

Andreas Filler

Elementare Lineare Algebra

Linearisieren und Koordinatisieren



Mathematik
Primarstufe und
Sekundarstufe I+II

Elementare Lineare Algebra

Mathematik Primarstufe und Sekundarstufe I + II

Herausgegeben von
Prof. Dr. Friedhelm Padberg
Universität Bielefeld

Bisher erschienene Bände (Auswahl):

Didaktik der Mathematik

- P. Bardy: Mathematisch begabte Grundschulkinder - Diagnostik und Förderung (P)
M. Franke: Didaktik der Geometrie (P)
M. Franke/S. Ruwisch: Didaktik des Sachrechnens in der Grundschule (P)
K. Hasemann: Anfangsunterricht Mathematik (P)
K. Heckmann/F. Padberg: Unterrichtsentwürfe Mathematik Primarstufe (P)
G. Krauthausen/P. Scherer: Einführung in die Mathematikdidaktik (P)
G. Krummheuer/M. Fetzner: Der Alltag im Mathematikunterricht (P)
F. Padberg: Didaktik der Arithmetik (P)
P. Scherer/E. Moser Opitz: Fördern im Mathematikunterricht der Primarstufe (P)

G. Hinrichs: Modellierung im Mathematikunterricht (P/S)

R. Danckwerts/D. Vogel: Analysis verständlich unterrichten (S)
G. Greefrath: Didaktik des Sachrechnens in der Sekundarstufe (S)
F. Padberg: Didaktik der Bruchrechnung (S)
H.-J. Vollrath/H.-G. Weigand: Algebra in der Sekundarstufe (S)
H.-J. Vollrath: Grundlagen des Mathematikunterrichts in der Sekundarstufe (S)
H.-G. Weigand/T. Weth: Computer im Mathematikunterricht (S)
H.-G. Weigand et al.: Didaktik der Geometrie für die Sekundarstufe I (S)

Mathematik

- F. Padberg: Einführung in die Mathematik I – Arithmetik (P)
F. Padberg: Zahlentheorie und Arithmetik (P)

K. Appell/J. Appell: Mengen – Zahlen – Zahlbereiche (P/S)
S. Krauter: Erlebnis Elementargeometrie (P/S)
H. Kütting/M. Sauer: Elementare Stochastik (P/S)
T. Leuders: Erlebnis Arithmetik (P/S)
F. Padberg: Elementare Zahlentheorie (P/S)
F. Padberg/R. Danckwerts/M. Stein: Zahlbereiche (P/S)

A. Büchter/H.-W. Henn: Elementare Analysis (S)
G. Wittmann: Elementare Funktionen und ihre Anwendungen (S)

P: Schwerpunkt Primarstufe
S: Schwerpunkt Sekundarstufe

Weitere Bände in Vorbereitung

Andreas Filler

Elementare Lineare Algebra

Linearisieren und Koordinatisieren

Autor

Prof. Dr. Andreas Filler
Mathematisch-Naturwissenschaftliche Fakultät II
Humboldt-Universität zu Berlin
Unter den Linden 6
10099 Berlin

Internetseite zu diesem Buch: <http://www.afiller.de/linalg>

Wichtiger Hinweis für den Benutzer

Der Verlag, der Autor und der Herausgeber haben alle Sorgfalt walten lassen, um vollständige und akkurate Informationen in diesem Buch zu publizieren. Der Verlag übernimmt weder Garantie noch die juristische Verantwortung oder irgendeine Haftung für die Nutzung dieser Informationen, für deren Wirtschaftlichkeit oder fehlerfreie Funktion für einen bestimmten Zweck. Der Verlag übernimmt keine Gewähr dafür, dass die beschriebenen Verfahren, Programme usw. frei von Schutzrechten Dritter sind. Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen usw. in diesem Buch 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. Der Verlag hat sich bemüht, sämtliche Rechteinhaber von Abbildungen zu ermitteln. Sollte dem Verlag gegenüber dennoch der Nachweis der Rechtsinhaberschaft geführt werden, wird das branchenübliche Honorar gezahlt.

Bibliografische Information der Deutschen Nationalbibliothek

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

Springer ist ein Unternehmen von Springer Science+Business Media
springer.de

© Spektrum Akademischer Verlag Heidelberg 2011

Spektrum Akademischer Verlag ist ein Imprint von Springer

11 12 13 14 15 5 4 3 2 1

Das Werk einschließlich aller seiner Teile ist urheberrechtlich geschützt. Jede Verwertung außerhalb der engen Grenzen des Urheberrechtsgesetzes ist ohne Zustimmung des Verlages unzulässig und strafbar. Das gilt insbesondere für Vervielfältigungen, Übersetzungen, Mikroverfilmungen und die Einspeicherung und Verarbeitung in elektronischen Systemen.

Planung und Lektorat: Dr. Andreas Rüdinger, Martina Mechler
Herstellung und Satz: Crest Premedia Solutions (P) Ltd, Pune, Maharashtra, India
Umschlaggestaltung: SpieszDesign, Neu-Ulm

ISBN 978-3-8274-2412-9

Vorwort

Liebe Leserin, lieber Leser,

Linearisieren im Sinne der Beschreibung inner- und außermathematischer Sachverhalte durch lineare Strukturen und die Untersuchung dieser Strukturen sowie *Koordinatisieren* als Herstellung von Beziehungen zwischen der Arithmetik bzw. der Algebra und der Geometrie gehören zu den Grundprinzipien mathematischen Arbeitens. Diese beiden Prinzipien durchdringen als „rote Fäden“ das gesamte vorliegende Buch. Sie sind auch die Leitideen für das Studium der Linearen Algebra der von der Deutschen Mathematiker-Vereinigung (DMV), der Gesellschaft für Didaktik der Mathematik (GDM) und dem Verein zur Förderung des mathematischen und naturwissenschaftlichen Unterrichts (MNU) erarbeiten Standardempfehlungen für die Lehrerbildung im Fach Mathematik. Auf deren Grundlage wurde das vorliegende Buch mit der Zielstellung verfasst, ein besonders für die Lehrerbildung geeignetes Lehrwerk der linearen Algebra zu schaffen. In den erwähnten Standards von DMV, GDM und MNU heißt es:

Das Streben der Mathematik nach Abstraktion (um größtmögliche Anwendbarkeit zu erhalten) und nach Klassifikation (z. B. unter dem Aspekt der Dimension) zeigt sich deutlich in der Linearen Algebra.

Das vorliegende Buch behandelt die in den Standards für Lehrerinnen und Lehrer der Primarstufe und der Sekundarstufe I aufgeführten Inhalte umfassend und wirft noch einen Blick darüber hinaus. Ich empfehle es auch Studierenden für das gymnasiale Lehramt, die damit einen „sanften“ Übergang von der Analytischen Geometrie der Schule zu den Vorlesungen in Linearer Algebra vollziehen und gleichzeitig Anregungen für ihren späteren Unterricht erhalten möchten. Hinweise auf Literatur zu weiterführenden Themen werden gegeben.

Mitunter wird die Lineare Algebra als „trockenes“ Gebiet der Mathematik angesehen. Mit dem Studium dieses Buches werden Sie erleben, dass dem nicht so ist. Ausgehend von Bekanntem werden Sie schrittweise Verallgemeinerungen vornehmen und damit in neue „Dimensionen“ vordringen – und das sogar im direkten Sinne des Wortes. So führen Systematisierungen und Verallgemeinerungen bereits aus der Schulmathematik bekannter Objekte und Verfahren der elementaren Algebra und der analytischen Geometrie zu Vektor- und affinen Räumen, in denen dann Geometrie in vier und mehr Dimensionen betrieben werden kann. Sie werden dabei erfahren, dass die Mathematik nicht nur Bekanntes beschreiben kann, sondern ganz neue, ihr eigene „Räume eröffnet“.

Besonderes Augenmerk habe ich beim Schreiben des Buches folgenden Aspekten zukommen lassen:

- *Zentrale Begriffe und Strukturen werden ausgehend von Vertrautem, von konkreten Beispielen und Spezialfällen entwickelt.* So ist es für ein Lehrwerk der

Linearen Algebra sicherlich ungewöhnlich, dass der Begriff „Vektorraum“ erst in der zweiten Hälfte des Buches eingeführt wird. Jedoch bereiten zuvor bereits zahlreiche Beispiele und schrittweise Verallgemeinerungen auf das Verständnis dieses zentralen Strukturbegriffs vor.

- *Bezüge zur Schulmathematik* werden an unterschiedlichen Stellen des Buches hergestellt. Sie dienen im ersten Kapitel bei der Behandlung der linearen Gleichungssysteme als Einstieg in die Lineare Algebra und treten danach immer wieder auf, beispielsweise im Zusammenhang mit der Vektorrechnung und der analytischen Beschreibung geometrischer Abbildungen sowie in einer Reihe von Beispielen. Ein wichtiges Anliegen des Buches besteht darin, „Schulmathematik von einem höheren Standpunkt“ aus zu betrachten.
- Großer Wert wird auf die *Anschaulichkeit* der geführten Betrachtungen gelegt. Dazu dienen zahlreiche Illustrationen sowie zusätzliche Materialien auf der Internetseite zu diesem Buch. Bereits bei der Behandlung der linearen Gleichungssysteme im ersten Kapitel nehmen geometrische Interpretationen und Veranschaulichungen einen wichtigen Stellenwert ein. Damit werden Zusammenhänge zwischen Geometrie und Algebra aufgezeigt, denen auch eine Schlüsselstellung für das Verständnis der behandelten mathematischen Inhalte zukommt: „Verstehen durch Vorstellen“.
- Um Mathematik zu verstehen, ist *aktive Auseinandersetzung* mit Begriffen, Verfahren sowie Sätzen und ihren Beweisen erforderlich. Das Buch enthält daher eine Vielzahl von *Aufgaben* zum besseren Verständnis und zur Festigung des Erlernten. Mitunter werden Empfehlungen gegeben, zur Vorbereitung von Verallgemeinerungen geeignete Aufgaben zu lösen, die zu wichtigen Begriffen oder Sätzen hinführen. Auf der Internetseite des Buches stehen (zum Teil recht ausführliche) Lösungen der Aufgaben zur Verfügung. Ich empfehle Ihnen aber, diese erst dann anzuschauen, nachdem Sie sich selbst intensiv mit den Aufgaben auseinandergesetzt haben.

Berechnungen in der Linearen Algebra sind oft sehr aufwändig. Dennoch ist es empfehlenswert, zu rechnerischen Verfahren exemplarisch jeweils einige Aufgaben „per Hand“ zu lösen, denn das Vollziehen von Lösungsverfahren trägt auch zum Verständnis mathematischer Inhalte bei. Darüber hinaus ist die Nutzung des Computers sinnvoll. Das Buch enthält Anleitungen, die Sie befähigen, das freie Computeralgebrasystem Maxima für Berechnungen und Visualisierungen in der Linearen Algebra zu nutzen. Beispieldateien stehen auf der Internetseite

<http://www.afiller.de/linalg>

zur Verfügung. Weiterhin finden Sie dort Lösungen der im Buch gestellten Aufgaben, interaktive Illustrationen sowie weitere Materialien.

Ich wünsche Ihnen viel Freude und Erfolg beim Studium der Linearen Algebra.

Januar 2011, Andreas Filler

Inhaltsverzeichnis

Vorwort	v
1 Lineare Gleichungssysteme	1
1.1 Lineare Gleichungssysteme mit zwei Variablen	2
1.1.1 Lösungsmengen linearer Gleichungen mit zwei Variablen ...	2
1.1.2 Lösungsmengen von Gleichungssystemen mit zwei Variablen	5
1.1.3 Lösungsverfahren	8
1.1.4 Sachsituationen, die auf LGS mit zwei Variablen führen ...	12
1.1.5 Aufgaben zu Abschnitt 1.1	13
1.2 Lineare Gleichungssysteme mit drei Variablen	15
1.2.1 Lösungsmengen linearer Gleichungen und Gleichungssysteme mit drei Variablen	15
1.2.2 Der Gauß-Algorithmus	22
1.2.3 Aufgaben zu Abschnitt 1.2	28
1.3 Verallgemeinerungen und einige Begriffe	29
1.3.1 Der Gauß-Algorithmus in Matrixschreibweise	29
1.3.2 Ränge der einfachen und erweiterten Koeffizientenmatrix ...	32
1.3.3 Ein Lösbarkeitskriterium für lineare Gleichungssysteme	34
1.3.4 Homogene und inhomogene lineare Gleichungssysteme	37
1.3.5 Aufgaben zu Abschnitt 1.3	39
1.4 Einige Anwendungen linearer Gleichungssysteme	40
1.4.1 Aufgaben zu Abschnitt 1.4	43
1.5 Lösen linearer Gleichungssysteme mithilfe des Computers	44
2 Koordinatengeometrie	49
2.1 Geraden in der Ebene	50
2.1.1 Koordinatengleichungen von Geraden	50
2.1.2 Winkel zwischen Geraden in der Ebene	53
2.1.3 Aufgaben zu Abschnitt 2.1	54
2.2 Abstände von Punkten; Kreisgleichungen	55
2.2.1 Abstände von Punkten in der Ebene und im Raum	55
2.2.2 Kreise	58
2.2.3 Lagebeziehungen von Kreisen und Geraden	61
2.2.4 Aufgaben zu Abschnitt 2.2	65
2.3 Kugeln und Kegel	67
2.3.1 Kugelgleichungen	67
2.3.2 Kegel	69
2.3.3 Aufgaben zu Abschnitt 2.3	71
2.4 Kurven zweiter Ordnung: Ellipsen, Hyperbeln und Parabeln	72
2.4.1 Ortsdefinitionen von Ellipsen, Hyperbeln und Parabeln	72
2.4.2 Gleichungen von Ellipsen, Hyperbeln und Parabeln	76
2.4.3 Tangenten an Ellipsen, Hyperbeln und Parabeln	81

2.4.4	Aufgaben zu Abschnitt 2.4	82
3	Vektoren	85
3.1	Vektoren in physikalischen und geometrischen Kontexten	86
3.1.1	Kräfte, Geschwindigkeiten und Verschiebungen	86
3.1.2	Pfeilklassen	89
3.1.3	Addition und skalare Multiplikation von Pfeilklassen	91
3.1.4	Aufgaben zu Abschnitt 3.1	99
3.2	Vektoren in arithmetischen Kontexten	100
3.2.1	Stücklisten, Farben	100
3.2.2	n -Tupel	103
3.2.3	Aufgaben zu Abschnitt 3.2	105
3.3	Zusammenhänge zwischen Pfeilklassen und n -Tupeln – Vektoren	106
3.3.1	Pfeilklassen im Koordinatensystem	106
3.3.2	Die Isomorphie zwischen der Menge der Pfeilklassen der Ebene bzw. des Raumes und $\mathbb{R}^2$ bzw. $\mathbb{R}^3$	109
3.3.3	Vektoren	110
3.3.4	Aufgaben zu Abschnitt 3.3	112
3.4	Linearkombinationen von Vektoren	113
3.4.1	Linearkombinationen	113
3.4.2	Kollineare und komplanare Vektoren	116
3.4.3	Anwendungen von Linearkombinationen in der Geometrie	118
3.4.4	Aufgaben zu Abschnitt 3.4	121
3.5	Das Skalarprodukt zweier Vektoren	122
3.5.1	Arithmetische Einführung des Skalarprodukts	122
3.5.2	Geometrische Deutung des Skalarprodukts	124
3.5.3	Anwendungen des Skalarprodukts in Geometrie und Physik	130
3.5.4	Aufgaben zu Abschnitt 3.5	131
3.6	Vektor- und Spatprodukt	133
3.6.1	Das Vektorprodukt	133
3.6.2	Das Spatprodukt – Berechnung von Volumina	136
3.6.3	Aufgaben zu Abschnitt 3.6	137
3.7	Vektorrechnung und -darstellung mithilfe des Computers	137
4	Vektorielle Raumgeometrie	139
4.1	Parameterdarstellungen von Geraden und Kurven	140
4.1.1	Beschreibung von Geraden durch Parametergleichungen	140
4.1.2	Parameter- und Koordinatengleichungen bzw. LGS	142
4.1.3	Lagebeziehungen und Schnittpunkte von Geraden	144
4.1.4	Parameterdarstellungen als Funktionen, Beschreibung von Bewegungen	146
4.1.5	Exkurs: Parameterdarstellungen einiger Kurven	147
4.1.6	Aufgaben zu Abschnitt 4.1	149
4.2	Ebenen	151

