$
$S(t) e^{-a t}$	$\frac{1}{a + i\omega} \quad (a > 0 \text{ bzw. } \text{Re } a > 0)$
$S(t) t^n \frac{e^{-a t}}{n!}$	$\frac{1}{(a + i\omega)^{n+1}} \quad (a > 0 \text{ bzw. } \text{Re } a > 0)$
$S(t) e^{-a t} \cos(\omega_0 t)$	$\frac{i\omega + a}{(i\omega + a)^2 + \omega_0^2} \quad (a > 0 \text{ bzw. } \text{Re } a > 0)$
$S(t) e^{-a t} \sin(\omega_0 t)$	$\frac{\omega_0}{(i\omega + a)^2 + \omega_0^2} \quad (a > 0 \text{ bzw. } \text{Re } a > 0)$
$e^{-a t }, \quad a > 0$	$\frac{2a}{a^2 + \omega^2}$
$e^{-a t } \cos(\omega_0 t), \quad a > 0$	$\frac{2a (\omega^2 + \omega_0^2 + a^2)}{(\omega^2 - \omega_0^2)^2 + a^2 (2\omega^2 + 2\omega_0^2 + a^2)}$
$e^{-a t^2}, \quad a > 0$	$\sqrt{\frac{\pi}{a}} e^{-\frac{\omega^2}{4a}}$
$\text{rect}\left(\frac{t}{T}\right) = \begin{cases} 1 & \text{für } t < T \\ 0 & \text{für } t > T \end{cases}$	$\frac{2 \sin(\omega T)}{\omega}$
$\Delta\left(\frac{t}{T}\right) = \begin{cases} 1 - \frac{ t }{T} & \text{für } t < T \\ 0 & \text{für } t > T \end{cases}$	$\frac{4 \sin^2\left(\frac{\omega T}{2}\right)}{T \omega^2}$

15.4 Aufgaben zur Fourier-Transformation

- 15.1 a) Bestimmen Sie die Fourier-Transformierte von

$$f_1(t) = \begin{cases} A & \text{für } -\frac{T}{2} < t < \frac{T}{2} \\ 0 & \text{sonst} \end{cases}$$

- b) Man diskutiere die Funktion $F(f_1)(\omega)$ für $A = \frac{1}{T}$

- c) Was ergibt sich für

$$f_2(t) = \begin{cases} A & \text{für } t_0 < t < t_0 + T \\ 0 & \text{sonst} \end{cases} ?$$

- 15.2 Man bestimme das Spektrum des Dreiecksignals

$$f(t) = \begin{cases} A \cdot (1 - |\frac{t}{T}|) & |t| \leq T \\ 0 & |t| > T \end{cases}$$

- 15.3 a) Berechnen Sie die Fourier-Transformierte der Funktion $e^{-\alpha|t|} \operatorname{sign}(t)$ (vgl. Abb. (a)).

- b) Man berechne die Fourier-Transformierte des $\cos^2$ -Impulses (vgl. (b)).

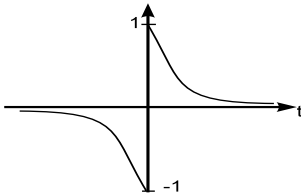


Abb. (a)

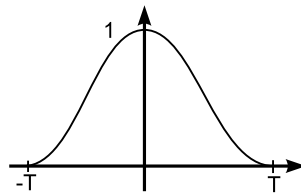


Abb. (b)

- 15.4 Man berechne mit MAPLE die Fourier-Transformierten von 15.1 - 15.3.

- 15.5 Geben Sie die Fourier-Transformierte der in Abb. (c) skizzierten Funktion $f(t)$ an.

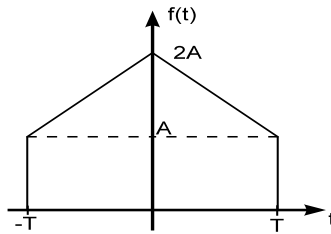


Abb. (c)

- 15.6 Zeigen Sie, dass die Fourier-Transformation eine lineare Transformation ist, d.h. das Superpositionsgesetz gültig ist:

$$\mathcal{F}(\alpha_1 f_1 + \alpha_2 f_2) = \alpha_1 \mathcal{F}(f_1) + \alpha_2 \mathcal{F}(f_2) .$$

- 15.7 Beweisen Sie

a) die Skalierungseigenschaft $\mathcal{F}(f(at))(\omega) = \frac{1}{|a|} \mathcal{F}(f(t))\left(\frac{\omega}{a}\right)$,

b) den Verschiebungssatz $\mathcal{F}(f(t - t_0))(\omega) = e^{-i\omega_0 t} \mathcal{F}(f(t))(\omega)$.