4.2.1	Parameter- und Koordinatengleichungen von Ebenen	151
4.2.2	Lagebeziehungen und Schnittpunkte von Geraden und Ebenen sowie von Ebenen und Ebenen	152
4.2.3	Aufgaben zu Abschnitt 4.2	155
4.3	Metrische Geometrie von Geraden und Ebenen	156
4.3.1	Normalengleichungen von Geraden und Ebenen	156
4.3.2	Schnittwinkel zwischen Geraden und Ebenen	159
4.3.3	Die Hessesche Normalform von Geraden- und Ebenengleichungen; Abstandsberechnungen	161
4.3.4	Aufgaben zu Abschnitt 4.3	165
5	Vektorräume	167
5.1	Der Begriff des Vektorraumes, Beispiele	168
5.1.1	Definition des Begriffs Vektorraum	168
5.1.2	Beispiele für Vektorräume	169
5.1.3	Aufgaben zu Abschnitt 5.1	171
5.2	Untervektorräume	172
5.2.1	Definition, Unterraumkriterium und Beispiele	172
5.2.2	Der Durchschnitt und die Summe zweier Unterräume	176
5.2.3	Aufgaben zu Abschnitt 5.2	177
5.3	Lineare Hüllen, Erzeugendensysteme, lineare Abhängigkeit	178
5.3.1	Lineare Hüllen von Vektormengen	178
5.3.2	Erzeugendensysteme	179
5.3.3	Lineare Abhängigkeit und Unabhängigkeit	182
5.3.4	Aufgaben zu Abschnitt 5.3	186
5.4	Basis und Dimension	187
5.4.1	Der Begriff der Basis	187
5.4.2	Koordinaten von Vektoren bezüglich Basen	188
5.4.3	Weitere Beispiele für Basen	189
5.4.4	Sätze über Basen von Vektorräumen	191
5.4.5	Der Begriff der Dimension	197
5.4.6	Die Dimensionsformel für lineare Unterräume	199
5.4.7	Aufgaben zu Abschnitt 5.4	201
5.5	Affine Punkträume	203
5.5.1	Der Begriff des affinen Raumes	203
5.5.2	Koordinatensysteme	205
5.5.3	Affine Unterräume und ihre Lagebeziehungen	210
5.5.4	Anwendung auf die Theorie der linearen Gleichungssysteme	214
5.5.5	Aufgaben zu Abschnitt 5.5	218
5.6	Euklidische Vektor- und Punkträume	219
5.6.1	Positiv definite symmetrische Bilinearformen	219
5.6.2	Beträge von Vektoren, Orthogonalität und Winkel	222
5.6.3	Orthonormalbasen und kartesische Koordinatensysteme	224

5.6.4	Aufgaben zu Abschnitt 5.6	226
6	Matrizen	227
6.1	Begriffsbestimmung, Ränge von Matrizen	228
6.1.1	Zeilen- und Spaltenvektoren, transponierte Matrizen	228
6.1.2	Zeilenrang und Spaltenrang einer Matrix	229
6.1.3	Aufgaben zu Abschnitt 6.1	234
6.2	Matrizenmultiplikation und -inversion	235
6.2.1	Einführungsbeispiel: Materialverflechtung	235
6.2.2	Matrizenmultiplikation – Definition und Rechenregeln	238
6.2.3	Anwendungsbeispiele zur Matrizenmultiplikation	241
6.2.4	Inverse Matrizen	245
6.2.5	Aufgaben zu Abschnitt 6.2	249
6.3	Ein Ausblick auf Determinanten	251
6.3.1	Aufgaben zu Abschnitt 6.3	255
6.4	Matrizenrechnung mithilfe des Computers	256
7	Lineare und affine Abbildungen	257
7.1	Abbildungen: Definition und einige Beispiele	258
7.1.1	Der Begriff „Abbildung“, Eigenschaften von Abbildungen	258
7.1.2	Koordinatenbeschreibungen geometrischer Abbildungen	259
7.1.3	Aufgaben zu Abschnitt 7.1	263
7.2	Lineare Abbildungen	264
7.2.1	Definition und einige Beispiele	264
7.2.2	Einige Eigenschaften linearer Abbildungen	267
7.2.3	Matrizielle Darstellung linearer Abbildungen	270
7.2.4	Nacheinanderausführung linearer Abbildungen	273
7.2.5	Isomorphismen	274
7.2.6	Aufgaben zu Abschnitt 7.2	277
7.3	Affine Abbildungen	279
7.3.1	Definition und Beispiele	279
7.3.2	Einige Eigenschaften affiner Abbildungen	284
7.3.3	Aufgaben zu Abschnitt 7.3	286
7.4	Kongruenz- und Ähnlichkeitsabbildungen	287
7.4.1	Isometrien	287
7.4.2	Kongruenzabbildungen	292
7.4.3	Ähnlichkeitsabbildungen	294
7.4.4	Aufgaben zu Abschnitt 7.4	296
	Literaturverzeichnis	297
	Index	298

1 Lineare Gleichungssysteme

Übersicht

1.1	Lineare Gleichungssysteme mit zwei Variablen	2
1.2	Lineare Gleichungssysteme mit drei Variablen	15
1.3	Verallgemeinerungen und einige Begriffe	29
1.4	Einige Anwendungen linearer Gleichungssysteme	40
1.5	Lösen linearer Gleichungssysteme mithilfe des Computers	44

Lineare Gleichungssysteme bilden in mindestens dreifacher Hinsicht ein Kernstück der Linearen Algebra:

- *Lösungsverfahren* für lineare Gleichungssysteme sind für die Lösung vieler Standardaufgaben der Linearen Algebra von Bedeutung und führen – angewendet auf allgemeine lineare Gleichungssysteme ohne vorgegebene Zahlenwerte – zu Erkenntnissen über Lösungsmengen, welche Begründungen für essenzielle Aussagen der Linearen Algebra bereit stellen.
- Lösungsmengen linearer Gleichungssysteme bilden besonders wichtige Beispiele für *Strukturen*, mit denen sich die Lineare Algebra befasst.
- Lineare Gleichungssysteme sind für *Anwendungen* der Linearen Algebra bedeutsam.

Aus diesen Gründen bilden lineare Gleichungssysteme den Ausgangspunkt dieses Buches. Zunächst werden – direkt anknüpfend an Unterrichtsinhalte der Sekundarstufen I und II – Lösungsverfahren für lineare Gleichungssysteme behandelt, und es werden anhand von Beispielen, Plausibilitätsbetrachtungen sowie anschaulich begründeten Überlegungen Erkenntnisse über die Strukturen ihrer Lösungsmengen gewonnen. Damit wird auch der Boden für erste Verallgemeinerungen gegen Ende dieses Kapitels (in dem Abschnitt 1.3) bereitet; dort können dann auch erstmals in diesem Buch einige Tatsachen exakt bewiesen werden.

Die Inhalte dieses Kapitels werden in fast allen weiteren Kapiteln des Buches aufgegriffen und in den Kapiteln 5 und 6 auf der Grundlage der bis dahin behandelten Grundlagen aus der Theorie der Vektorräume zu allgemeinen Strukturbetrachtungen erweitert.

1.1 Lineare Gleichungssysteme mit zwei Variablen

1.1.1 Lösungsmengen linearer Gleichungen mit zwei Variablen

Lineare Gleichungen mit mehr als einer Variablen ergeben sich aus vielen – teilweise bereits recht einfachen – Anwendungskontexten.

Beispiel 1.1

Auf einem Mobiltelefon befindet sich ein Prepaid-Guthaben von 20 €. Eine SMS kostet 0,15 €, eine Minute Telefonieren 0,20 €. Wie viele Minuten lang kann mit dem Guthaben telefoniert und wie viele SMS können verschickt werden?

Die gegebene Anwendungssituation lässt sich durch die Gleichung

$$0,15 \cdot x + 0,2 \cdot y = 20 \quad (1.1)$$

beschreiben, wobei x die Anzahl der SMS und y die Anzahl der Gesprächsminuten angibt. Lösungen dieser Gleichung sind keine einzelnen Zahlen, sondern Zahlenpaare $(x; y)$, zum Beispiel $(0; 100)$, $(4; 97)$, $(8; 94)$, ... Diese Zahlenpaare lassen sich in einem Koordinatensystem darstellen, siehe Abb. 1.1. ■

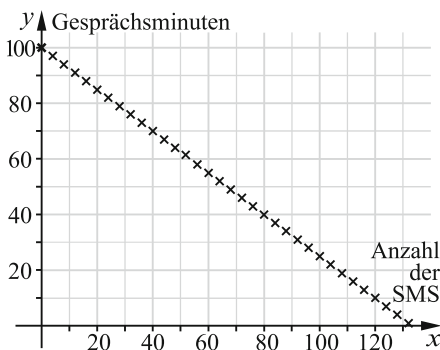


Abb. 1.1: Einige Lösungen einer linearen Gleichung mit zwei Variablen

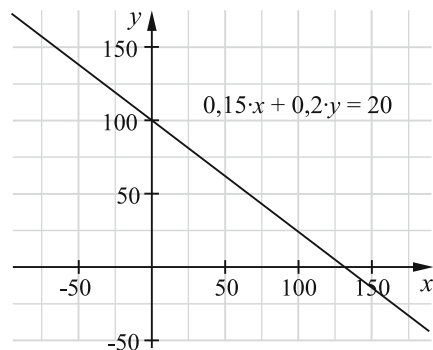


Abb. 1.2: Lösungsmenge einer linearen Gleichung mit zwei Variablen

Es zeigt sich, dass die den Lösungen zugeordneten Punkte auf einer Strecke liegen. Aufgrund der Anwendungssituation in Beispiel 1.1 sind nur Paare natürlicher Zahlen sinnvolle Lösungen. Betrachtet man jedoch – unabhängig von dem in dem Beispiel behandelten Kontext – alle Lösungen der Gleichung (1.1), d. h. alle Paare reeller Zahlen, welche die Gleichung erfüllen, so wird diese Lösungsmenge durch eine Gerade dargestellt, siehe Abb. 1.2. Diese Gerade ist zugleich Graph der linearen Funktion, deren Funktionsterm sich ergibt, wenn die Gleichung (1.1) nach y umgestellt wird, also der Funktion f mit $f(x) = -\frac{3}{4} \cdot x + 100$.

Allgemein können lineare Gleichungen mit zwei Variablen in der Form

$$a \cdot x + b \cdot y = c \quad \text{mit } a, b, c \in \mathbb{R} \text{ und } x, y \in \mathbb{R} \quad (1.2)$$

geschrieben werden. Eine Umstellung nach y und somit eine Darstellung der Lösungsmenge als Graph einer linearen Funktion in x ist nur möglich, falls $b \neq 0$ ist. Für $b=0$ und $a \neq 0$ nimmt Gleichung (1.2) die Gestalt $a \cdot x = c$ bzw. $x = \frac{c}{a}$ an. Die den Lösungen der Gleichung zugeordneten Punkte liegen somit ebenfalls auf einer Geraden, die allerdings parallel zur y -Achse verläuft.

Falls in einer Gleichung der Form (1.2) sowohl $a = 0$ als auch $b = 0$ ist, so nimmt die Gleichung die Gestalt $0 = c$ an. Falls dabei $c = 0$ ist, so ist jedes Paar $(x; y)$ mit $x, y \in \mathbb{R}$ eine Lösung der Gleichung; für $c \neq 0$ existiert keine Lösung, da die Gleichung unabhängig von x und y eine falsche Aussage ist. Nimmt man diese beiden Sonderfälle von den Betrachtungen aus, so gilt für die Lösungsmengen linearer Gleichungen mit zwei Variablen:

- Lösungsmengen linearer Gleichungen der Form $ax + by = c$ (mit $a, b, c \in \mathbb{R}$, wobei nicht zugleich $a=0$ und $b=0$ ist) lassen sich durch Geraden darstellen.
- Von der zweidimensionalen Grundmenge aller Paare $(x; y)$ mit $x, y \in \mathbb{R}$ lässt eine derartige Gleichung eine eindimensionale Lösungsmenge übrig. Die durch die Gleichung gegebene Bedingung schränkt also die Dimension der Menge der grundsätzlich in Frage kommenden Lösungen um Eins ein.¹

Die Lösungsmenge der Gleichung (1.1) lässt sich durch Umstellen der Gleichung nach y in der Form

$$L = \left\{ (x; y) \mid x \in \mathbb{R}; y = -\frac{3}{4}x + 100 \right\}$$

schreiben. Die (allgemeine) Gleichung (1.2) kann nach y umgestellt werden, falls $b \neq 0$ ist; dann ergibt sich $y = \frac{-ax+c}{b}$ und die Lösungsmenge lässt sich durch

$$L = \left\{ (x; y) \mid x \in \mathbb{R}; y = \frac{-ax+c}{b} \right\} \quad (1.3)$$

darstellen. Hat in (1.2) allerdings b den Wert Null, so kann diese Darstellung der Lösungsmenge nicht verwendet werden; es ist in diesem Falle

$$L = \left\{ (x; y) \mid y \in \mathbb{R}; x = \frac{c}{a} \right\}. \quad (1.4)$$

Für eine einheitliche Darstellung der Lösungsmenge einer linearen Gleichung der Form $ax + by = c$ (mit $a, b, c \in \mathbb{R}$ und $a \neq 0$ oder² $b \neq 0$) wird ein *Parameter* t eingeführt, der beispielsweise mit einer der beiden Variablen übereinstimmen kann (aber nicht muss). So lässt sich (1.3) in der Form

$$L = \left\{ (x; y) \mid \begin{array}{l} x = t \\ y = -\frac{a}{b} \cdot t + \frac{c}{b} \end{array}; t \in \mathbb{R} \right\} \quad (1.5)$$

¹Der Begriff der Dimension wird hier zunächst rein anschaulich verwendet (Geraden als ein-, Ebenen als zweidimensionale Objekte, der Raum als dreidimensionales Objekt). Eine präzise Definition des Dimensionsbegriffs erfolgt in dem Kapitel 5.

²Das Wort *oder* wird in der Mathematik in dem Sinne gebraucht, dass auch beide Bedingungen erfüllt sein dürfen. Die Formulierung „ $a \neq 0$ oder $b \neq 0$ “ schließt also den Fall „ $a \neq 0$ und $b \neq 0$ “ mit ein.

und (1.4) in der Form

$$L = \left\{ (x; y) \mid \begin{array}{l} x = \frac{c}{a} \\ y = t \end{array} ; t \in \mathbb{R} \right\} \quad (1.6)$$

schreiben. Diese Schreibweise kann verallgemeinert und damit vereinheitlicht werden. Gibt man die Lösungsmenge einer linearen Gleichung mit zwei Variablen x, y durch

$$L = \left\{ (x; y) \mid \begin{array}{l} x = m_1 \cdot t + n_1 \\ y = m_2 \cdot t + n_2 \end{array} ; t \in \mathbb{R} \right\} \quad (1.7)$$

an, so sind damit beide Fälle erfasst. m_1, m_2, n_1 und n_2 sind dabei feste reelle Zahlen, die von den Koeffizienten der jeweiligen Gleichung abhängen; t durchläuft den gesamten Bereich der reellen Zahlen. Die Lösungsmengen (1.5) und (1.6) sind Spezialfälle von (1.7) mit $m_1 = 1, m_2 = -\frac{a}{b}, n_1 = 0, n_2 = \frac{c}{b}$ für (1.5) sowie $m_1 = 0, m_2 = 1, n_1 = \frac{c}{a}, n_2 = 0$ für (1.6).

Für die Darstellung von Lösungsmengen wird oft die übersichtliche Spaltenschreibweise verwendet:

$$\begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} m_1 \\ m_2 \end{pmatrix} + \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}. \quad (1.8)$$

Diese Schreibweise gibt auch eine Beschreibung der durch die Lösungsmenge beschriebenen Geraden. Durch $\begin{pmatrix} m_1 \\ m_2 \end{pmatrix}$ ist ein Richtungsvektor dieser Geraden gegeben, $(n_1; n_2)$ sind die Koordinaten eines festen Punktes der Lösungsgeraden.

Auf den Begriff des Vektors wird in den Kapiteln 3 und 5 ausführlich eingegangen. Hier wird dieser Begriff zunächst in einem anschaulichen, aus der Schule bekannten, Sinne verwendet.

Beispiel 1.2

- Die lineare Gleichung

$$7x - 4y = 5$$

hat die Lösungsmenge

$$L_1 = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} \mid \begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} 1 \\ \frac{7}{4} \end{pmatrix} + \begin{pmatrix} 0 \\ -\frac{5}{4} \end{pmatrix} ; t \in \mathbb{R} \right\}.$$

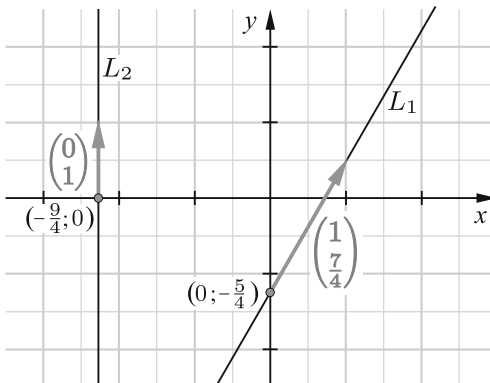


Abb. 1.3: Graphische Darstellungen der Lösungsmengen der linearen Gleichungen aus dem Beispiel 1.2.

- Die lineare Gleichung

$$4x = -9$$

hat die Lösungsmenge

$$L_2 = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} \mid \begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} + \begin{pmatrix} -\frac{9}{4} \\ 0 \end{pmatrix}; t \in \mathbb{R} \right\}.$$

Die Lösungsmengen beider Gleichungen sind mit den zugehörigen Richtungsvektoren $\begin{pmatrix} m_1 \\ m_2 \end{pmatrix}$ in Abb. 1.3 dargestellt. ■

Zusammenfassung:

Lösungsmengen linearer Gleichungen mit zwei Variablen

Lineare Gleichungen der Form

$$ax + by = c$$

mit $a, b, c \in \mathbb{R}$ und $a \neq 0$ oder $b \neq 0$ besitzen einparametrische Lösungsmengen, welche sich geometrisch als Geraden interpretieren lassen. Die Lösungsmengen können in der Form

$$L = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} \mid \begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} m_1 \\ m_2 \end{pmatrix} + \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}; t \in \mathbb{R} \right\}$$

dargestellt werden. Dabei sind m_1, m_2, n_1 und n_2 feste reelle Zahlen, die von den Koeffizienten der jeweiligen Gleichung abhängen; der Parameter t durchläuft den gesamten Bereich der reellen Zahlen.