15.8 Zeigen Sie durch vollständige Induktion, dass

$$\mathcal{F}((-it)^n f) = \frac{d^n}{d\omega^n} \mathcal{F}(f)(\omega) = F^{(n)}(\omega),$$

wenn $F(\omega) = \mathcal{F}(f)(\omega)$.

15.9 Bestimmen Sie unter Verwendung der Eigenschaften der Fourier-Transformation die Transformaten von

a) $\delta(t)$ b) $\delta(t - t_0)$ c) $\frac{i}{2} (\delta(t + t_0) - \delta(t - t_0))$ d) $\sin(\omega_0 t)$

15.10 Man zeige, dass die folgenden Gleichungen gültig sind

a) $\mathcal{F}(e^{iat})(\omega) = 2\pi \delta(\omega - a)$

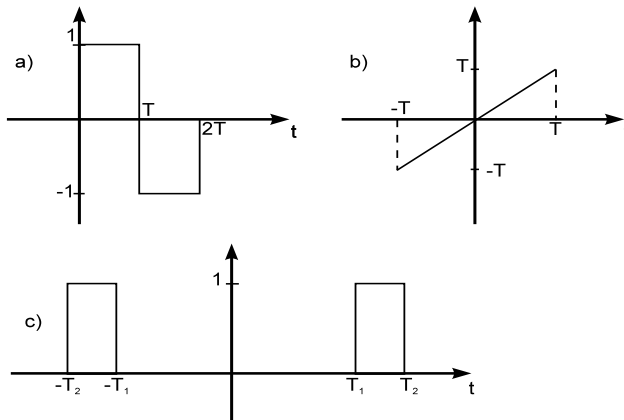
b) $\delta(t - t_0) * f(t) = f(t - t_0)$

15.11 Wie lautet die Faltung des Rechteckimpulses

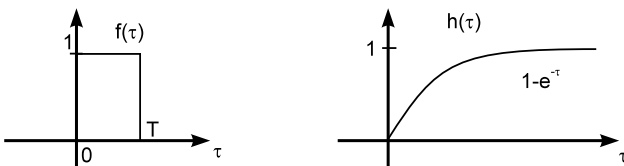
$$\text{rect}\left(\frac{2t}{T}\right) = \begin{cases} 1 & |t| \leq \frac{T}{2} \\ 0 & |t| > \frac{T}{2} \end{cases}$$

mit sich selbst? (Skizze!)

15.12 Berechnen Sie die Fourier-Transformaten der unten gezeichneten Funktionen



15.13 Bestimmen Sie mit Hilfe einer graphischen Skizze die Faltung $f * h$ der Funktionen für a) $t \leq 0$ b) $0 \leq t \leq T$ c) $T \leq t$.



15.14 Gegeben ist die Differenzialgleichung

$$y''(t) - 4y(t) = f(t).$$

Bestimmen Sie die Fourier-Transformierte von $y(t)$ sowie $y(t)$ in Form eines Faltungsintegrals.

Literaturverzeichnis

Das folgende Literaturverzeichnis enthält eine (keineswegs vollständige) Aufstellung von Lehrbüchern zur Ergänzung und Vertiefung der Ingenieurmathematik, Aufgabensammlungen, Handbücher sowie Literatur über MAPLE und über das Textverarbeitungssystem L^AT_EX.

Lehrbücher Ingenieurmathematik:

- Ayres, F.: Differential- und Integralrechnung. McGraw-Hill 1975.
- Brauch, W., Dreyer, H.J., Haacke, W.: Mathematik für Ingenieure. Teubner, Stuttgart 1990.
- Bronstein, I.N., Semendjajew, K.A.: Taschenbuch der Mathematik. Harri Deutsch, Thun/Frankfurt 1989.
- Burg, K., Haf, W., Wille, F.: Höhere Mathematik für Ingenieure I-IV. Teubner, Stuttgart 1985-90.
- Engeln-Müllges, G., Reutter, F.: Formelsamml. zur Numerischen Mathematik. BI Wissenschaftsverlag, Mannheim 1985.
- Fetzer, A., Fränkel, H.: Mathematik 1+2. Springer 1997+99.
- v. Finckenstein, K.: Grundkurs Mathematik für Ingenieure. Teubner, Stuttgart 1986.
- Fischer, G.: Lineare Algebra. Vieweg, Braunschweig 1986.
- Forster, O.: Analysis 1. Vieweg, Braunschweig 1983.
- Hainzel, J.: Mathematik für Naturwissenschaftler. Teubner, Stuttgart 1985.
- Hohloch, E., Kümmerer, H.: Brücken zur Mathematik 1-7, Cornelsen 1989-96.
- Meyberg, K., Vachenauer, P.: Höhere Mathematik 1+2. Springer 1999+97.
- Munz, C.D., Westermann, T.: Numerische Behandlung gewöhnlicher und partieller Differenzialgleichungen. Springer 2006.
- Papula, L.: Mathematik für Ingenieure 1+2. Vieweg, Braunschweig 1988.
- Spiegel, M.R.: Höhere Mathematik für Ingenieure und Naturwissenschaftler. McGraw-Hill 1978.
- Stingl, P.: Mathematik für Fachhochschulen. Carl Hanser 1992.
- Werner, W.: Mathematik lernen mit Maple (Band 1+2). dpunkt 1996+98.
- Westermann, T., Buhmann, W., Diemer, L., Endres, E., Laule, M., Wilke, G.: Mathematische Begriffe visualisiert mit MAPLE. Springer 2001.

Literatur zur Physik und Systemtheorie:

- Crawford, F.S.: Schwingungen und Wellen. Berkley Physik Kurs 3. Vieweg, Braunschweig 1979.
- Gerthsen, C., Vogel, H.: Physik. Springer 1993.
- Hering, E., Martin, R., Stohrer, M.: Physik für Ingenieure. Springer 1999.
- Mildenberger, O.: System- und Signaltheorie. Vieweg, Braunschweig 1989.
- Vielhauer, P.: Passive Lineare Netzwerke. Hüthig-Verlag 1974.

Literatur zu MAPLE:

- Burkhardt, W.: Erste Schritte mit Maple. Springer 1996.
- Char, B.W. et al.: Maple9 Learning Guide. Maple Inc. 2003.
- Devitt, J.S.: Calculus with Maple V. Brooks/Cole 1994.
- Dodson, C.T.J., Gonzalez, E.A.: Experiments In Mathematics Using Maple. Springer 1995.
- Ellis, W. et al.: Maple V Flight Manual. Brooks/Cole 1996.
- Heal, K.M. et. al.: Maple V: Learning Guide. Springer 1996.
- Heck, A.: Introduction to Maple. Springer 2003.
- Heinrich, E., Janetzko, H.D.: Das Maple Arbeitsbuch. Vieweg, Braunschweig 1995.
- Kofler, M. et al.: Maple: Einführung, Anwendung, Referenz. Addison-Wesley 2001.
- Komma, M.: Moderne Physik mit Maple. Int. Thomson Publishing 1996.
- Lopez, R.J.: Maple via Calculus. Birkhäuser, Boston 1994.
- Maple 11 Advanced Programming Guide. Maplesoft, Waterloo 2007.
- Maple 11 User Manual, Maplesoft. Waterloo 2007.
- Monagan, M.B. et al.: Maple9 Programming. Maple Inc. 2003.
- Westermann, T.: Mathematische Probleme lösen mit Maple. Springer 2008.

Literatur zu L^AT_EX:

- Dietsche, L., Lammarsch, J.: Latex zum Loslegen. Springer 1994.
- Kopka, H.: Latex. Addison-Wesley 1994.