1.1.2 Lösungsmengen linearer Gleichungssysteme mit zwei Gleichungen und zwei Variablen

Wie in dem vorangegangenen Abschnitt festgestellt wurde, lassen sich Lösungsmengen linearer Gleichungen geometrisch als Geraden darstellen, wobei jede Lösung jeweils durch einen Punkt repräsentiert wird. Lösungen $(x; y)$ eines linearen Gleichungssystems mit zwei Variablen

$$\begin{aligned} a_{11}x + a_{12}y &= b_1 \\ a_{21}x + a_{22}y &= b_2 \end{aligned} \tag{1.9}$$

müssen jede der beiden Gleichungen erfüllen, die entsprechenden Punkte müssen also beiden durch die Gleichungen beschriebenen Geraden angehören.

Beispiel 1.3

Die Lösungsmengen der beiden Gleichungen des Gleichungssystems

$$\begin{aligned} 4x - 5y &= 13 \\ 3x + 4y &= 3 \end{aligned}$$

werden durch zwei Geraden repräsentiert, siehe Abb. 1.4. Lösungen des gesamten Systems müssen beiden Lösungsmengen angehören. Das Gleichungssystem hat somit genau eine Lösung, sie entspricht dem Schnittpunkt der beiden Geraden. Durch die Ermittlung dieses Schnittpunktes ist es möglich, ein lineares Gleichungssystem näherungsweise zu lösen. ■

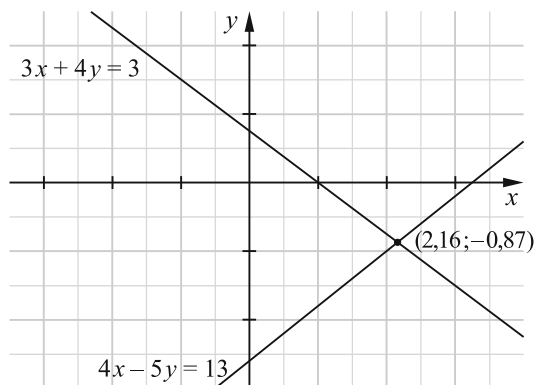


Abb. 1.4: Graphische Lösung eines linearen Gleichungssystems mit 2 Gleichungen und 2 Variablen

In den meisten Fällen haben lineare Gleichungssysteme mit zwei Gleichungen und zwei Variablen wie in Beispiel 1.3 genau eine Lösung. Jedoch können auch Fälle auftreten, in denen derartige Gleichungssysteme unlösbar sind oder unendlich viele Lösungen besitzen.

Beispiel 1.4

Das lineare Gleichungssystem (LGS)

$$\begin{aligned} x + y &= -1 \\ -3x - 3y &= 9 \end{aligned}$$

besitzt *keine Lösung*. Die beiden Gleichungen beschreiben zueinander parallele Geraden, siehe Abb. 1.5. Es existiert also kein Schnittpunkt der beiden Geraden und somit kein Paar $(x; y)$, welches beide Gleichungen erfüllt und somit Lösung des LGS sein könnte. ■

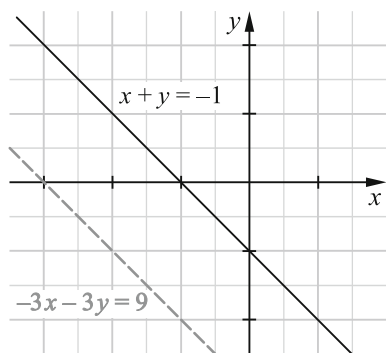


Abb. 1.5: LGS mit leerer Lösungsmenge

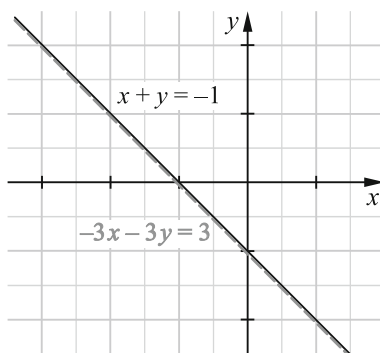


Abb. 1.6: LGS mit unendlich vielen Lösungen

Beispiel 1.5

Das lineare Gleichungssystem

$$\begin{aligned}x + y &= -1 \\ -3x - 3y &= 3\end{aligned}$$

besitzt *unendlich viele Lösungen*. Die beiden Gleichungen beschreiben jeweils dieselbe Gerade, siehe Abb. 1.6. Somit haben beide Gleichungen dieselbe unendliche (eiparametrische) Lösungsmenge, welche zugleich die Lösungsmenge des gesamten LGS ist. ■

Ein näherer Blick auf die in den Beispielen 1.4 und 1.5 betrachteten Gleichungssysteme offenbart die Ursachen für die Unlösbarkeit bzw. das Vorliegen einer unendlichen Lösungsmenge. Multipliziert man die erste Gleichung aus dem Beispiel 1.5 mit -3 , so ergibt sich gerade die zweite Gleichung. Somit sind beide Gleichungen äquivalent, besitzen also dieselbe Lösungsmenge, mit anderen Worten: die zweite Gleichung enthält keine zusätzliche Bedingung, die nicht schon in der ersten Gleichung enthalten ist (bzw. umgekehrt). Somit reduziert sich dieses LGS auf eine einzige Gleichung und hat deshalb eine unendliche (eiparametrische) Lösungsmenge, vgl. Abschnitt 1.1.1.

Im Falle des LGS aus Beispiel 1.4 fällt auf, dass die linke Seite der ersten Gleichung das -3 -fache der linken Seite der zweiten Gleichung ist. Gleichwertig zu dem in Beispiel 1.4 betrachteten LGS ist das Gleichungssystem

$$\begin{aligned}-3x - 3y &= 3 \\ -3x - 3y &= 9\end{aligned}$$

Wenn diese Gleichungssystem Lösungen $(x; y)$ hätte, so müsste für diese *zugleich* $-3x - 3y = 3$ und $-3x - 3y = 9$ sein. Dies ist jedoch nicht möglich, denn es hätte $3 = 9$ zur Folge. Somit kann dieses LGS keine Lösungen besitzen.

Zusammenfassung: Lösungsmengen linearer Gleichungssysteme mit zwei Variablen und zwei Gleichungen

Lineare Gleichungssysteme der Form

$$\begin{aligned}\text{I} \quad & a_{11}x + a_{12}y = b_1 \\ \text{II} \quad & a_{21}x + a_{22}y = b_2\end{aligned}$$

bei denen jede Gleichung für sich betrachtet eine eiparametrische Lösungsmenge besitzt, also $a_{11} \neq 0$ oder $a_{12} \neq 0$ sowie $a_{21} \neq 0$ oder $a_{22} \neq 0$ gilt, können folgendermaßen geartete Lösungsmengen besitzen:

- Es existiert genau eine Lösung. Dies ist der Fall, falls die linke Seite der Gleichung II kein Vielfaches der linken Seite der Gleichung I ist (d. h. dass die linke Seite von II nicht durch Multiplikation mit einer reellen Zahl aus der linken Seite von I entsteht). Die Lösungsmengen der beiden Gleichungen werden durch Geraden repräsentiert, die genau einen Schnittpunkt besitzen, vgl. Beispiel 1.3 und Abb. 1.4.

- Die Lösungsmenge des LGS ist leer, falls die linke Seite der Gleichung II ein Vielfaches der linken Seite der Gleichung I ist, die gesamte Gleichung II jedoch kein Vielfaches von Gleichung I ist. In diesem Falle können die Lösungsmengen der einzelnen Gleichungen durch zwei parallele Geraden dargestellt werden; eine gemeinsame Lösung existiert nicht, vgl. Beispiel 1.4 und Abb. 1.5.
- Es gibt unendlich viele Lösungen, falls die Gleichung II durch Multiplikation mit einer reellen Zahl aus Gleichung I entsteht. Die beiden Gleichungen sind in diesem Falle äquivalent; eine Gleichung ist verzichtbar, vgl. Beispiel 1.5 und Abb. 1.6.

1.1.3 Lösungsverfahren für lineare Gleichungssysteme mit zwei Gleichungen und zwei Variablen

Ein nahe liegendes Lösungsverfahren für LGS mit zwei Gleichungen und zwei Variablen ist das *Gleichsetzungsverfahren*. Hierbei werden beide Gleichungen nach derselben Variablen umgestellt und dann gleich gesetzt.

Beispiel 1.6

Gleichsetzungsverfahren

In dem linearen Gleichungssystem aus Beispiel 1.3

$$4x - 5y = 13$$

$$3x + 4y = 3$$

werden beide Gleichungen nach y umgestellt:

$$y = \frac{4}{5}x - \frac{13}{5}$$

$$y = -\frac{3}{4}x + \frac{3}{4}.$$

Durch Gleichsetzen ergibt sich daraus

$$\frac{4}{5}x - \frac{13}{5} = -\frac{3}{4}x + \frac{3}{4}$$

und nach Vereinfachung dieser Gleichung

$$31x = 67 \quad \text{bzw.} \quad x = \frac{67}{31}.$$

Setzt man diesen Wert für x in eine der nach y umgestellten Gleichungen ein, so ergibt sich der Wert für y :

$$y = \frac{4}{5} \cdot \frac{67}{31} - \frac{13}{5} = \frac{268}{155} - \frac{403}{155} = -\frac{135}{155} = -\frac{27}{31}.$$

Um die Richtigkeit des Ergebnisses zu überprüfen, kann der für x gefundene Wert auch noch in die zweite Gleichung eingesetzt werden – es muss sich derselbe Wert für y ergeben. Außerdem sollten für eine Probe des Ergebnisses die für x und y ermittelten Werte noch in das ursprüngliche LGS eingesetzt werden.

Als Lösung des LGS ergibt sich also $(\frac{67}{31}; -\frac{27}{31})$. Diese Lösung wurde näherungsweise bereits graphisch ermittelt, siehe Abb. 1.4 auf S. 6. ■

Ein weiteres Lösungsverfahren für LGS mit zwei Gleichungen und zwei Variablen ist das *Einsetzungsverfahren*. Hierbei wird eine der beiden Gleichungen nach einer Variablen umgestellt und der dabei erhaltene Term in die andere Gleichung eingesetzt.

Beispiel 1.7

Einsetzungsverfahren

Es wird erneut das LGS aus den Beispielen 1.3 und 1.6 betrachtet und die zweite Gleichung nach x umgestellt:³

$$x = -\frac{4}{3}y + 1.$$

Der erhaltene Term wird nun in die erste Gleichung an Stelle von x eingesetzt:

$$4\left(-\frac{4}{3}y + 1\right) - 5y = 13.$$

Durch Vereinfachung dieser Gleichung und Auflösung nach y ergibt sich:

$$\begin{aligned} -\frac{16}{3}y + 4 - 5y &= 13 \\ -\frac{31}{3}y &= 9 \\ y &= -\frac{27}{31}. \end{aligned}$$

Dieser Wert wird nun in die nach x umgestellte zweite Gleichung eingesetzt:

$$x = -\frac{4}{3} \cdot \left(-\frac{27}{31}\right) + 1 = \frac{36}{31} + 1 = \frac{67}{31}.$$

Als Lösung des LGS ergibt sich also $\left(\frac{67}{31}; -\frac{27}{31}\right)$. Auch bei diesem Verfahren ist eine Probe durch Einsetzen der Lösung in das gegebene LGS sinnvoll. ■

Ein drittes Lösungsverfahren ist das *Additionsverfahren*. Dieses ist von besonderer Bedeutung, da seine Vorgehensweise auch auf lineare Gleichungssysteme mit mehr als zwei Gleichungen und mehr als zwei Variablen angewendet werden kann. Bei dem Additionsverfahren werden die beiden Gleichungen mit geeigneten Zahlen multipliziert, so dass nach der Addition beider Gleichungen in der entstehenden Gleichung eine Variable verschwindet (eliminiert wird).

Beispiel 1.8

Additionsverfahren

Wir betrachten nochmals das LGS aus den Beispielen 1.3, 1.6, 1.7 und überlegen, mit welchen Zahlen die Gleichungen multipliziert werden können, so dass nach Addition beider Gleichungen eine Gleichung entsteht, in der x nicht auftritt:

$$\begin{array}{ll} \text{I} & 4x - 5y = 13 \quad | \cdot (-3) \\ \text{II} & 3x + 4y = 3 \quad | \cdot 4 \end{array}.$$

³Es kann bei dem Einsetzungsverfahren wahlweise die erste oder die zweite Gleichung nach einer Variablen umgestellt werden, wobei auch die Wahl der Variablen frei ist, falls alle Koeffizienten von Null verschieden sind. Besonders sinnvoll ist die Verwendung des Einsetzungsverfahrens jedoch bei Gleichungssystemen, in denen ein Koeffizient Null ist und somit in einer Gleichung eine Variable nicht auftritt. In diesem Falle ist der für diese Variable in die andere Gleichung einzusetzende Term besonders einfach.

Durch Multiplikation der Gleichungen mit den Faktoren (-3) bzw. 4 ergibt sich:

$$\begin{array}{rcl} \text{I}' & -12x + 15y & = -39 \\ \text{II}' & 12x + 16y & = 12 \end{array}$$

Im nächsten Schritt wird die ursprüngliche Gleichung I beibehalten und die zweite Gleichung durch die Summe $\text{I}' + \text{II}'$ der multiplizierten Gleichungen ersetzt:

$$\begin{array}{rcl} \text{I}'' & 4x - 5y & = 13 \\ \text{II}'' & 31y & = -27 \end{array}$$

An dieser Stelle kann nun y bestimmt und, analog zum Einsetzungsverfahren, in die erste Gleichung eingesetzt werden, womit sich dann x berechnen lässt. Mit Blick auf die spätere Erweiterung des Additionsverfahrens auf LGS mit mehr als zwei Gleichungen und Variablen ist es jedoch sinnvoll, einen weiteren Umformungsschritt vorzunehmen, bei dem y aus der ersten Gleichung eliminiert wird:

$$\begin{array}{rcl} \text{I}'' & 4x - 5y & = 13 \quad | \cdot 1 \\ \text{II}'' & 31y & = -27 \quad | \cdot \frac{5}{31} \end{array}$$

Es werden nun die derart multiplizierten Gleichungen addiert; die entstehende Gleichung ersetzt I'' :

$$\begin{array}{rcl} \text{I}''' & 4x & = \frac{268}{31} \\ \text{II}''' & 31y & = -27 \end{array}$$

Als Lösung des LGS ergibt sich daraus wiederum $(\frac{67}{31}; -\frac{27}{31})$. ■

Das Additionsverfahren wurde in dem Beispiel 1.8 recht ausführlich dargestellt. Um das Verfahren kürzer aufzuschreiben, wird meist darauf verzichtet, den ersten Umformungsschritt (I' und II') zu notieren. Die Multiplikation der Gleichungen sowie die Addition werden also in einem einzigen Schritt vorgenommen.

Beispiel 1.9

Additionsverfahren in Kurzschreibweise

$$\begin{array}{rclclcl} 2x + 3y = 4 & | \cdot (-3) & 2x + 3y = 4 & | \cdot 1 & 2x & = \frac{32}{5} \\ 3x + 2y = 8 & | \cdot 2 & -5y = 4 & | \cdot \frac{3}{5} & -5y & = 4 \end{array}$$

Das Gleichungssystem hat somit die (eindeutige) Lösung $(\frac{16}{5}; -\frac{4}{5})$. ■

Die durch die Umformungsschritte des Additionsverfahrens erfolgenden Veränderungen der Gleichungen lassen sich durch die Darstellung der durch sie beschriebenen Geraden sichtbar machen; Abb. 1.7 zeigt dies für das in dem Beispiel 1.9 gelöste LGS. Die Vereinfachung des Gleichungssystems kommt dabei dadurch zum Ausdruck, dass die durch die Gleichungen beschriebenen Geraden in eine spezielle Lage, parallel zu den Koordinatenachsen, überführt werden.

Durch die Darstellungen in Abb. 1.7 wird deutlich, dass sich die Lösungsmengen der einzelnen Gleichungen durchaus verändern (im Gegensatz zu Äquivalenzumformungen von Gleichungen); der Schnittpunkt, welcher die Lösung des Gleichungssystems repräsentiert, jedoch unverändert bleibt.

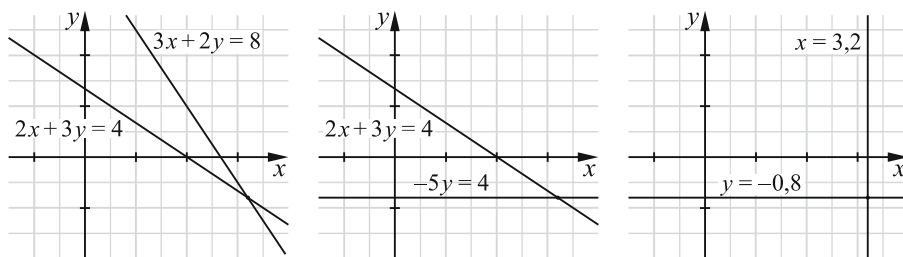


Abb. 1.7: Wirkungen der Schritte des Additionsverfahrens auf die durch die Gleichungen eines linearen Gleichungssystems (siehe Beispiel 1.9) bestimmten Geraden

Neben den in den Beispielen 1.8 und 1.9 vorgenommenen Umformungen kann es bei bestimmten Gleichungssystemen sinnvoll sein, Gleichungen zu vertauschen. So sollten z. B. bei dem LGS

$$\begin{aligned} -14y &= -11 \\ 5x - 9y &= 12 \end{aligned}$$

die Gleichungen I und II vertauscht werden, womit der erste Schritt des Additionsverfahrens überflüssig wird. Eine Veränderung der Lösungsmenge erfolgt durch eine Vertauschung der beiden Gleichungen natürlich nicht.

Bei LGS mit 2 Gleichungen und 2 Variablen, die *nicht lösbar* sind oder eine *unendliche Lösungsmenge* besitzen, wird dies bereits nach den ersten Schritten des Additionsverfahrens deutlich.