Sachverzeichnis

Äquivalenzumformungen, 29

Ableitung, 245

der Umkehrfunktion, 255

elementarer Funktionen, 247

gemischte, 405

höhere, 248

numerische, 246

partielle, 401

Abstand, 18, 390

Ebene-Ebene, 73

Ebene-Gerade, 72

Gerade-Gerade, 65

Punkt-Ebene, 72

Punkt-Gerade, 65

Addition

komplexe, 200

Matrizen, 101

Vektoren, 43, 50

Additionstheoreme, 181, 380

Additivität des Integrals, 312

Amplitudenspektrum, 601, 618

Anordnung reeller Zahlen, 17

Anpassung

exponentielle, 444

logarithmische, 443

Potenz-, 444

Aperiodischer Grenzfall, 539

Arbeitsintegral, 333

Areafunktionen, 258

Arkusfunktionen, 182–185, 293

Assoziativgesetz, 14

Asymptoten, 164

Ausblendeigenschaft, 639

Balkenbiegung, 279

Basis, 87

Bernoullische Ungleichung, 18

Betrag, 18, 50, 618

eines Vektors, 43

komplexer, 195, 196

Betragsfunktion, 141

Bijektivität, 151

Bildfunktion, 563, 564

Bildungsgesetz bei Folgen, 223

Bildvektor, 109

Binomialkoeffizient, 11

Binomischer Lehrsatz, 12

Bogenmaß, 175

Charakteristisches Polynom, 511,
531–533, 535

Cramersche Regel, 123

Definitionsbereich, 140, 392

Definitionslücken, 163, 235

Deltafunktion, 637, 638

Ausblendeigenschaft, 639

Determinante, 114

Entwicklungssatz, 118

n-reihige, 118

zweireihige, 115

Dezimalzahlen, 431

Differenzenformeln, 246

Differenzenquotient, 245

Differenzial, 265

abhängiges, 265, 423

als lineare Näherung, 423

einer Funktion, 265

totales, 424, 425

unabhängiges, 265, 423

Differenzialgleichungen, 475

1. Ordnung, 476, 490

gewöhnliche, 475

homogene, 480, 525

inhomogene, 480, 482, 525

lineare 1. Ordnung, 480, 483

lineare DG Systeme, 501

n-ter Ordnung, 550

nichtlineare, 494

Ordnung der DG, 475

Differenzialquotient, 245

Differenzialrechnung, 244

Differenziation, 245

komplexwertiger Funktionen, 382

Differenziationsregeln

Faktorregel, 249

implizite, 260

Kettenregel, 252

logarithmische, 258

Potenzregel, 250

Produktregel, 250

Quotientenregel, 251

Summenregel, 249

- Differenzierbarkeit, 245
- Dimension, 89
- Diracfunktion, 638
- Diskretisierung, 498
- Diskriminante, 20
- Distributivgesetz
 - Matrizen, 105
 - Vektoren, 2D, 45
- divergent, 225, 345
 - bestimmt, 346
- Divergenz, 225
- Dividierte Differenzen, 159
- Division
 - komplexe, 203
- Doppelintegral, 453
 - Berechnung von, 455
- Dreifachintegral, 465
- Durchschnitt von Mengen, 4
- Dämpfung
 - starke, 540
- e, 227
- Ebenengleichung, 67
- Effektivwert, 335
- Eigenfrequenzen, 522
- Eigenraum, 512
- Eigenvektor, 508, 511
- Eigenwerte, 508
- Eineindeutigkeit, 151
- Einheitsvektor, 50
- Einschließungsalgorithmen, 238
- Einschwingvorgang, 489
- Einweggleichrichter, 602
- Elektrische Netzwerke, 26, 562
- Elektrischer Vierpol, 111
- Elektrisches Feld, 264
- Elektrostatisches Potenzial, 394
- Elemente einer Menge, 3
- Energieintegral, 333
- Entladekurve, 171
- Entwicklungspunkt, 355
- Entwicklungssatz nach Laplace, 118
- Erweiterung, stetige, 237
- Erzeugendensystem, 83
- Erzeugnis von Vektoren, 81
- Erzwungene Schwingung, 552
- Euler-Verfahren, 498, 499
- Eulersche Formel, 196, 378
- Eulersche Zahl, 227
- Exponentialform
 - komplexe, 196
- Exponentialfunktion, 142, 170, 282
 - allgemeine, 174
- Extremalwerte, relative, 271, 432
- Extremwertaufgaben, 277
- Fadenpendel, 132, 523
- Fakultät, 8
- Falk-Schema, 104
- Faltungsintegral, 630
- Faltungstheorem, 630
- Federpendel, 524
- Fehler
 - absoluter, 269, 428
 - relativer, 269, 429
- Fehlerfortpflanzung nach Gauß, 429
- Fehlerrechnung, 268, 428
- Flächenberechnung, 459
- Flächenmoment, 462
- Flächenberechnung, 329
- Folgen
 - Exponentialfolge, 227
 - Funktionsgrenzwerte, 229
 - Limesrechenregeln, 228
- Folnglieder, 223
- Fourier-Analyse, 587
- Fourier-Integral, 615
- Fourier-Koeffizienten, 591
 - komplexe, 606
- Fourier-Reihe, 586, 591, 609
 - 2w-periodische, 591
 - komplexe, 606
 - p-periodische, 599
- Fourier-Transformation, 613
 - der Ableitung, 629
 - der Deltafunktion, 639
 - der n-te Ableitung, 629
 - Faltungstheorem, 630
 - Frequenzverschiebung, 626
 - inverse, 619
 - Linearität, 623
 - Modulation, 627
 - Skalierung, 625
 - Symmetrie, 623
 - Zeitverschiebung, 626
- Fourier-Transformierte, 615