Beispiel 1.10

Auf die LGS aus den Beispielen 1.4 und 1.5 wird das Additionsverfahren angewendet:

$$\begin{array}{rcl} x + y & = & -1 \quad | \cdot 3 \\ -3x - 3y & = & 9 \quad | \cdot 1 \\ \hline x + y & = & -1 \\ 0 & = & 6 \end{array}$$

Es entsteht eine Gleichung, die offenbar eine falsche Aussage ist. Das gesamte LGS ist somit unlösbar, denn jede Lösung müsste dazu führen, dass *beide* Gleichungen erfüllt sind.

$$\begin{array}{rcl} x + y & = & -1 \quad | \cdot 3 \\ -3x - 3y & = & 3 \quad | \cdot 1 \\ \hline x + y & = & -1 \\ 0 & = & 0 \end{array}$$

Es entsteht eine Gleichung, die immer erfüllt ist. Die Lösungsmenge der ersten Gleichung ist zugleich die Lösungsmenge des LGS. Eine der Gleichungen des gegebenen LGS ist somit „überflüssig“. ■

Äquivalenzumformungen linearer Gleichungssysteme

Äquivalenzumformungen eines LGS sind Umformungen, welche die Lösungsmenge des Systems nicht verändern (Äquivalenz: Gleichwertigkeit):

- Vertauschen zweier Gleichungen,
- Multiplikation einer Gleichung mit einer reellen Zahl (außer Null),
- Addition zweier Gleichungen.

Falls ein lineares Gleichungssystem noch nicht in der Form

$$a_{11}x + a_{12}y = b_1$$

$$a_{21}x + a_{22}y = b_2$$

vorliegt, ist auch die Addition gleicher Terme auf beiden Seiten einer Gleichung zulässig. Grundsätzlich sind äquivalente Umformungen einzelner Gleichungen auch Äquivalenzumformungen in Bezug auf ein LGS. Die umgekehrte Aussage gilt jedoch nicht: die für ein Gleichungssystem zulässige Äquivalenzumformung der Addition zweier Gleichungen führt im Allgemeinen zu Einzelgleichungen mit veränderten Lösungsmengen, vgl. Abb. 1.7.

1.1.4 Sachsituationen, die auf LGS mit zwei Variablen führen

Häufig werden Lineare Gleichungssysteme für die Lösung von *Mischungsproblemen* verwendet.

Beispiel 1.11

Aus zwei Gold-Silber-Legierungen, in denen sich die Metallmassen wie 2 : 3 bzw. 3 : 7 verhalten, sind 8 kg einer neuen Legierung mit dem Verhältnis 5 : 11 herzustellen. Wie viel Kilogramm der Legierungen sind dabei zu verwenden?

Im Folgenden werden Lösungsschritte bei der Bearbeitung dieser Aufgabe beschrieben, die für viele Text- bzw. Sachaufgaben sinnvoll sind, die auf das Lösen linearer Gleichungssysteme führen.

- Entnahme der wesentlichen Informationen aus dem Text
 - In der ersten Legierung verhält sich die Gold- zur Silbermasse wie 2 : 3, d. h. ein Kilogramm der Legierung enthält $\frac{2}{5}$ kg Gold und $\frac{3}{5}$ kg Silber.
 - Ein Kilogramm der zweiten Legierung enthält $\frac{3}{10}$ kg Gold und $\frac{7}{10}$ kg Silber.
 - Ein Kilogramm der durch Mischung herzustellenden Legierung soll $\frac{5}{16}$ kg Gold und $\frac{11}{16}$ kg Silber enthalten.
 - Es sollen 8 kg der neuen Legierung hergestellt werden, hierin müssen also $\frac{5}{16} \cdot 8 \text{ kg} = \frac{40}{16} \text{ kg} = \frac{5}{2} \text{ kg}$ Gold und $\frac{88}{16} \text{ kg} = \frac{11}{2} \text{ kg}$ Silber enthalten sein.
- Festlegen der Variablen
 - x – benötigte Masse (in kg) der ersten Legierung
 - y – benötigte Masse (in kg) der zweiten Legierung
- Aufstellen von Gleichungen
 - Gold: $\frac{2}{5}x + \frac{3}{10}y = \frac{5}{2}$
 - Silber: $\frac{3}{5}x + \frac{7}{10}y = \frac{11}{2}$

Es empfiehlt sich bei diesem LGS, beide Gleichungen mit 10 zu multiplizieren:

 - Gold: $4x + 3y = 25$
 - Silber: $6x + 7y = 55$

- Lösen des Gleichungssystems mithilfe des Additionsverfahrens

$$\begin{array}{rclclcl} 4x + 3y = 25 & | \cdot (-3) & 4x + 3y = 25 & | \cdot 1 & 4x & = & 4 \\ 6x + 7y = 55 & | \cdot 2 & 5y = 35 & | \cdot (-\frac{3}{5}) & 5y & = & 35 \end{array}$$

Als Lösung ergibt sich damit $x = 1$ und $y = 7$.

- Überprüfung des Ergebnisses

$$\text{Gold:} \quad \frac{2}{5} \cdot 1 + \frac{3}{10} \cdot 7 = \frac{4}{10} + \frac{21}{10} = \frac{25}{10} = \frac{40}{16} = \frac{5}{2}$$

$$\text{Silber:} \quad \frac{3}{5} \cdot 1 + \frac{7}{10} \cdot 7 = \frac{6}{10} + \frac{49}{10} = \frac{55}{10} = \frac{88}{16} = \frac{11}{2}$$

- Es werden also 1 kg der 2:3-Legierung und 7 kg der 3:7-Legierung benötigt. ■

Das folgende Beispiel ist weniger als echte Anwendung, denn als eine „Knobelaufgabe“ anzusehen.

Beispiel 1.12

Eine Aufgabe aus vergangenen Zeiten: Eine Köchin hat 18 Heringe für 18 Groschen gekauft. Es waren gute Heringe für die Herrschaft, das Stück für 5 Groschen, und schlechte Heringe für das Gesinde, 5 Stück für einen Groschen. Wie viele gute und wie viele schlechte Heringe waren es?

Wird die Anzahl der guten Heringe mit x und die der schlechten Heringe mit y bezeichnet, so lässt sich aus dem Text das folgende Gleichungssystem aufstellen:

$$\text{Anzahl der Heringe:} \quad x + y = 18$$

$$\text{Kosten:} \quad 5x + \frac{1}{5}y = 18.$$

Aufgrund der sehr einfachen Struktur der ersten Gleichung empfiehlt es sich, für die Lösung dieses LGS das Einsetzungsverfahren zu verwenden. Die nach y umgestellte erste Gleichung wird in die zweite Gleichung eingesetzt:

$$5x + \frac{1}{5}(18 - x) = 18$$

$$\text{bzw.} \quad 25x + 18 - x = 90.$$

Auflösen nach x ergibt $x = 3$ und Einsetzen in die erste Gleichung $y = 15$. Die Köchin hat also 3 gute und 15 schlechte Heringe gekauft. Eine Überprüfung dieses Ergebnisses anhand des Aufgabentextes ergibt einen Preis von $5 \cdot 3$ Groschen und 15 Fünftel Groschen, insgesamt also 18 Groschen. ■

1.1.5 Aufgaben zu Abschnitt 1.1

1. Stellen Sie die Lösungsmengen der folgenden Gleichungen graphisch dar. Geben Sie die Lösungsmengen außerdem in der Form

$$L = \left\{ \begin{pmatrix} x \\ y \end{pmatrix} \mid \begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} m_1 \\ m_2 \end{pmatrix} + \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}; t \in \mathbb{R} \right\}$$

an (ermitteln Sie hierzu m_1 , m_2 , n_1 und n_2 , vgl. Beispiel 1.2).

$$\text{a) } 5x + 4y = 13 \quad \text{b) } 45x - 84y = 132 \quad \text{c) } 5 \cdot (x - 2y) + 10y = \frac{45}{2}$$

2. Lösen Sie die folgenden linearen Gleichungssysteme mithilfe des Einsetzungs- und mithilfe des Gleichsetzungsverfahrens. Vergleichen Sie Ihre Lösungen.

a)
$$\begin{aligned} 3x - 15y &= 180 \\ 2x + 4y &= 64 \end{aligned}$$

b)
$$\begin{aligned} \frac{1}{5}x - \frac{1}{6}y &= 2 \\ \frac{2}{3}x + \frac{3}{4}y &= 1 \end{aligned}$$

3. Lösen Sie die folgenden LGS mithilfe des Additionsverfahrens. Stellen Sie außerdem die durch die Gleichungen bestimmten Geraden graphisch dar.

a)
$$\begin{aligned} 3x + 2y &= 149 \\ 2x - y &= 55 \end{aligned}$$

b)
$$\begin{aligned} 5x + 3y &= 19 \\ 6x - 2y &= 6 \end{aligned}$$

c)
$$\begin{aligned} 36x + 15y &= 8 \\ 3x + \frac{5}{4}y &= \frac{3}{4} \end{aligned}$$

d)
$$\begin{aligned} x - 2y &= -2 \\ -\frac{1}{2}x + y &= 1 \end{aligned}$$

4. Es sei c eine beliebige, aber feste reelle Zahl. Lösen Sie das folgende LGS in Abhängigkeit von c .

$$\begin{aligned} x + y &= c \\ x - y &= 3 \end{aligned}$$

5. Begründen Sie, dass ein lineares Gleichungssystem der Form

$$\begin{aligned} a_{11}x + a_{12}y &= 0 \\ a_{21}x + a_{22}y &= 0 \end{aligned}$$

für beliebige $a_{11}, a_{12}, a_{21}, a_{22}$ lösbar ist. Welcher Zusammenhang muss zwischen a_{11}, a_{12}, a_{21} und a_{22} bestehen, damit das LGS unendlich viele Lösungen besitzt?

6. Ein Händler mischt zwei Sorten Tee. Nimmt er von der ersten Sorte 120 kg und von der zweiten Sorte 90 kg, so kostet ihn 1 kg der Mischung 25 €. Mischt er dagegen 10 kg der ersten mit 150 kg der zweiten Sorte, so kostet ihn 1 kg der Mischung 30 €. Wie teuer war 1 kg jeder Sorte?
7. Zwei Stahlsorten enthalten 12% bzw. 30% Nickel. Man will aus beiden Stählen einen Stahl mit einem Nickelgehalt von 25% und einer Masse von 180 kg schmelzen. Wieviel Kilogramm Stahl jeder Sorte werden benötigt, wenn man bei der Rechnung von Verlusten beim Schmelzen absieht?
8. Anja, Berta und Claudius fahren mit dem Rad von demselben Ausgangspunkt aus zur Schule. Anja startet um 7.30 Uhr, Berta um 7.33 Uhr, Claudius um 7.35 Uhr. Anja fährt im Durchschnitt 3 km/h langsamer als Berta, Claudius um 3 km/h schneller. Alle drei kommen zugleich an. Wie lange braucht Berta und wie schnell fährt sie? Wie lang ist der Weg zur Schule?
9. Familie Wegert kauft für den Garten 4 Apfelhalbstämme und 2 Apfelhochstämme. Sie bezahlt 192 €. Ihre Nachbarn kaufen 3 Apfelhalbstämme und einen Apfelhochstamm für 129,50 €. Wie viel kostet jede Sorte?
10. Aus einem alten Rechenbuch: Eine Anzahl von Personen hatte in einem Gasthaus eine Zeche zu bezahlen. Zahlte jede 4,35 Mark, so fehlten noch 20 Pf an der ganzen Summe. Zahlte jede 4,40 Mark, so waren es 20 Pf zuviel. Wie groß war die zu zahlende Summe? Wie viele Personen waren es?

1.2 Lineare Gleichungssysteme mit drei Variablen

1.2.1 Lösungsmengen linearer Gleichungen und Gleichungssysteme mit drei Variablen

Lineare Gleichungen mit 3 Variablen lassen sich allgemein in der Form

$$ax + by + cz = d \quad \text{mit } a, b, c, d \in \mathbb{R} \text{ und } x, y, z \in \mathbb{R} \quad (1.10)$$

schreiben. Als Lösungen kommen Tripel $(x; y; z)$ reeller Zahlen in Frage, die oft auch in der Spaltenschreibweise $\begin{pmatrix} x \\ y \\ z \end{pmatrix}$ geschrieben werden. Die *Menge aller prinzipiell möglichen Lösungen* einer linearen Gleichung mit drei Variablen ist also die Menge aller Tripel reeller Zahlen. Diese Menge wird mit $\mathbb{R}^3$ bezeichnet; geometrisch wird sie durch den (dreidimensionalen) Raum veranschaulicht. Tripel reeller Zahlen und somit Lösungen linearer Gleichungen in drei Variablen können in einem dreidimensionalen Koordinatensystem dargestellt werden.

Beispiel 1.13

Es werden einige Lösungen der Gleichung

$$-3x + \frac{3}{2}y - 5z = -1$$

ermittelt, indem die Gleichung nach z umgestellt wird:

$$z = -\frac{3}{5}x + \frac{3}{10}y + \frac{1}{5}.$$

Durch Einsetzen von Werten für x und y und Berechnung der zugehörigen Werte für z lassen sich Lösungstriple der Gleichung ermitteln. Nach Einzeichnen der zugehörigen Punkte in ein dreidimensionales Koordinatensystem wird sichtbar, dass die zu den Lösungen gehörenden Punkte in einer Ebene liegen, siehe Abb. 1.8. Diese Ebene ist der Graph der linearen Funktion f (in zwei Variablen) mit $f(x, y) = -\frac{3}{5}x + \frac{3}{10}y + \frac{1}{5}$.

Ohne technische Hilfsmittel ist die graphische Darstellung von Lösungsmengen linearer Gleichungen bzw. LGS mit drei Variablen recht mühsam. In dem Abschnitt 1.5 wird auf hierfür geeignete Computersoftware eingegangen. ■

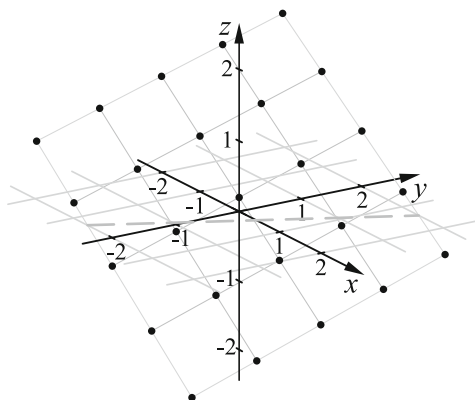


Abb. 1.8: Graphische Darstellung einiger Lösungstriple einer linearen Gleichung mit 3 Variablen

Das Beispiel 1.13 ist in dem Sinne verallgemeinerbar, dass die Lösungsmenge jeder linearen Gleichung der Form 1.10 durch eine Ebene dargestellt werden kann, falls nicht zugleich $a=0$ und $b=0$ und $c=0$ ist.⁴

Die Lösungsmenge der Gleichung aus Beispiel 1.13 lässt sich in der Form

$$L = \{(x; y; z) \mid x, y \in \mathbb{R}; z = -\frac{3}{5}x + \frac{3}{10}y + \frac{1}{5}\} \quad (1.11)$$

schreiben. Allerdings ist eine derartige Darstellung nicht auf alle Gleichungen der Form (1.10) übertragbar. Da der Koeffizient c auch den Wert Null haben kann, ist eine Auflösung nach z nicht in jedem Falle möglich. Da allerdings $a \neq 0$ oder $b \neq 0$ oder $c \neq 0$ verlangt wird, kann die Lösungsmenge jeder linearen Gleichung (1.10) in *mindestens* einer der drei folgenden Schreibweisen dargestellt werden:

$$\begin{aligned} L &= \{(x; y; z) \mid x, y \in \mathbb{R}; z = -\frac{a}{c}x - \frac{b}{c}y + \frac{d}{c}\}, \\ L &= \{(x; y; z) \mid x, z \in \mathbb{R}; y = -\frac{a}{b}x - \frac{c}{b}z + \frac{d}{b}\}, \\ L &= \{(x; y; z) \mid y, z \in \mathbb{R}; x = -\frac{b}{a}y - \frac{c}{a}z + \frac{d}{a}\}. \end{aligned} \quad (1.12)$$

In jedem Falle treten in der Lösungsmenge zwei unabhängige Variablen auf.

Für eine einheitliche Darstellung der Lösungsmengen linearer Gleichungen mit drei Variablen werden zwei *Parameter* s und t eingeführt (vgl. Abschnitt 1.1.1, S. 3f.), die mit zwei der Variablen übereinstimmen können. Die Lösungsmenge (1.12) einer linearen Gleichung der Gestalt (1.10) mit $a \neq 0$ oder $b \neq 0$ oder $c \neq 0$ lässt sich damit in jedem Falle folgendermaßen schreiben:

$$L = \left\{ (x; y; z) \left| \begin{array}{l} x = m_1 \cdot s + n_1 \cdot t + p_1 \\ y = m_2 \cdot s + n_2 \cdot t + p_2 \\ z = m_3 \cdot s + n_3 \cdot t + p_3 \end{array} \right. ; s, t \in \mathbb{R} \right\}. \quad (1.13)$$

Dabei sind $m_1, m_2, m_3, n_1, n_2, n_3$ sowie p_1, p_2 und p_3 feste reelle Zahlen, die von den Koeffizienten der jeweiligen Gleichung abhängen; die Parameter s und t durchlaufen jeweils den gesamten Bereich der reellen Zahlen.