- Freie gedämpfte Schwingung, 537
- Freier Fall mit Luftwiderstand, 495
- Frequenzbereich, 615
- Frequenzverschiebung, 626
- Fundamentalsatz
 - der Algebra, 207
 - der Diff. u. Integralrechnung, 302
 - für LGS, 128
- Fundamentalsystem, 513, 533, 535
 - komplexes, 531
 - reelles, 532
- Funktionen, 140
 - Arkus-, 182
 - Betrags-, 141
 - diskrete, 224
 - einer Variablen, 140
 - Exponential-, 170
 - gebrochenrationale, 162
 - Integral-, 301, 304
 - komplexe Exponential-, 377
 - komplexwertige, 376
 - Kosinus-, 175
 - Kosinus-Hyperbolikus, 381
 - Kotangens-, 180
 - lineare, 427
 - Logarithmus-, 172
 - rationale, 162
 - reellwertige, 140
 - Sinus-, 175
 - Sinus-Hyperbolikus, 381
 - Stamm-, 305
 - Tangens-, 180
 - Umkehr-, 148
 - von n Variablen, 391
- Funktionenreihe, 355
- Funktionseigenschaften, 144
- Funktionsgrenzwert, 231

- Ganzrationale Funktion, 152
- Gauß-Algorithmus, 26, 30
- Gauß-Jordan-Verfahren, 106
- Gaußsche Zahlenebene, 194
- Gaußsches Eliminationsverfahren, 30
- Gebietsintegral, 453
 - dreidimensionales, 465
- gebrochenrational
 - echt, 321
 - unecht, 321
- Gebrochenrationale Funktionen, 162
- Gedämpfte Schwingung, 538
- Gekoppelte Pendel, 501
 - ohne Reibung, 518
- Geometrie
 - Ebene, 67
 - Gerade, 61
 - Hesse-Normalform, 68
 - Schnittpunkt Gerade-Ebene, 73
 - Schnittwinkel Gerade-Ebene, 73, 74
 - Schnittwinkel Geraden, 66
- Geometrische Summe, 10
- Gerade, 48
- Geradengleichung, 61
- Gestaffeltes System, 30
- Gewöhnliche Differenzialgleichungen, 475
- Gleichspannungsanteil, 601
- Gleichungen, 19
 - Ungleichungen, 23
- Gleichungssystem
 - homogenes, 29
 - inhomogenes, 29
 - lineares, 26, 28
- Gradient, 411
- Gradmaß, 175
- Graph, 140, 392
- Grenzwert, 225, 231, 233
 - linksseitiger, 231
 - rechtsseitiger, 231
- Grundschwingungen, 519

- Halbwertszeit, 173
- Harmonische Schwingungen, 209, 585
- Harmonisches Pendel, 267
- Hauptdiagonale, 100
- Hauptsatz der Differenzial- und Integralrechnung, 309
- Hesse-Normalform, 48, 68
- Hessesche Matrix, 420, 436
- Homogene DG
 - n -ter Ordnung, 529, 533, 535
- Homogene LDGS, 504
- Hooksches Gesetz, 263
- Horner-Schema, 155
- Hospitalsche Regeln, 284
- Hyperbelfunktionen, 257
- Häufungspunkt, 226

Höhenlinie, 393

Imaginäre Einheit, 193, 199

Imaginärteil, 194

Implizite Differenziation, 260

Impulsfunktion, 638

Induktion, vollständige, 6

Induktionsgesetz, 263

Inhomogene DG, 482, 545

 n-ter Ordnung, 529, 541

Injektivität, 151

Integral

 bestimmtes, 298, 309

 Riemann, 297

 unbestimmtes, 301

 uneigentliches, 327

Integralfunktion, 301, 304

Integration

 Integrationskonstante, 306

 komplexwertiger Funktionen, 383

 partielle, 313

Integrationsregeln

 Additivität, 312

 Faktorregel, 311

 Partialbruchzerlegung, 321

 partielle Integration, 313

 Substitutionsregel, 315

 Summenregel, 311

Interpolationspolynom

 Lagranges, 153

 Newtonsches, 159

Intervalle, 19

Inverse Matrix, 105, 121

Inverses Element, 14

Iteration, 238

Kartesisches Produkt, 5

Kern, 125

Kettenregel, 252

Kirchhoffsche Gesetze, 26, 111

Knotensatz, 26, 111

Koeffizienten bei LGS, 28

Koeffizientenmatrix, 29

Koeffizientenvergleich, 154

Kommutativgesetz, 14

Komplement von Mengen, 4

Komplexe Amplitude, 210

Komplexe Fourier-Reihe, 606

Komplexe Umformungen, 197

Komplexe Zahlen, 193

Komplexes Fundamentalsystem, 531

Konjugiert komplexe Zahl, 198

konvergent, 225, 345

 absolut, 345

Konvergenz, 225

Konvergenzbereich, 355

Konvergenzkriterien, 350

Konvergenzradius, 357, 377

Koordinatensystem

 kartesisches, 41

Korrespondenz, 564, 615

Kosinusfunktion, 175

Kosinushyperbolikus, 257, 293

Kosinustransformierte, 620

Kotangensfunktion, 180

Kotangenshyperbolikus, 258

Kräfteaddition, 43

Kräfteparallelogramm, 43

Kreuzprodukt, 55, 124

Kriechfall, 540

Krümmung

 Links-, 270

 Rechts-, 270

Kurvendiskussion, 274

Körper, 15

l'Hospitalsche Regeln, 284

Lagrange Interpolation, 153

Langzeitverhalten, 489

Laplace-Transformation, 561

 Additionssatz, 570

 Anwendungen, 577

 der Ableitung, 571

 der n-ten Ableitung, 572

 Eigenschaften, 569

 inverse, 568

 Linearität, 569

 Rücktransformation, 574

Laplace-Transformierte, 564, 566

Laplacescher Entwicklungssatz, 118

LDGS, 501

 homogene, 504

 Lösungen, 508

 zweiter Ordnung, 517

LGS, 28

Limes, 225

Limesrechenregeln, 228

- Linear unabhängige Funktionen, 505
- Lineare Abbildungen, 109
- Lineare Abhängigkeit, 84
- Lineare Differenzialgleichung, 475
- Lineare Differenzialgleichung 1.
 - Ordnung, 480, 483
- Lineare Differenzialgleichungssysteme, 501
- Lineare Gleichungssysteme
 - Lösbarkeit, 124
- Lineare Unabhängigkeit, 84, 129
- Linearfaktor, 155
- Linearisierung, 267, 419, 421
 - von Funktionen, 426, 427
- Linearität, 623
- Linearkombination, 81
- Logarithmische Differenziation, 258
- Logarithmus, 16
 - zur Basis b , 16
- Logarithmusfunktion, 172
- Lokale Extrema, 432
 - hinreichende Bedingung, 435
 - notwendige Bedingung, 433
- Lorentz-Kraft, 58
- Lösung
 - partikuläre, 484, 490, 491, 541, 542, 549
 - spezielle, 484
- Lösungen von LDGS, 508
- Lösungs-Fundamentalsystem, 506, 529

- Majorante, 350
- Majorantenkriterium, 350
- Maschensatz, 26, 111
- Matrix, 29
- Matrizelemente, 100
- Matrizen, 99
 - Addition, 101
 - Determinante, 115
 - Diagonalmatrix, 100
 - Einheitsmatrix, 100
 - Falk-Schema, 104
 - Gauß-Jordan-Verfahren, 106
 - Hauptdiagonale, 100
 - Inverse Matrix, 106
 - Multiplikation, 102
 - Nullmatrix, 101
 - Produkt, 103
 - quadratische, 100
 - Rang, 126
 - reguläre, 106
 - Sarrussche Regel, 120
 - symmetrische, 100
 - transponierte, 102
 - Umkehrmatrix, 106
- Maximum
 - relatives, 271, 432, 435
- Mengen, 3
- Mengenoperationen, 4
- Methode der kleinsten Quadrate, 440
- Minimum
 - relatives, 271, 432, 435
- Minorantenkriterium, 349
- Mittelwert
 - integraler, 303
 - linearer, 334
 - quadratischer, 335
- Mittelwerteigenschaft, 590, 619
- Mittelwertsatz, 283, 419
- Modulation, 627
- Moivresche Formel, 205
- Monotonie, 146
- Monotoniekriterium, 226, 227
- Monotonieverhalten, 270
- Multiplikation
 - komplexe, 201
 - Matrizen, 102

- Natürliche Zahlen, 5
- Newton-Verfahren, 159, 289, 290
- Nichtlineare Differenzialgleichungen, 494
- Normalform
 - algebraische, 195, 199
 - Exponentialform, 196, 199
 - trigonometrische, 196, 199
 - Umformungen, 197
- Nullfolge, 226
- Nullphase, 178
- Nullraum, 125
- Nullstellen, 144, 163
 - Polynome, 156
- Numerische Differenziation, 246
- Numerische Integration, 302
- Näherungsausdrücke, 421
- Näherungspolynome, 371