Beispiel 1.14

Die Lösungsmenge der Gleichung $-3x + \frac{3}{2}y - 5z = -1$ aus dem Beispiel 1.13 lässt sich zunächst durch (1.11) angeben. Setzt man $s=x$ und $t=y$, so ergibt sich

$$L = \left\{ (x; y; z) \left| \begin{array}{l} x = 1 \cdot s \\ y = 1 \cdot t \\ z = -\frac{3}{5} \cdot s + \frac{3}{10} \cdot t + \frac{1}{5} \end{array} \right. ; s, t \in \mathbb{R} \right\}.$$

Die Parameter können auch anders festgelegt werden. Mit $s=y$ und $t=z$ ist:

$$L = \left\{ (x; y; z) \left| \begin{array}{l} x = \frac{1}{2} \cdot s - \frac{5}{3} \cdot t + \frac{1}{3} \\ y = 1 \cdot s \\ z = 1 \cdot t \end{array} \right. ; s, t \in \mathbb{R} \right\}.$$

⁴Falls in einer Gleichung der Form (1.10) $a=b=c=0$ ist, so hat die Gleichung die Gestalt $0=d$. Für $d=0$ ist dann jedes Tripel $(x; y; z) \in \mathbb{R}^3$ eine Lösung der Gleichung; für $d \neq 0$ existiert keine Lösung, da die Gleichung unabhängig von x, y und z eine falsche Aussage ist.

Ohne dies an dieser Stelle nachweisen zu können, sei erwähnt, dass durch beide Darstellungen in der Tat dieselbe Lösungsmenge beschrieben wird. ■

Wie bereits in dem Abschnitt 1.1.1 erläutert wurde, wird für die Darstellung von Lösungsmengen aufgrund der Übersichtlichkeit oft die Spaltenschreibweise verwendet. Die Lösungsmenge (1.13) lässt sich dann mit

$$\begin{pmatrix} x \\ y \\ z \end{pmatrix} = s \cdot \begin{pmatrix} m_1 \\ m_2 \\ m_3 \end{pmatrix} + t \cdot \begin{pmatrix} n_1 \\ n_2 \\ n_3 \end{pmatrix} + \begin{pmatrix} p_1 \\ p_2 \\ p_3 \end{pmatrix}. \quad (1.14)$$

angeben. Auch hier besteht ein Bezug zur geometrischen Interpretation der Lösungsmenge als Ebene im Raum: Die Vektoren $\begin{pmatrix} m_1 \\ m_2 \\ m_3 \end{pmatrix}$ und $\begin{pmatrix} n_1 \\ n_2 \\ n_3 \end{pmatrix}$ spannen diese Ebene auf (sind Richtungsvektoren); $(p_1; p_2; p_3)$ sind die Koordinaten eines festen Punktes der „Lösungsebene“.

Beispiel 1.15

Wir greifen nochmals die Gleichung aus dem Beispiel 1.13 und die beiden in dem Beispiel 1.14 herausgearbeiteten Darstellungen ihrer Lösungsmenge auf. In Spaltenform lassen sich diese folgendermaßen schreiben:

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \left| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = s \cdot \begin{pmatrix} 1 \\ 0 \\ -\frac{3}{5} \end{pmatrix} + t \cdot \begin{pmatrix} 0 \\ 1 \\ \frac{3}{10} \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ \frac{1}{5} \end{pmatrix}; s, t \in \mathbb{R} \right\} \quad \text{bzw.}$$

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \left| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = s \cdot \begin{pmatrix} \frac{1}{2} \\ 1 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} -\frac{5}{3} \\ 0 \\ 1 \end{pmatrix} + \begin{pmatrix} \frac{1}{3} \\ 0 \\ 0 \end{pmatrix}; s, t \in \mathbb{R} \right\}$$

Abb. 1.9 zeigt für jede dieser Darstellungen die Richtungsvektoren der die Lösungsmenge repräsentierenden Ebene. ■

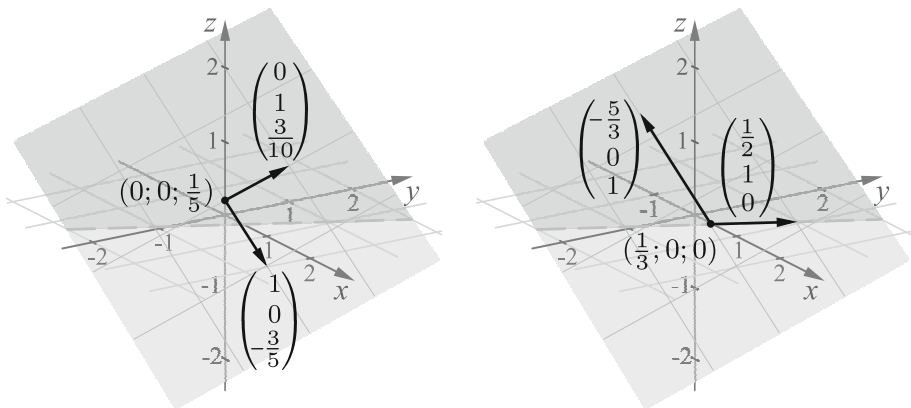


Abb. 1.9: Darstellungen der Lösungsmenge einer linearen Gleichung durch eine Ebene mit Richtungsvektoren

Zusammenfassung:**Lösungsmengen linearer Gleichungen mit drei Variablen**

Lineare Gleichungen der Form

$$ax + by + cz = d$$

mit $a, b, c, d \in \mathbb{R}$ und $a \neq 0$ oder $b \neq 0$ oder $c \neq 0$ besitzen zweiparametrische Lösungsmengen, welche sich geometrisch als Ebenen interpretieren lassen.

Die Lösungsmengen können in der Form

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mid \begin{pmatrix} x \\ y \\ z \end{pmatrix} = s \cdot \begin{pmatrix} m_1 \\ m_2 \\ m_3 \end{pmatrix} + t \cdot \begin{pmatrix} n_1 \\ n_2 \\ n_3 \end{pmatrix} + \begin{pmatrix} p_1 \\ p_2 \\ p_3 \end{pmatrix} ; s, t \in \mathbb{R} \right\}$$

dargestellt werden. Dabei sind $m_1, m_2, m_3, n_1, n_2, n_3$ sowie p_1, p_2 und p_3 feste reelle Zahlen, die von den Koeffizienten der jeweiligen Gleichung abhängen; die Parameter s und t durchlaufen jeweils den gesamten Bereich der reellen Zahlen.

Vergleicht man Lösungsmengen linearer Gleichungen mit drei Variablen und linearer Gleichungen mit zwei Variablen (siehe Abschnitt 1.1.1), so lässt sich eine Analogie feststellen:

Die Grundmenge aller in Frage kommenden Lösungen ist bei Gleichungen mit zwei Variablen $\mathbb{R}^2$ (die Menge aller Paare $(x; y)$ mit $x, y \in \mathbb{R}$), also *zweidimensional*. Lösungsmengen linearer Gleichungen mit zwei Variablen sind (außer bei solchen Gleichungen, die sich so umformen lassen, dass beide Variablen darin nicht auftreten) einparametrisch und werden durch *eindimensionale* geometrische Objekte (Geraden) repräsentiert. Hingegen ist die Grundmenge bei Gleichungen mit drei Variablen $\mathbb{R}^3$, also *dreidimensional*, und die Lösungen sind (von den genannten Ausnahmen abgesehen) zweiparametrisch und können durch Ebenen, also *zweidimensionale* Objekte dargestellt werden. In beiden Fällen ist die *Dimension der Lösungsmenge um Eins geringer als die Dimension der Grundmenge* – eine Gleichung „verringert“ gewissermaßen „die Dimension“ der Grundmenge um Eins. Diese – bislang nur anschaulich anhand von Beispielen gewonnene – Erkenntnis lässt sich verallgemeinern: Eine lineare Gleichung in n Variablen hat die Grundmenge $\mathbb{R}^n$ und als Lösungsmenge im Allgemeinen einen $n-1$ -dimensionalen Unterraum von $\mathbb{R}^n$, auch als Hyperebene bezeichnet.⁵

⁵Der Begriff „Unterraum“ (bzw. Teilraum) wird hier nur erwähnt und erst in dem Kapitel 5 behandelt, da hierzu Grundlagen aus der Theorie der Vektorräume erforderlich sind. In „populärwissenschaftlichem“ Sinne stellen Hyperebenen Verallgemeinerungen von Ebenen in folgendem Sinne dar: Gewöhnliche Ebenen sind zweidimensionale lineare Unterstrukturen des dreidimensionalen Raumes, Hyperebenen sind analog dazu lineare Unterstrukturen eines beliebig dimensionalen Raumes, wobei die Dimension einer Hyperebene stets um Eins geringer ist als die Dimension des Raumes zu dem sie eine Hyperebene ist.

Alle *Lösungen linearer Gleichungssysteme mit drei Variablen und zwei Gleichungen* müssen beiden Lösungsmengen der Einzelgleichungen angehören. Dabei können drei Fälle auftreten:

- a) Die beiden Ebenen, welche die Lösungsmengen der beiden Gleichungen darstellen, schneiden sich in einer Geraden. Diese Gerade stellt dann die Lösungsmenge des Gleichungssystems dar.
- b) Die den beiden Gleichungen des LGS zugeordneten Ebenen sind parallel. Die beiden Gleichungen besitzen keine gemeinsamen Lösungen und die Lösungsmenge des Systems ist somit leer.
- c) Die beiden Gleichungen haben dieselbe Lösungsmenge. Die Lösungsmenge des Gleichungssystems ist mit dieser Lösungsmenge ebenfalls identisch und wird durch eine Ebene dargestellt.

Beispiel 1.16

Lineare Gleichungssysteme mit drei Variablen und zwei Gleichungen

a) LGS mit einparametrischer Lösungsmenge	b) LGS mit leerer Lösungsmenge	c) LGS mit zweiparametrischer Lösungsmenge
$\begin{aligned}\frac{3}{2}x - 2y + 4z &= 1 \\ x + y - 3z &= 2\end{aligned}$	$\begin{aligned}-3x - 3y + 9z &= 30 \\ x + y - 3z &= 2\end{aligned}$	$\begin{aligned}-3x - 3y + 9z &= -6 \\ x + y - 3z &= 2\end{aligned}$

Abb. 1.10 zeigt die Ebenen, welche durch die Gleichungen dieser LGS beschrieben werden. Für das Beispiel c) fallen diese beiden Ebenen zusammen. Multipliziert man die erste Gleichung in c) mit $-\frac{1}{3}$, so ergibt sich die zweite Gleichung. Diese enthält also keine „neue“ Bedingung an die Lösungsmenge. Eine der beiden Gleichungen ist „verzichtbar“ und die Lösungsmenge des LGS c) ist mit der Lösungsmenge jeder der beiden Gleichungen identisch. Hingegen ergibt sich durch Multiplikation der ersten Gleichung in dem Beispiel b) mit $-\frac{1}{3}$ die Gleichung $x + y - 3z = -10$. Für jedes Lösungstriplet $(x; y; z)$ des LGS b) müsste also sowohl $x + y - 3z = -10$ als auch $x + y - 3z = 2$ gelten. Dazu müsste aber $-10 = 2$ sein, also kann es keine Lösung dieses LGS geben. ■

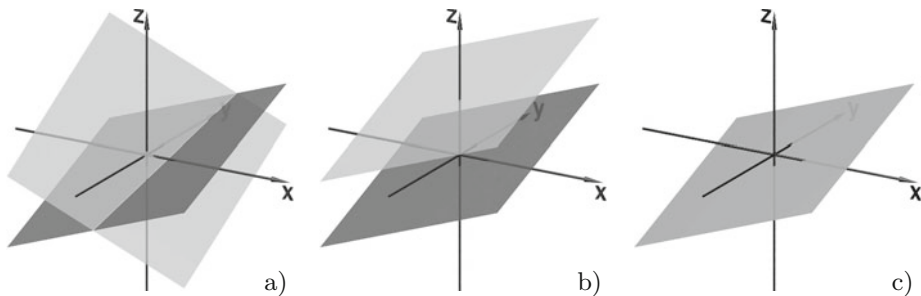


Abb. 1.10: Darstellung von Lösungsmengen linearer Gleichungssysteme mit drei Variablen und zwei Gleichungen durch Ebenen (Beispiel 1.16)

Die hellgraue Ebene stellt jeweils die Lösungsmenge der ersten Gleichung, die dunklere Ebene die der zweiten Gleichung dar, bei c) fallen beide Ebenen zusammen.

Anhand der untersuchten Fälle und ihrer geometrischen Interpretation wurde deutlich, dass lineare Gleichungssysteme mit drei Variablen und zwei Gleichungen in keinem Falle eindeutig lösbar sein (also genau ein Lösungstriplett besitzen) können, denn dazu müssten zwei Ebenen genau einen Schnittpunkt haben. Kommt jedoch eine dritte Gleichung hinzu (durch die dann auch eine dritte Ebene beschrieben wird), so ist eine eindeutige Lösung möglich.

Beispiel 1.17

Das lineare Gleichungssystem

$$\frac{3}{2}x - 2y + 4z = 1$$

$$2x - \frac{9}{2}y - 2z = -\frac{1}{2}$$

$$x + y - 3z = 2$$

wird durch drei Ebenen repräsentiert, die sich in genau einem Punkt schneiden, vgl. Abb. 1.11. Dieser Schnittpunkt stellt die einzige Lösung des LGS dar, denn jede Lösung des LGS muss alle drei Gleichungen erfüllen und somit in allen drei Ebenen liegen. Der Schnittpunkt hat die Koordinaten $(\frac{22}{15}; \frac{11}{15}; \frac{1}{15})$, dadurch ist auch die Lösung des LGS gegeben. Ein rechnerisches Verfahren zur Bestimmung dieser Lösung wird im folgenden Abschnitt behandelt. ■

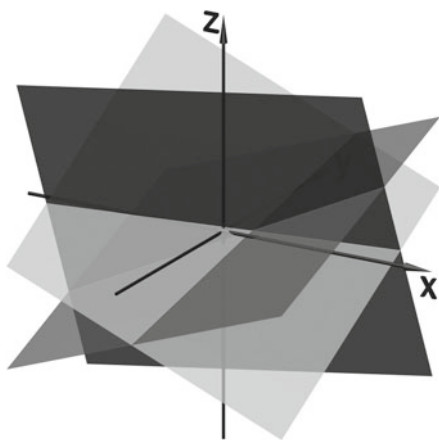


Abb. 1.11: Darstellung der Lösungsmengen der Gleichungen eines LGS mit drei Gleichungen und drei Variablen (Bsp. 1.17)

Die hellgraue Ebene stellt die Lösungsmenge der ersten Gleichung, die dunkelgraue Ebene die der zweiten Gleichung und die mittelgraue Ebene die der dritten Gleichung dar.

Farbige Darstellungen und zugehörige Videos, welche die drei Ebenen aus verschiedenen Blickrichtungen zeigen, enthält die Internetseite www.afiller.de/linalg.

In dem Abschnitt 1.5 werden Hinweise zur Darstellung von Lösungsmengen linearer Gleichungen bzw. Gleichungssysteme mithilfe des Computers gegeben. Derartige Darstellungen lassen sich auch drehen, womit eine Abschätzung der Lösung möglich ist.

Nicht alle linearen Gleichungssysteme mit drei Gleichungen und drei Variablen sind eindeutig lösbar.

Beispiel 1.18

Wir betrachten die Lösungsmengen der folgenden LGS (siehe dazu Abb. 1.12).

I	$3x - 4y + 8z = 16$	I	$3x - 4y + 8z = 2$
a) II	$6x - y + 5z = -4$	b) II	$6x - y + 5z = 4$
III	$-3x - 3y + 3z = -8$	III	$-3x - 3y + 3z = -2$

In Abb. 1.12 ist erkennbar, dass das LGS a) offenbar keine Lösung besitzt, während LGS b) eine unendliche (eiparametrische) Lösungsmenge hat. Rechnerisch wird dies durch folgende Überlegung plausibel:

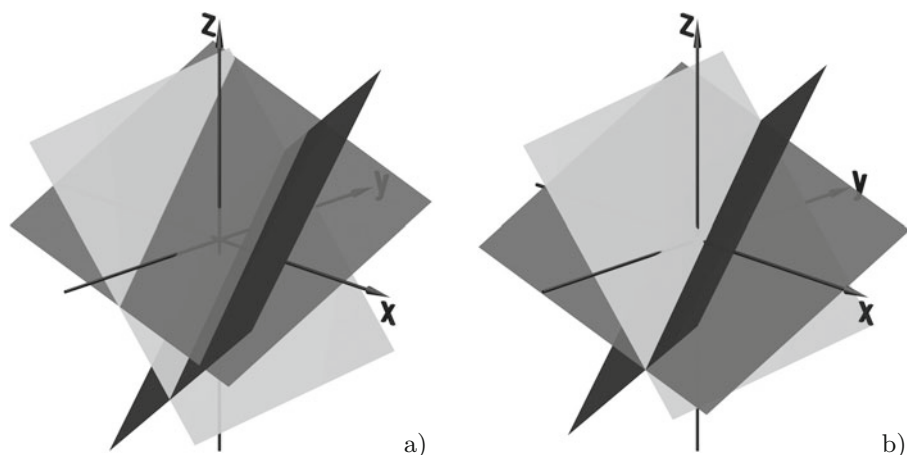


Abb. 1.12: Darstellung von Lösungsmengen linearer Gleichungssysteme mit drei Variablen und drei Gleichungen (Beispiel 1.18)

Der Lösungsmenge von Gleichung I entspricht jeweils eine mittelgraue, der von Gleichung II eine hellgraue und der Lösungsmenge von Gleichung III eine dunkle Ebene.

Subtrahiert man in a) die Gleichung II von I, so ergibt sich $-3x - 3y + 3z = 20$. Nach Gleichung III muss jedoch $-3x - 3y + 3z = -8$. Diese Bedingung kann kein Lösungstriplett $(x; y; z)$ erfüllen, denn dazu müsste $20 = -8$ sein.

In Beispiel b) ergibt sich durch I–II die Gleichung $-3x - 3y + 3z = -2$, also Gleichung III. Damit stellt III keine „neue“ Bedingung an Lösungen des Systems, das nur aus den Gleichungen I und II besteht. ■

Zusammenfassung: Lösungsmengen linearer Gleichungssysteme mit drei Variablen und drei Gleichungen

Lineare Gleichungssysteme der Form

$$\text{I} \quad a_{11}x + a_{12}y + a_{13}z = b_1$$

$$\text{II} \quad a_{21}x + a_{22}y + a_{23}z = b_2$$

$$\text{III} \quad a_{31}x + a_{32}y + a_{33}z = b_3$$

bei denen jede Gleichung für sich betrachtet eine einparametrische Lösungsmenge besitzt, können folgendermaßen geartete Lösungsmengen besitzen:

- Es existiert *genau eine Lösung*. Die drei Ebenen, welche die Lösungsmengen der Gleichungen darstellen, haben genau einen Schnittpunkt.
- Die *Lösungsmenge des LGS ist leer*. Mindestens zwei zu den Lösungsmengen der Gleichungen des LGS gehörenden Ebenen sind parallel oder die drei Schnittgeraden der Ebenen sind parallel (siehe Abb. 1.12 a).
- Das LGS hat eine *unendliche (einparametrische) Lösungsmenge*. Alle drei Ebenen schneiden sich in einer Geraden (siehe Abb. 1.12 b).
- Das LGS hat eine *unendliche (zweiparametrische) Lösungsmenge*. Alle zu den Lösungsmengen der Gleichungen gehörenden Ebenen sind identisch.

1.2.2 Der Gauß-Algorithmus

Ein sehr gebräuchliches Lösungsverfahren für lineare Gleichungssysteme ist das Gaußsche Eliminationsverfahren (kurz: Gauß-Algorithmus). Vor allem kompliziertere LGS werden heute zwar bevorzugt mithilfe des Computers gelöst, vgl. Abschnitt 1.5. Dennoch ist der Gauß-Algorithmus von hoher Bedeutung, denn er ermöglicht nicht nur das Lösen konkreter Gleichungssysteme, sondern führt auch zu Erkenntnissen über Strukturen von Lösungsmengen verschiedener LGS.

Der Gauß-Algorithmus ist gewissermaßen eine Verallgemeinerung des Additionsverfahrens (Abschnitt 1.1.3) für lineare Gleichungssysteme mit beliebig vielen Variablen und Gleichungen. Das Grundprinzip dieses Algorithmus besteht darin, Gleichungen so mit Zahlen zu multiplizieren, dass sich nach Addition der dadurch entstehenden Gleichungen neue Gleichungen ergeben, in denen einzelne Variablen nicht mehr auftreten. Das Ziel ist also die schrittweise Elimination von möglichst vielen Variablen aus möglichst vielen Gleichungen, wodurch gegebene LGS eine Form erhalten, in der ihre Lösungen leicht „abzulesen“ sind.

Die Schritte des Gauß-Algorithmus sind *Äquivalenzumformungen*, also Umformungen, welche die Lösungsmenge eines gegebenen LGS nicht ändern (siehe dazu auch Abschnitt 1.1.3, vor allem S. 11):

- Vertauschen zweier Gleichungen,
- Multiplikation einer Gleichung mit einer reellen Zahl (außer Null),
- Addition einer Gleichung zu einer anderen Gleichung.

Häufig wird statt der Addition zweier Gleichungen auch angegeben:

- Ersetzen einer Gleichung durch die Summe eines Vielfachen (Faktor $\neq 0$) von ihr und eines Vielfachen einer anderen Gleichung.

Diese Äquivalenzumformung fasst zwei der darüber angegebenen Umformungen zusammen: Multiplikation von Gleichungen mit reellen Zahlen (außer Null) und Addition von Gleichungen. Bei der Durchführung des Gauß-Algorithmus ist die Zusammenfassung sinnvoll, um weniger Schritte aufschreiben zu müssen.

Beispiel 1.19

Wir lösen das bereits in dem Beispiel 1.17 betrachtete lineare Gleichungssystem

$$\begin{array}{ll} \text{I} & \frac{3}{2}x - 2y + 4z = 1 \\ \text{II} & 2x - \frac{9}{2}y - 2z = -\frac{1}{2} \\ \text{III} & x + y - 3z = 2 \end{array}$$

Vor der Durchführung des Gauß-Algorithmus empfiehlt es sich, die Gleichungen so zu vervielfachen, dass alle Koeffizienten ganzzahlig werden – dies erleichtert die folgenden Rechnungen erheblich.

$$\begin{array}{llll} \text{I} & 3x - 4y + 8z = 2 & | \cdot (-4) & | \cdot (-1) \\ \text{II} & 4x - 9y - 4z = -1 & | \cdot 3 & \\ \text{III} & x + y - 3z = 2 & & | \cdot 3 \end{array}$$

Im ersten Umformungsschritt des Gauß-Algorithmus wird die Gleichung I übernommen. Die Gleichungen II und III werden derart multipliziert und mit geeigneten Vielfachen der Gleichung I addiert, dass Gleichungen entstehen, in denen die Variable x nicht mehr auftritt. Die Zahlen rechts neben den Gleichungen geben an, mit welchen Zahlen die Gleichungen addiert werden. In den dadurch entstehenden Gleichungen

$$\text{II}' : (-4) \cdot \text{I} + 3 \cdot \text{II} \quad \text{und} \quad \text{III}' : (-1) \cdot \text{I} + 3 \cdot \text{III}$$

tritt x nicht mehr auf:

$$\begin{array}{rclcl} \text{I}' & 3x - 4y + 8z & = & 2 & \\ \text{II}' & -11y - 44z & = & -11 & | \cdot 7 \\ \text{III}' & 7y - 17z & = & 4 & | \cdot 11 \end{array}$$

Im nächsten Schritt werden die ersten beiden Gleichungen beibehalten. Durch geeignete Multiplikation der Gleichungen II' und III' wird eine Gleichung erzeugt, in der auch y nicht mehr vorkommt:

$$\begin{array}{rclcl} \text{I}'' & 3x - 4y + 8z & = & 2 & \\ \text{II}'' & -11y - 44z & = & -11 & \\ \text{III}'' & -495z & = & -33 & \end{array}$$

Das Gleichungssystem befindet sich jetzt in der so genannten Dreiecksform bzw. *Zeilenstufenform*. Der „*einfache Gauß-Algorithmus*“ endet an dieser Stelle. Der aus der Gleichung III'' entnehmbare Wert $z = \frac{1}{15}$ kann dann in II'' eingesetzt werden, womit sich $y = \frac{11}{15}$ ergibt. Durch Einsetzen beider Werte in I'' wird $x = \frac{22}{15}$ ermittelt. Das LGS ist damit gelöst, und es hat sich ergeben, dass es genau eine Lösung besitzt: $(\frac{22}{15}; \frac{11}{15}; \frac{1}{15})$.

Häufig ist es sinnvoll, statt die Lösungen wie beschrieben durch „Einsetzen von unten nach oben“ zu entnehmen, den Gauß-Algorithmus systematisch weiterzuführen. Die folgenden Schritte gehören zu dem erweiterten Gauß-Algorithmus bzw. *Gauß-Jordan-Algorithmus*. Die Gleichungen werden dabei so umgeformt, dass in jeder Gleichung so wenige Variablen wie möglich auftreten. In dem hier betrachteten LGS wird das nur noch eine Variable in jeder Gleichung sein.

Wir vereinfachen vor der Weiterführung des Algorithmus die Gleichungen II'' und III'' und führen das Verfahren mit dem folgenden LGS fort, dessen Gleichungen keine neuen Bezeichnungen erhalten, da lediglich Vereinfachungen *innerhalb* einzelner Gleichungen vorgenommen wurden.

$$\begin{array}{rclcl} \text{I}'' & 3x - 4y + 8z & = & 2 & | \cdot (-15) \\ \text{II}'' & y + 4z & = & 1 & | \cdot (-15) \\ \text{III}'' & 15z & = & 1 & | \cdot 4 \quad | \cdot 8 \end{array}$$

Als Ergebnis der Umformungen erhalten wir:

$$\begin{array}{rclcl} \text{I}''' & -45x + 60y & = & -22 & | \cdot 1 \\ \text{II}''' & -15y & = & -11 & | \cdot 4 \\ \text{III}''' & 15z & = & 1 & \end{array}$$

Mit dem letzten Schritt wird y aus Gleichung I''' eliminiert:

$$\begin{array}{rcl} I'''' & -45x & = -66 \\ II'''' & -15y & = -11 \\ III'''' & 15z & = 1 \end{array}$$

Mit dieser Umformung wurde das LGS in die *Diagonalform* überführt. Der erweiterte Gauß-Algorithmus bzw. Gauß-Jordan-Algorithmus wurde damit vollständig durchgeführt. Die Lösungsmenge des LGS lässt sich direkt aus den Gleichungen I'''' (die sich zu $15x = 22$ vereinfacht), II'''' und III'''' ablesen:

$$L = \left\{ \left(\frac{22}{15}; \frac{11}{15}; \frac{1}{15} \right) \right\}.$$

Eine Zusammenfassung der Umformungsschritte und eine Visualisierung des Gauß-Algorithmus anhand des behandelten Beispiels gibt Abb. 1.13. Der Gauß-Algorithmus verändern die drei Gleichungen schrittweise so, dass zugehörige Ebenen entstehen, die zunächst zu Koordinatenachsen und schließlich sogar zu Koordinatenebenen parallel sind. Sind alle drei durch das Gleichungssystem

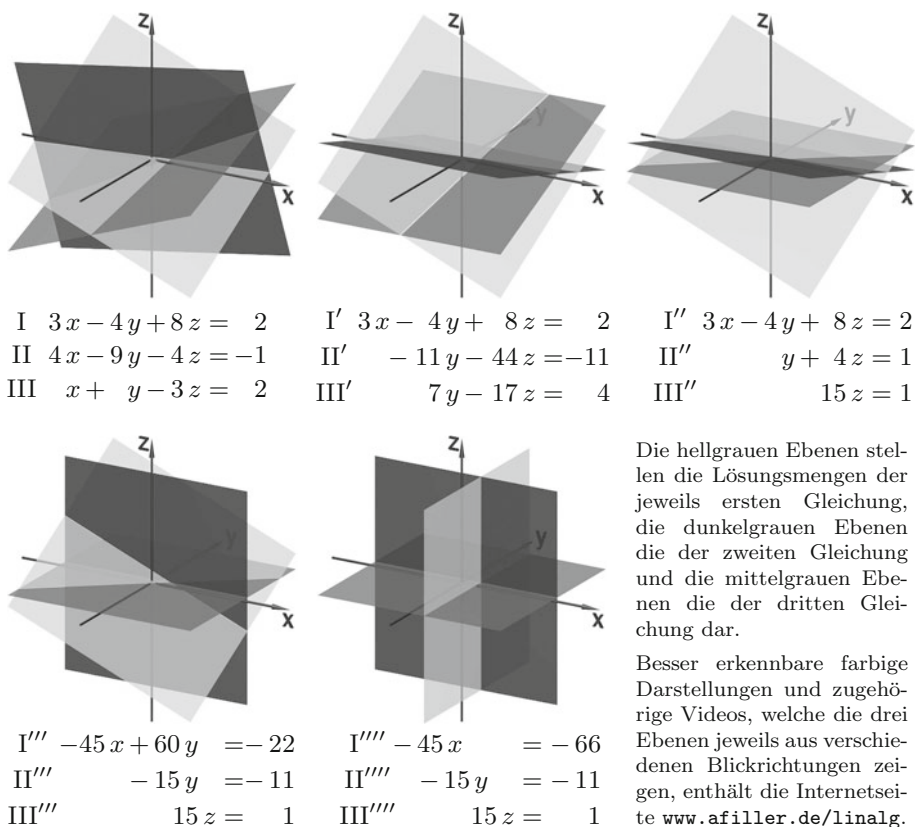


Abb. 1.13: Darstellung der Lösungsmengen der Gleichungen eines LGS mit drei Gleichungen und drei Variablen nach den Lösungsschritten des Gauß-Jordan-Algorithmus

gegebenen Ebenen zu jeweils einer Koordinatenebene parallel, so befindet sich das System in Diagonalform.

Es ist in Abb. 1.13 erkennbar, dass die durch den Gauß-Algorithmus vorgenommen Umformungen den Schnittpunkt der drei Ebenen unverändert lassen, die Ebenen selbst aber verändern. Die Umformungsschritte sind somit Äquivalenzumformungen in Bezug auf das gesamte Gleichungssystem, nicht aber Äquivalenzumformungen in Bezug auf die einzelnen Gleichungen, deren – durch die jeweiligen Ebenen repräsentierten – Lösungsmengen sich erkennbar verändern.

Mitunter kann es bei der Durchführung des Gauß-Algorithmus notwendig sein, Gleichungen zu vertauschen. Wenn in einem LGS der Form

$$\begin{array}{ll} \text{I} & a_{11} x + a_{12} y + a_{13} z = b_1 \\ \text{II} & a_{21} x + a_{22} y + a_{23} z = b_2 \\ \text{III} & a_{31} x + a_{32} y + a_{33} z = b_3 \end{array}$$

der Koeffizient a_{11} den Wert Null hat, so kann der Gauß-Algorithmus nicht in der in Beispiel 1.19 dargestellten Weise durchgeführt werden. Vertauscht man aber die Gleichung I mit einer Gleichung, in welcher der zu x gehörige Koeffizient von Null verschieden ist, so vereinfacht sich der erste Umformungsschritt sogar. Ist hingegen $a_{11}=0$, $a_{21}=0$ und $a_{31}=0$, so tritt x in dem LGS nicht auf und es handelt sich um ein LGS in zwei Variablen.

Mitunter ist die Notwendigkeit, Gleichungen zu vertauschen, anhand des zu lösenden Gleichungssystems noch nicht erkennbar, sondern ergibt sich erst während der Durchführung des Algorithmus.

Beispiel 1.20

Es wird ein Gleichungssystem mithilfe des Gauß-Algorithmus gelöst.

$$\begin{array}{llll} \text{I} & 2x - y - 5z = 4 & | \cdot (-2) & | \cdot (-5) \\ \text{II} & 4x - 2y + 3z = 8 & | \cdot 1 & \\ \text{III} & 5x + 3y - 7z = -1 & & | \cdot 2 \\ \\ \text{I}' & 2x - y - 5z = 4 & & \\ \text{II}' & & 13z = 0 & \\ \text{III}' & & 11y + 11z = -22 & \end{array}$$

Da in der Gleichung II' der Koeffizient vor y verschwunden ist, kann der Gauß-Algorithmus nicht wie in dem Beispiel 1.19 weitergeführt werden. Durch Vertauschung der Gleichungen II' und III' erhält das LGS aber bereits die Zeilenstufenform:

$$\begin{array}{llll} \text{I}' & 2x - y - 5z = 4 & & \\ \text{II}' & & 11y + 11z = -22 & \\ \text{III}' & & 13z = 0 & \end{array}$$

Durch Weiterführung des Gauß-Jordan-Algorithmus ließe sich dieses LGS nun in die Diagonalform bringen. Es ist aber auch möglich, die Lösung dieses (recht einfachen) LGS von unten nach oben aus den Gleichungen „abzulesen“:

$z = 0$, $y = -2$, $x = 1$. ■

Der Gauß-Algorithmus lässt sich nicht nur auf lineare Gleichungssysteme anwenden, die eindeutig bestimmte Lösungen besitzen, sondern auch auf nicht lösbare LGS und LGS mit unendlichen Lösungsmengen.

Beispiel 1.21

Wir betrachten das bereits in dem Beispiel 1.18 als *nicht lösbar* erkannte LGS a (siehe Abb. 1.12 a auf S. 21), und wenden darauf den Gauß-Algorithmus an:

$$\begin{array}{llll}
 \text{I} & 3x - 4y + 8z = 16 & | \cdot (-2) & | \cdot 1 \\
 \text{II} & 6x - y + 5z = -4 & | \cdot 1 & \\
 \text{III} & -3x - 3y + 3z = -8 & & | \cdot 1 \\
 \\
 \text{I}' & 3x - 4y + 8z = 16 & & \\
 \text{II}' & 7y - 11z = -36 & | \cdot 1 & \\
 \text{III}' & -7y + 11z = 8 & | \cdot 1 & \\
 \\
 \text{I}'' & 3x - 4y + 8z = 16 & & \\
 \text{II}'' & 7y - 11z = -36 & & \\
 \text{III}'' & 0 = -28 & &
 \end{array}$$

Der Gauß-Algorithmus hat einen Widerspruch innerhalb des Gleichungssystems zu Tage gefördert, das LGS besitzt keine Lösungen. ■

Beispiel 1.22

Der Gauß-Algorithmus wird nun auf das in dem Beispiel 1.18 b betrachtete LGS mit *unendlicher Lösungsmenge* angewendet (siehe dazu auch Abb. 1.12 b).

$$\begin{array}{llll}
 \text{I} & 3x - 4y + 8z = 2 & | \cdot (-2) & | \cdot 1 \\
 \text{II} & 6x - y + 5z = 4 & | \cdot 1 & \\
 \text{III} & -3x - 3y + 3z = -2 & & | \cdot 1 \\
 \\
 \text{I}' & 3x - 4y + 8z = 2 & & \\
 \text{II}' & 7y - 11z = 0 & | \cdot 1 & \\
 \text{III}' & -7y + 11z = 0 & | \cdot 1 & \\
 \\
 \text{I}'' & 3x - 4y + 8z = 2 & & \\
 \text{II}'' & 7y - 11z = 0 & & \\
 \text{III}'' & 0 = 0 & &
 \end{array}$$

Durch die Schritte des Gauß-Algorithmus ist die dritte Gleichung gewissermaßen „verschwunden“. Das LGS enthält nicht mehr Informationen als ein LGS mit zwei Gleichungen, eine der drei Gleichungen des ursprünglichen LGS war „überflüssig“. Man sagt deshalb, dass der *Rang* des hier betrachteten linearen Gleichungssystems zwei ist.