- Optimierungsprobleme, 277

- Ordnung der Differenzialgleichung, 475
- Ortsvektor, 42, 50
- p/q-Lösungsformel, 20
- Partialbruchzerlegung, 321
- Partialsumme, 344
- Partiell differenzierbar, 401
- Partielle Ableitung, 401
 - 1. Ordnung, 401
 - von f nach x , 404
 - zweiter Ordnung, 404
- Partielle Integration, 313
- Peanosche Axiome, 6
- Pendel, harmonisches, 267
- Periode, 177
- Periodizität, 147
- Permutation, 10
- Phase, 178, 618
- Phasenspektrum, 601, 618
- Phasenverschiebung, 179
- Plancksches Strahlungsgesetz, 287
- Plattenkondensator, 264
- Pole, 163
- Polygonzugverfahren, 499
- Polynomdivision, 157
- Polynome, 152
- Potenz, 15
 - komplexe, 205
- Potenzanpassung, 444
- Potenzfunktion, 167
 - allgemeine, 174
- Potenzreihe, 355
 - Eigenschaften, 360
 - geometrische, 356
 - komplexe, 376
- Potenzreihenentwicklung, 375
- Primzahlen, 9
- Prinzip der kleinsten Quadrate, 439
- Produktregel, 250
- Projektion eines Vektors, 53
- Quadratfunktion, 142
- Quadratische Gleichungen, 20
- Quotientenkriterium, 351
 - Limesform, 352
- Quotientenregel, 251
- Radioaktiver Zerfall, 171, 479
- Raketengleichung, 331
- Rang, 126
- Rationale Funktionen, 162
- RC-Kreis, 486, 492
- RCL-Kreis, 524
- RCL-Wechselstromkreis, 214, 277
- Realteil, 194
- Rechengesetze
 - für Vektorprodukt, 56
 - komplexe, 200
 - reeller Zahlen, 14
 - Vektoren, 77
 - Vektoren, 2D, 42
 - Vektoren, 3D, 50
- Rechteckimpuls
 - modulierter, 627
- Reduktion einer DG, 526
- Reelle Zahlen, 13
- Reelles Fundamentalsystem, 532
- Regeln
 - Substitutionsregel, 319
 - von l'Hospital, 284
- Regressionsgerade, 439, 441, 442
- Reihe, 345
 - alternierende, 353
 - alternierende harmonische, 354
 - arithmetische, 347
 - geometrische, 346
 - harmonische, 348, 349
 - MacLaurinsche, 365
 - Taylor-Reihe, 364
 - unendliche, 345
- rekursive Folge, 228
- Relative Extremwerte, 271, 432
 - notwendige Bedingung, 433
- relatives
 - Maximum, 271
 - Minimum, 271
- Restglied, 420
- Richtungsableitung, 414
- Richtungsvektor, 41, 50
- Riemann-Integral, 297
- RL-Kreis, 477, 481, 486
- Rohstoffkette, 110
- S-Multiplikation, 77
- Sarrus, 120
- Sattelpunkt, 272, 434, 435
- Satz von Fourier, 591, 599