Wir führen mit den verbliebenen zwei Gleichungen den Gauß-Jordan-Algorithmus fort und eliminieren y aus der Gleichung I'' :

$$\begin{array}{llll}
 \text{I}'' & 3x - 4y + 8z = 2 & | \cdot 7 & \\
 \text{II}'' & 7y - 11z = 0 & | \cdot 4 & \\
 \\
 \text{I}''' & 21x & + 12z = 14 & \\
 \text{II}''' & 7y - 11z = 0 & &
 \end{array}$$

Dieses lineare Gleichungssystem kann nicht weiter vereinfacht werden. Um die Lösungsmenge anzugeben, wird die in beiden Gleichungen noch auftretende Variable z als Parameter t gesetzt. Es ist dann:

$$\begin{aligned}x &= \frac{-12t+14}{21} = -\frac{4}{7}t + \frac{2}{3} \\y &= \frac{11}{7}t \\z &= t.\end{aligned}$$

Die Lösungsmenge des LGS ist somit einparametrig und lässt sich folgendermaßen angeben:

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mid \begin{pmatrix} x \\ y \\ z \end{pmatrix} = t \cdot \begin{pmatrix} -\frac{4}{7} \\ \frac{11}{7} \\ 1 \end{pmatrix} + \begin{pmatrix} \frac{2}{3} \\ 0 \\ 0 \end{pmatrix} ; t \in \mathbb{R} \right\}.$$

(Zur Darstellung von Lösungsmengen mit Parametern in Spaltenschreibweise siehe die Abschnitte 1.1.1 und 1.2.1.) ■

Der Gauß-Algorithmus lässt sich auch auf LGS mit drei Variablen anwenden, die von vornherein nur aus zwei Gleichungen bestehen.

Beispiel 1.23

Wir lösen nun das LGS a aus dem Beispiel 1.16 auf S. 19, multiplizieren aber zuvor die erste Gleichung mit 2, damit keine Brüche mehr auftreten.

$$\begin{array}{llll} \text{I} & 3x - 4y + 8z = 2 & | \cdot (-1) \\ \text{II} & x + y - 3z = 2 & | \cdot 3 \\ \\ \text{I}' & 3x - 4y + 8z = 2 & | \cdot 7 \\ \text{II}' & 7y - 17z = 4 & | \cdot 4 \\ \\ \text{I}'' & 21x - 12z = 30 \\ \text{II}'' & 7y - 17z = 4 \end{array}$$

Eine weitere Elimination von Variablen ist nicht mehr möglich. Die noch in beiden Gleichungen auftretende Variable z wird als Parameter t gesetzt.

$$\begin{aligned}x &= \frac{12t+30}{21} = \frac{4}{7}t + \frac{10}{7} \\y &= \frac{17t+4}{7} = \frac{17}{7}t + \frac{4}{7} \\z &= t\end{aligned}$$

Daraus ergibt sich die folgende Lösungsmenge:

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mid \begin{pmatrix} x \\ y \\ z \end{pmatrix} = t \cdot \begin{pmatrix} \frac{4}{7} \\ \frac{17}{7} \\ 1 \end{pmatrix} + \begin{pmatrix} \frac{10}{7} \\ \frac{4}{7} \\ 0 \end{pmatrix} ; t \in \mathbb{R} \right\}.$$

Die Lösungsmenge ist also einparametrig und wird geometrisch durch eine Gerade repräsentiert, siehe Abb. 1.10 a auf S. 19. ■

Die Beispiele 1.22 und 1.23 unterscheiden sich qualitativ nicht, beide Lösungsmengen sind einparametrig. Bedeutend für die Gestalt der Lösungsmenge eines LGS ist also nicht die Anzahl der gegebenen Gleichungen, sondern die Anzahl der bei der Durchführung des Gauß-Algorithmus nicht „verschwindenden“ Gleichungen – der *Rang* des LGS. Dieser beträgt in beiden Beispielen zwei.

1.2.3 Aufgaben zu Abschnitt 1.2

1. Geben Sie die Lösungsmengen der folgenden Gleichungssysteme (vgl. Beispiel 1.16 auf S. 19) an.

$$\begin{array}{l} \text{a)} \quad -3x - 3y + 9z = 30 \\ \quad \quad x + y - 3z = 2 \end{array}$$

$$\begin{array}{l} \text{b)} \quad -3x - 3y + 9z = -6 \\ \quad \quad x + y - 3z = 2 \end{array}$$

2. Lösen Sie die folgenden linearen Gleichungssysteme.

$$\begin{array}{l} \text{a)} \quad 2x + 3y - z = 1 \\ \quad \quad x + 3y + z = 2 \\ \quad \quad -2x - 2y + 4z = 4 \end{array}$$

$$\begin{array}{l} \text{b)} \quad 3x + y + 2z = -1 \\ \quad \quad 9x - 5y - 3z = 5 \\ \quad \quad 2x - 3y - 2z = 0 \end{array}$$

3. Bestimmen Sie die Lösungsmengen der folgenden LGS und interpretieren Sie diese geometrisch.

$$\begin{array}{l} \text{a)} \quad 2x + 3y = 5 \\ \quad \quad x + y = 2 \\ \quad \quad 3x + y = 1 \end{array}$$

$$\begin{array}{l} \text{b)} \quad 2x - y + 2z = 1 \\ \quad \quad x - 2y + 3z = 1 \\ \quad \quad 6x + 3y - 2z = 1 \\ \quad \quad x - 5y + 7z = 2 \end{array}$$

4. Bestimmen Sie die Lösungsmenge des folgenden linearen Gleichungssystems in Abhängigkeit von dem Parameter $r \in \mathbb{R}$:

$$\begin{array}{l} rx + y + z = 1 \\ x + ry + z = 1 \\ x + y + rz = 1 \end{array}$$

Hinweis: Nehmen Sie eine Fallunterscheidung vor. Überlegen Sie dazu, für welche(s) r das LGS nicht lösbar ist und für welche(s) r es unendlich viele Lösungen besitzt.

5. Geben Sie ein lineares Gleichungssystem mit 3 Gleichungen und drei Variablen an, das nicht lösbar ist, weil zwei der Gleichungen durch zueinander parallele Ebenen dargestellt werden.

Überlegen Sie dazu, welcher Zusammenhang zwischen zwei Gleichungen bestehen muss, damit die zugehörigen Ebenen parallel sind.

6. Die beiden folgenden Aufgaben stammen aus einem der bekanntesten Mathematikbücher aller Zeiten, der 1767/68 erschienenen „Vollständigen Anleitung zur Algebra“ von LEONHARD EULER.

- a) Einer kauft 12 Stück Tuch für 140 Reichsthaler, und zwar 2 weiße, 3 schwarze und 7 blaue. Ein Stück schwarzes Tuch kostet 2 Reichsthaler mehr als ein weißes, und ein blaues 3 Reichsthaler mehr als ein schwarzes. Nun ist die Frage, wie viel jedes gekostet?

- b) Drei haben ein Haus gekauft für 100 Reichsthaler; der erste begehrt vom andern $\frac{1}{2}$ seines Geldes, so könnte er das Haus allein bezahlen; der andere begehrt vom dritten $\frac{1}{3}$ seines Geldes, so könnte er das Haus allein bezahlen; der dritte begehrt vom ersten $\frac{1}{4}$ seines Geldes, so möchte er das Haus allein bezahlen. Wie viel hat jeder Geld gehabt?

1.3 Verallgemeinerungen und einige Begriffe

1.3.1 Der Gauß-Algorithmus in Matrixschreibweise

Der Gauß-Algorithmus lässt sich für lineare Gleichungssysteme mit beliebig vielen Gleichungen und beliebig vielen Variablen durchführen. Treten mehr als drei Variablen auf, so werden diese meist $x_1, x_2, x_3, x_4, \dots$ genannt. Um Schreibarbeit zu sparen und eine übersichtliche Darstellung zu erreichen, wird gerade bei größeren LGS die Matrixschreibweise bevorzugt. Dabei werden lediglich die Koeffizienten und die absoluten Glieder eines Gleichungssystems in einer Tabelle (Matrix) notiert.

Beispiel 1.24

Dem links angegebenen linearen Gleichungssystem entspricht die rechts dargestellte Matrixschreibweise:

$$\begin{array}{rrcr} 2x_1 + 7x_2 + 9x_3 + x_4 = & 1 \\ 4x_1 + 14x_2 + 8x_3 + 3x_4 = & 6 \\ x_1 + 3x_2 + 5x_3 - 3x_4 = & -13 \\ 10x_1 + 5x_2 - x_3 - 4x_4 = & -1 \end{array} \quad \left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 4 & 14 & 8 & 3 & 6 \\ 1 & 3 & 5 & -3 & -13 \\ 10 & 5 & -1 & -4 & -1 \end{array} \right)$$

Diese Schreibweise wird nun bei der Lösung des LGS verwendet.

$$\left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 4 & 14 & 8 & 3 & 6 \\ 1 & 3 & 5 & -3 & -13 \\ 10 & 5 & -1 & -4 & -1 \end{array} \right) \begin{array}{l} | \cdot (-2) \quad | \cdot (-1) \quad | \cdot (-5) \\ | \cdot 1 \\ | \cdot 2 \\ | \cdot 1 \end{array}$$

$$\left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 0 & 0 & -10 & 1 & 4 \\ 0 & -1 & 1 & -7 & -27 \\ 0 & -30 & -46 & -9 & -6 \end{array} \right)$$

Für die Fortsetzung des Gauß-Algorithmus werden die zweite und die vierte Zeile vertauscht, dies entspricht der Vertauschung zweier Gleichungen.

$$\left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 0 & -30 & -46 & -9 & -6 \\ 0 & -1 & 1 & -7 & -27 \\ 0 & 0 & -10 & 1 & 4 \end{array} \right) \begin{array}{l} | \cdot (-1) \\ | \cdot 30 \end{array}$$

$$\left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 0 & -30 & -46 & -9 & -6 \\ 0 & 0 & 76 & -201 & -804 \\ 0 & 0 & -10 & 1 & 4 \end{array} \right) \begin{array}{l} | \cdot 5 \\ | \cdot 38 \end{array}$$

$$\left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 0 & -30 & -46 & -9 & -6 \\ 0 & 0 & 76 & -201 & -804 \\ 0 & 0 & 0 & -967 & -3868 \end{array} \right)$$

Die Lösung könnte nun von unten beginnend durch Einsetzen in die Gleichungen entnommen werden. Wir führen hier jedoch den Gauß-Jordan-Algorithmus vollständig durch und bringen das LGS bzw. die Matrix in die Diagonalform.

Vorher empfiehlt es sich, die Matrix, wenn möglich, zu vereinfachen (im vorliegenden Beispiel durch Multiplikation der vierten Zeile mit $\frac{1}{967}$).

$$\begin{aligned}
 & \left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 1 \\ 0 & -30 & -46 & -9 & -6 \\ 0 & 0 & 76 & -201 & -804 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right) \begin{array}{l} | \cdot 1 \\ | \cdot 1 \\ | \cdot 201 \quad | \cdot 9 \quad | \cdot (-1) \end{array} \\
 & \left(\begin{array}{cccc|c} 2 & 7 & 9 & 0 & -3 \\ 0 & -30 & -46 & 0 & 30 \\ 0 & 0 & 76 & 0 & 0 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right) \begin{array}{l} | \cdot 76 \\ | \cdot 76 \\ | \cdot 46 \quad | \cdot (-9) \end{array} \\
 & \left(\begin{array}{cccc|c} 2 & 7 & 0 & 0 & -3 \\ 0 & -2280 & 0 & 0 & 2280 \\ 0 & 0 & 76 & 0 & 0 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right) \begin{array}{l} \text{wird ver-} \\ \text{einfacht zu} \end{array} \left(\begin{array}{cccc|c} 2 & 7 & 0 & 0 & -3 \\ 0 & 1 & 0 & 0 & -1 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right) \begin{array}{l} | \cdot 1 \\ | \cdot (-7) \end{array} \\
 & \left(\begin{array}{cccc|c} 2 & 0 & 0 & 0 & 4 \\ 0 & 1 & 0 & 0 & -1 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right)
 \end{aligned}$$

Aus diesem Ergebnis lassen sich die Werte der Variablen direkt ablesen:

$$\begin{aligned}
 2x_1 &= 4 \\
 x_2 &= -1 \\
 x_3 &= 0 \\
 x_4 &= 4
 \end{aligned}$$

Das gegebene LGS ist also eindeutig lösbar und hat die Lösungsmenge

$$L = \{(2; -1; 0; 4)\}.$$

Einfache und erweiterte Koeffizientenmatrix

Allgemein lässt sich ein beliebiges lineares Gleichungssystem mit n Variablen und m Gleichungen in der Form

$$\begin{aligned}
 a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= b_1 \\
 a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= b_2 \\
 \vdots & \quad \quad \quad \vdots \\
 a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= b_m
 \end{aligned} \tag{1.15}$$

oder durch die Matrix

$$\left(\begin{array}{cccc|c} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & b_m \end{array} \right) \tag{1.16}$$

darstellen. Die Matrix (1.16) nennt man *erweiterte Koeffizientenmatrix* des LGS (1.15). Derjenige Teil dieser Matrix, der nur die Koeffizienten $a_{11} \dots a_{mn}$ enthält, heißt *einfache Koeffizientenmatrix*:

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}. \quad (1.17)$$

In Kurzschreibweise lässt sich die erweiterte Koeffizientenmatrix durch die einfache Koeffizientenmatrix $\mathbf{A}$ und den Spaltenvektor der absoluten

Glieder $\vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix}$ angeben:

$$(\mathbf{A} | \vec{b}). \quad (1.18)$$

Einige Merkmale der einfachen und der erweiterten Koeffizientenmatrix erlauben wesentliche Aussagen über die Lösbarkeit und die Struktur der Lösungsmenge eines linearen Gleichungssystems. Diese Merkmale offenbaren sich jedoch erst nach der Überführung dieser Matrizen in die Zeilenstufenform.

Beispiel 1.25

Auf das durch die folgende Matrix gegebene LGS wird der Gauß-Algorithmus angewendet.

$$\left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 1 & -3 & -4 & -1 & -5 \\ 4 & 6 & 11 & 5 & 16 \end{array} \right) \begin{array}{l} | \cdot (-1) \quad | \cdot (-4) \\ \\ | \cdot 1 \\ | \cdot 1 \end{array} \quad \left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & -6 & -9 & -3 & -12 \\ 0 & -6 & -9 & -3 & -12 \end{array} \right) \begin{array}{l} \\ | \cdot 3 \quad | \cdot 3 \\ | \cdot 1 \\ | \cdot 1 \end{array}$$

$$\left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right) \begin{array}{l} \\ \text{bzw. in Gleichungs-} \\ \text{schreibweise} \end{array} \quad \begin{array}{l} x_1 + 3x_2 + 5x_3 + 2x_4 = 7 \\ 2x_2 + 3x_3 + x_4 = 4 \\ 0 = 0 \\ 0 = 0 \end{array}$$

Es sind zwei Zeilen der erweiterten Koeffizientenmatrix „verschwunden“ (wurden in Nullzeilen umgeformt). Um die Lösungsmenge anzugeben, werden zwei der Variablen als Parameter gesetzt ($s = x_3$, $t = x_4$). Damit ist

$$\begin{aligned} x_1 &= -\frac{1}{2}s - \frac{1}{2}t + 1 \\ x_2 &= -\frac{3}{2}s - \frac{1}{2}t + 2 \\ x_3 &= s \\ x_4 &= t \end{aligned}$$

und die Lösungsmenge lässt sich folgendermaßen angeben:

$$L = \left\{ \left(\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \right) \middle| \left(\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \right) = s \cdot \begin{pmatrix} -\frac{1}{2} \\ -\frac{3}{2} \\ 1 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} -\frac{1}{2} \\ -\frac{1}{2} \\ 0 \\ 1 \end{pmatrix} + \begin{pmatrix} 1 \\ 2 \\ 0 \\ 0 \end{pmatrix} ; s, t \in \mathbb{R} \right\}. \quad \blacksquare$$

Beispiel 1.26

Das folgende Beispiel unterscheidet sich von dem Beispiel 1.25 auf den ersten Blick kaum, der Gauß-Algorithmus fördert jedoch erhebliche Unterschiede zu Tage.

$$\begin{array}{ccc}
 \left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 1 & -3 & -4 & -1 & -6 \\ 4 & 6 & 11 & 5 & 16 \end{array} \right) & \begin{array}{l} | \cdot (-1) \quad | \cdot (-4) \\ \\ | \cdot 1 \\ \\ | \cdot 1 \end{array} & \left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & -6 & -9 & -3 & -13 \\ 0 & -6 & -9 & -3 & -12 \end{array} \right) \begin{array}{l} \\ | \cdot 3 \quad | \cdot 3 \\ | \cdot 1 \\ | \cdot 1 \end{array} \\
 \\
 \left(\begin{array}{cccc|c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right) & \text{bzw. in Gleichungs-} & \begin{array}{l} x_1 + 3x_2 + 5x_3 + 2x_4 = 7 \\ 2x_2 + 3x_3 + x_4 = 4 \\ -1 = 0 \\ 0 = 0 \end{array} \\
 & \text{schreibweise} &
 \end{array}$$