- Satz von Laplace, 564
- Satz von Rolle, 283
- Satz von Schwarz, 406
- Satz von Steiner, 468
- Satz von Taylor, 418
- Scheinwerferregelung, 373
- Schnittkurvendiagramm, 397
- Schwarz, 406
- Schwebung, 519
- Schwerpunkt, 336, 467
 - ebene Fläche, 460
 - Koordinaten, 337
- Schwingungen, 209
 - harmonische, 585
- Schwingungsformen, 519, 522
- si-Funktion, 616
- Signalanalyse, 586
- Sinusfunktion, 142, 175
 - allgemeine, 177
- Sinushyperbolikus, 257, 293
- Sinustransformierte, 620
- Skalarprodukt, 51
 - 2D, 44
- Skalierung, 625
- Spaltenrang, 125
- Spaltenraum, 125
- Spaltenvektor, 99
- Spannungsintegral, 332
- Spatprodukt, 59
- Spektralbereich, 615
- Spektralfunktion, 615
- Spektrbreite, 628
- Spektrum
 - Amplitudenspektrum, 601, 618
 - diskretes, 601
 - Phasenspektrum, 601, 618
- Stammfunktion, 305
- Stationärer Punkt, 434
- Steinerscher Satz, 468
- stetig, 235
- stetige Erweiterung, 237
- Stetigkeit, 235, 398
 - Delta-Epsilon-Stetigkeit, 400
 - stückweise stetig, 563
- Streckenzugverfahren, 498
- Störfunktion, 480, 491
- Stückweise Stetigkeit, 590
- Substitutionsregel, 315
- Subtraktion
 - komplexe, 200
 - Vektoren, 2D, 43
 - Vektoren, 3D, 50
- Summe
 - unendliche Reihe, 345
- Superposition, 80, 210
- Surjektivität, 151
- Symmetrie, 145, 623
- Systeme
 - homogen, 503
 - inhomogen, 503
- T-periodische Signale, 599
- Tangensfunktion, 180
- Tangenshyperbolikus, 257, 293
- Tangentialebene, 408, 409, 421
- Taylor
 - Polynom, 363
 - Satz von, 363
 - Taylorsche Formel, 363
- Taylor-Reihen, 366 - 369
- Teilsommen, 344
- Totale Differenzierbarkeit, 408
- Totales Differenzial, 424, 425
- Trägheitsmoment, 467
- Trennung der Variablen, 480, 494
- Trigonometrische Funktionen, 175
- Überlagerung von Schwingungen, 209
- Umkehrfunktion, 148
- Umkehrmatrix, 105
- Ungleichungen, 23
- Untervektorraum, 80
- Variable
 - abhängige, 141
 - unabhängige, 141
- Variation der Konstanten, 482
- Vektoren, 41
 - Addition, 43, 50
 - Assoziativgesetz, 45
 - Betrag, 43, 50
 - Drehimpuls, 58
 - Drehmoment, 58
 - Einheitsvektor, 44, 50
 - Geraden-Darstellung, 48
 - Hesse-Normalform, 48
 - Komponenten, 42

- Koordinatensystem, 42
- Kräfteaddition, 43
- Kreuzprodukt, 55
- Länge, 43, 50
- Linearkombination, 44, 51
- Lorentz-Kraft, 58
- Multiplikation mit Skalar, 42, 50, 56
- Normalen-Einheitsvektor, 47
- Orthonormalsystem, 52
- Ortsvektor, 42, 50
- Projektion, 53, 54
- Rechtssystem, 59
- Richtungskosinus, 52
- Richtungsvektor, 42, 50
- Skalarprodukt, 44, 51
- Spatprodukt, 59
- Streckung, 42
- Vektorprodukt, 55
- Winkel, 46
- Vektoren, 2D, 42
- Vektoren, 3D, 50
- Vektoren, n-dimensional
 - äußere Verknüpfung, 77
 - Addition, 77
 - Basis, 87
 - Dimension, 89
 - Erzeugendensystem, 83
 - Erzeugnis, 82
 - innere Verknüpfung, 77
 - linear abhängig, 84
 - linear unabhängig, 84
 - Linearkombination, 81
 - Nullvektor, 77
 - S-Multiplikation, 77
 - Superposition, 80
 - Untervektorraum, 80
 - Vektorraum, 76, 78
- Vektorprodukt, 55, 124
- Vektorraum, 78
- Venn-Diagramm, 4
- Vereinigung von Mengen, 4
- Vollständige Induktion, 6
- Volumen, 463, 467
- Wachstum
 - höchstens exponentielles, 563
- Wechselspannung, 487
- Weg-Zeit-Gesetz, 139, 243, 524
- Wendepunkt, 272
- Wertebereich, 140
- Wheatstonesche Brückenschaltung, 269
- Widerstand
 - Blind-, 215
 - komplexer, 214
 - ohmscher, 214
 - reeller Schein-, 215
 - Wirk-, 215
- Widerstandsanpassung, 278
- Wiensche Verschiebungsgesetz, 288
- Winkelargument
 - komplexes, 196
- Winkelfunktionen, 175
- Wronski-Determinante, 529
- Wurzelfunktion, 142, 168
- Wurzelgleichungen, 21
- Wurzeln, 291
 - Einheitswurzel, 206
 - komplexe, 206
- Wurzelziehen
 - babylonisches, 228, 292
- Wärmestrahlung, 287
- Zahlen
 - komplex konjugierte, 198, 199
 - komplexe, 193
 - natürliche, 5
 - reelle, 13
- Zahlenebene
 - Gaußsche, 194, 199
- Zahlenfolge
 - reelle, 223
- Zahlengerade, 13
- Zeiger
 - komplexer, 195
- Zeilenrang, 125
- Zeilenumformungen
 - elementare, 29
- Zeilenvektor, 99
- Zeitfunktion, 563
- Zeitverschiebung, 626
- Zerfallsgesetz, 479
- Zielbereich, 140
- Zustandsgleichung, 390, 426
- Zwischensumme, 298