Das lineare Gleichungssystem ist nicht lösbar. In der Matrixschreibweise drückt sich dies dadurch aus, dass die einfache Koeffizientenmatrix eine „echte Zeile“ (damit ist eine Zeile gemeint, die nicht nur Nullen enthält) weniger behalten hat als die erweiterte Koeffizientenmatrix. ■

1.3.2 Ränge der einfachen und erweiterten Koeffizientenmatrix

Die Beispiele des vorangegangenen Abschnitts legen die Vermutung nahe, dass die Anzahl der nach Durchführung des Gauß-Algorithmus verbleibenden „echten“ Zeilen (Zeilen, die keine Nullzeilen sind) der einfachen und der erweiterten Koeffizientenmatrix eines linearen Gleichungssystems Aufschluss über die Lösbarkeit und die Qualität seiner Lösungsmenge geben. Um hierzu klare Aussagen zu machen, wird der Begriff des *Ranges* eingeführt. Wir betrachten die erweiterte Koeffizientenmatrix:

$$(\mathbf{A} | \vec{b}) = \left(\begin{array}{cccc|c} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & b_m \end{array} \right).$$

Nach der Durchführung des Gauß-Algorithmus nimmt diese Matrix allgemein folgende Gestalt an (Zeilenstufenform):

$$\begin{array}{ll}
 \begin{array}{l} 1. \text{ Zeile} \\ 2. \text{ Zeile} \\ \vdots \\ r. \text{ Zeile} \\ (r+1). \text{ Zeile} \\ \vdots \\ m. \text{ Zeile} \end{array} & \left(\begin{array}{cccc|cccc} a_{11}^{(z)} & \dots & \dots & a_{1r}^{(z)} & a_{1r+1}^{(z)} & \dots & a_{1n}^{(z)} & b_1^{(z)} \\ 0 & a_{22}^{(z)} & \dots & a_{2r}^{(z)} & a_{2r+1}^{(z)} & \dots & a_{2n}^{(z)} & b_2^{(z)} \\ \vdots & \ddots & \ddots & \vdots & \vdots & & \vdots & \vdots \\ 0 & \dots & 0 & a_{rr}^{(z)} & a_{rr+1}^{(z)} & \dots & a_{rn}^{(z)} & b_r^{(z)} \\ 0 & \dots & 0 & 0 & 0 & \dots & 0 & b_{r+1}^{(z)} \\ 0 & \dots & 0 & 0 & 0 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & \dots & 0 & 0 & 0 & \dots & 0 & 0 \end{array} \right) \quad (1.19)
 \end{array}$$

Bemerkungen zu (1.19):

- Die Koeffizienten haben sich gegenüber der Ausgangsmatrix i. Allg. verändert und werden in (1.19) deshalb mit $a_{11}^{(z)}, \dots$ bezeichnet.
- Nicht in jedem Falle treten in der Zeilenstufenform der einfachen Koeffizientenmatrix Nullzeilen auf. Im Falle $r = m$ entspricht die r -te Zeile der m -ten Zeile der Ausgangsmatrix, Nullzeilen existieren dann nicht.

- Die Koeffizienten $a_{11}^{(z)}, a_{22}^{(z)}, \dots, a_{rr}^{(z)}$ müssen alle von Null verschieden sein, da (1.19) durch Anwendung des Gauß-Algorithmus entstanden ist. Hingegen können die weiteren Koeffizienten $a_{ij}^{(z)}$ (mit $j \neq i$) sowie $b_1^{(z)}, \dots, b_{r+1}^{(z)}$ beliebige Werte haben, einschließlich Null.
- Bei der Durchführung des Gauß-Algorithmus ist die genaue Reihenfolge der Umformungsschritte nicht zwingend vorgegeben. So lassen sich z. B. Zeilen vertauschen, um Rechenerleichterungen zu erreichen, auch wenn dies im Sinne der korrekten Durchführung des Algorithmus nicht notwendig ist. Somit können aus demselben Gleichungssystem bzw. aus derselben erweiterten Koeffizientenmatrix nach Durchführung des Gauß-Algorithmus *unterschiedliche Koeffizienten in der Zeilenstufenform* (1.19) auftreten.
- Da alle Umformungsschritte des Gauß-Algorithmus Äquivalenzumformungen sind, müssen alle möglichen Matrizen in Zeilenstufenform, die aus demselben Gleichungssystem entstehen können, dieselbe Lösungsmenge besitzen. Insbesondere folgt daraus dass die *Anzahl r der Zeilen der einfachen Koeffizientenmatrix, in denen nicht nur Nullen stehen, nur von der Ausgangsmatrix $(\mathbf{A} | \vec{b})$ abhängt*. Außerdem hängt es nur von der Ausgangsmatrix ab, ob in der erweiterten Koeffizientenmatrix $b_{r+1}^{(z)} = 0$ oder $b_{r+1}^{(z)} \neq 0$ ist. Diese beiden Eigenschaften sind somit unabhängig von den konkret vorgenommenen Umformungen bei der Durchführung des Gauß-Algorithmus.

Die letzte Bemerkung rechtfertigt die folgenden Definitionen.

Rang der einfachen und Rang der erweiterten Koeffizientenmatrix

Es sei ein lineares Gleichungssystem der Form (1.15) mit einer erweiterten Koeffizientenmatrix $(\mathbf{A} | \vec{b})$ der Form (1.16) gegeben.

- Die Anzahl r der nach Überführung in die Zeilenstufenform (1.19) verbliebenen Zeilen der einfachen Koeffizientenmatrix $\mathbf{A}$, die von Null verschiedene Glieder enthalten, heißt *Rang der einfachen Koeffizientenmatrix*: $\text{rg } \mathbf{A}$.
- Die Anzahl der in der Zeilenstufenform (1.19) verbliebenen Zeilen der erweiterten Koeffizientenmatrix, die von Null verschiedene Glieder enthalten, heißt *Rang der erweiterten Koeffizientenmatrix*: $\text{rg } (\mathbf{A} | \vec{b})$.

Bemerkungen:

- Der Rang der erweiterten Koeffizientenmatrix ist stets gleich dem Rang r der einfachen Koeffizientenmatrix oder $r + 1$ – in Abhängigkeit davon, ob $b_{r+1}^{(z)} = 0$ oder $b_{r+1}^{(z)} \neq 0$ ist.
- In den Beispielen 1.22 und 1.23 auf S. 26f. wurde bereits von dem *Rang eines linearen Gleichungssystems* gesprochen. In beiden Fällen waren – wie eine

Übertragung der dort betrachteten LGS in die Matrizenschreibweise zeigt – die Ränge der einfachen und der erweiterten Koeffizientenmatrix gleich. Falls dies der Fall ist, so stimmen der Rang der einfachen und der Rang der erweiterten Koeffizientenmatrix mit der Anzahl der nach der Durchführung des Gauß-Algorithmus verbleibenden Gleichungen eines LGS überein. Man kann in diesem Falle auch einfach vom Rang eines LGS sprechen.

- Der Begriff „Rang einer Matrix“ wird nach der Behandlung grundlegender Elemente der Vektorrechnung (in den Kapiteln 3 und 5) dann in dem Kapitel 6 noch allgemeiner und eleganter definiert werden; die hier betrachtete Bedeutung des Begriffes bleibt dabei aber erhalten.

Als Beispiele für die allgemeine Darstellung (1.19) betrachten wir nochmals die in den Beispielen 1.24, 1.25 und 1.26 aufgetretenen erweiterten Koeffizientenmatrizen nach Überführung in die Zeilenstufenform und lesen die Ränge ab.

	Beispiel 1.24	Beispiel 1.25	Beispiel 1.26
$(\mathbf{A} \vec{b})$	$\left(\begin{array}{cccc c} 2 & 7 & 9 & 1 & 1 \\ 0 & -30 & -46 & -9 & -6 \\ 0 & 0 & 76 & -201 & -804 \\ 0 & 0 & 0 & 1 & 4 \end{array} \right)$	$\left(\begin{array}{cccc c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right)$	$\left(\begin{array}{cccc c} 1 & 3 & 5 & 2 & 7 \\ 0 & 2 & 3 & 1 & 4 \\ 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 & 0 \end{array} \right)$
$\text{rg } \mathbf{A}$	4	2	2
$\text{rg } (\mathbf{A} \vec{b})$	4	2	3

1.3.3 Ein Lösbarkeitskriterium für lineare Gleichungssysteme

Anhand der Ränge der einfachen und der erweiterten Koeffizientenmatrix lässt sich eine Aussage über die Lösbarkeit linearer Gleichungssysteme formulieren:

Lösbarkeitskriterium für lineare Gleichungssysteme

Ein lineares Gleichungssystem ist genau dann lösbar, falls der Rang seiner einfachen Koeffizientenmatrix mit dem Rang der erweiterten Koeffizientenmatrix übereinstimmt.

Das obige Lösbarkeitskriterium heißt auch *Satz von Kronecker-Capelli*. Es beinhaltet als „genau-dann-wenn-Aussage“ zwei Teilaussagen:

- Wenn* für ein lineares Gleichungssystem der Rang der einfachen mit dem der erweiterten Koeffizientenmatrix übereinstimmt, *dann* ist das LGS lösbar.
- Wenn* ein LGS lösbar ist, *dann* muss der Rang seiner einfachen mit dem der erweiterten Koeffizientenmatrix übereinstimmen. Mit anderen Worten: Wenn der Rang der einfachen *nicht* mit dem der erweiterten Koeffizientenmatrix übereinstimmt, so ist das LGS *nicht* lösbar.

Die zweite Teilaussage lässt sich leicht begründen. Stimmen nämlich die Ränge der erweiterten und der einfachen Koeffizientenmatrix nicht überein, so ist der Rang der erweiterten Koeffizientenmatrix größer als r . Dann muss die erweiterte Koeffizientenmatrix in der Zeilenstufenform (1.19) $r+1$ Zeilen haben, die keine Nullzeilen sind, somit ist $b_{r+1}^{(z)} \neq 0$. Dies bedeutet jedoch, dass das Gleichungssystem in der Zeilenstufenform eine Gleichung der Form $0 = b_{r+1}^{(z)}$ mit $b_{r+1}^{(z)} \neq 0$ enthält – dies ist ein Widerspruch, ein solches LGS ist also nicht lösbar.

Zur Begründung der Teilaussage (i) überlegt man sich, dass die Matrix (1.19) durch Weiterführung des Gauß-Jordan-Algorithmus in die folgende Form überführt werden kann:

$$\left(\begin{array}{cccc|cccc} a_{11}^{(d)} & 0 & \dots & \dots & 0 & a_{1\ r+1}^{(d)} & \dots & a_{1n}^{(d)} & b_1^{(d)} \\ 0 & a_{22}^{(d)} & \ddots & & \vdots & a_{2\ r+1}^{(d)} & \dots & a_{2n}^{(d)} & b_2^{(d)} \\ \vdots & \ddots & \ddots & \ddots & \vdots & \vdots & & \vdots & \vdots \\ \vdots & & \ddots & \ddots & 0 & \vdots & & \vdots & \vdots \\ 0 & \dots & \dots & 0 & a_{rr}^{(d)} & a_{r\ r+1}^{(d)} & \dots & a_{rn}^{(d)} & b_r^{(d)} \end{array} \right) \quad (1.20)$$

(Die $r+1$ -te bis m -te Zeile in 1.19 wurden weggelassen, da wegen der Voraussetzung $b_{r+1}^{(z)} = 0$ ist, es sich also um Nullzeilen handelt.)

Das durch (1.20) gegebene lineare Gleichungssystem ist lösbar. Um seine Lösungsmenge anzugeben, müssen die Variablen $x_{r+1}, \dots, x_n$ als Parameter gesetzt werden: $t_1 := x_{r+1}$, $t_2 := x_{r+2}, \dots$, $t_{n-r} := x_n$. Die Lösungsmenge des LGS lässt sich dann angeben durch:

$$\begin{aligned} x_1 &= -t_1 \frac{a_{1\ r+1}^{(d)}}{a_{11}^{(d)}} - t_2 \frac{a_{1\ r+2}^{(d)}}{a_{11}^{(d)}} - \dots - t_{n-r} \frac{a_{1n}^{(d)}}{a_{11}^{(d)}} + \frac{b_1^{(d)}}{a_{11}^{(d)}} \\ x_2 &= -t_1 \frac{a_{2\ r+1}^{(d)}}{a_{22}^{(d)}} - t_2 \frac{a_{2\ r+2}^{(d)}}{a_{22}^{(d)}} - \dots - t_{n-r} \frac{a_{2n}^{(d)}}{a_{22}^{(d)}} + \frac{b_2^{(d)}}{a_{22}^{(d)}} \\ &\vdots \\ x_r &= -t_1 \frac{a_{r\ r+1}^{(d)}}{a_{rr}^{(d)}} - t_2 \frac{a_{r\ r+2}^{(d)}}{a_{rr}^{(d)}} - \dots - t_{n-r} \frac{a_{rn}^{(d)}}{a_{rr}^{(d)}} + \frac{b_r^{(d)}}{a_{rr}^{(d)}} \\ x_{r+1} &= t_1 \\ &\vdots \\ x_n &= t_{n-r} \end{aligned} \quad (1.21)$$

mit $t_1, \dots, t_{n-r} \in \mathbb{R}$. Diese Überlegungen führen zu einer weiteren fundamentalen Aussage über die Lösungsmengen linearer Gleichungssysteme.

Rang eines LGS und Anzahl der Parameter der Lösungsmenge

Ein lösbares lineares Gleichungssystem in n Variablen mit dem Rang r der einfachen und der erweiterten Koeffizientenmatrix besitzt eine $n-r$ -parametrische Lösungsmenge.

Als Spezialfall dieser Aussage ergibt sich, dass ein lineares Gleichungssystem in n Variablen mit dem Rang n der einfachen und der erweiterten Koeffizientenmatrix eine *Lösungsmenge ohne Parameter* besitzt. In diesem Falle hat die erweiterte Koeffizientenmatrix nach der Durchführung des Gauß-Jordan-Algorithmus (1.20) die Gestalt

$$\left(\begin{array}{cccccc|c} a_{11}^{(d)} & 0 & \dots & \dots & 0 & b_1^{(d)} \\ 0 & a_{22}^{(d)} & \ddots & & \vdots & b_2^{(d)} \\ \vdots & & \ddots & \ddots & \vdots & \vdots \\ \vdots & & & \ddots & 0 & \vdots \\ 0 & \dots & \dots & 0 & a_{nn}^{(d)} & b_n^{(d)} \end{array} \right)$$

und die Lösung des LGS ist

$$x_1 = \frac{b_1^{(d)}}{a_{11}^{(d)}}; \quad x_2 = \frac{b_2^{(d)}}{a_{22}^{(d)}}; \dots; \quad x_n = \frac{b_n^{(d)}}{a_{nn}^{(d)}}.$$

Das Gleichungssystem ist für $n = \operatorname{rg}(\mathbf{A} | \vec{b}) = \operatorname{rg} \mathbf{A}$ also *eindeutig lösbar*.

Es empfiehlt sich, das in diesem Abschnitt formulierte Lösbarkeitskriterium und die Strukturaussage über die Anzahl der Parameter der Lösungsmenge eines LGS anhand der zuvor behandelten Beispiele zu überprüfen. Der Leserin bzw. dem Leser sei es überlassen, für die Beispiele in den Abschnitten 1.1, 1.2 und 1.3.1 Koeffizientenmatrizen aufzustellen, deren Ränge zu bestimmen und mit der Anzahl der Parameter in den ermittelten Lösungsmengen zu vergleichen. Aus allen in diesem Kapitel behandelten Beispielen lösbarer Gleichungssysteme (bei denen $\operatorname{rg}(\mathbf{A} | \vec{b}) = \operatorname{rg} \mathbf{A}$ ist und deshalb einfach von den Rängen der entsprechenden LGS gesprochen wird) ergibt sich die folgende Zusammenstellung.

Anzahl der Variablen	Rang des LGS	Dimension der Lösungsmenge (Anzahl der Parameter)
1	1	0
2	2	0
2	1	1
3	3	0
3	2	1
3	1	2
4	4	0
4	2	2

Diese Ergebnisse scheinen den herausgearbeiteten Zusammenhang zwischen dem Rang eines linearen Gleichungssystems und der Anzahl der Parameter seiner Lösungsmenge zu bestätigen. Jedoch ist damit kein allgemein gültiger Beweis erbracht, und auch die zuvor geführten Plausibilitätsbetrachtungen stellen keinen exakten und vollständigen Beweis dar. Der auf S. 35 formulierte Zusammenhang wird deshalb in den Kapiteln 5 und 6 zu einer der zentralen Strukturaussagen der linearen Algebra „ausgebaut“.

1.3.4 Homogene und inhomogene lineare Gleichungssysteme

In diesem Abschnitt werden spezielle LGS untersucht, für die sich besondere Aussagen über Lösungsmengen treffen und recht leicht beweisen lassen. Diese Aussagen führen dann auch zu weiteren strukturellen Erkenntnissen über Lösungsmengen beliebiger linearer Gleichungssysteme.

Definition 1.1

Ein lineares Gleichungssystem heißt *homogen*, falls seine sämtlichen absoluten Glieder Null sind, das LGS also die Gestalt

$$\begin{aligned} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n &= 0 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n &= 0 \\ \vdots & \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n &= 0 \end{aligned} \quad (1.22)$$

besitzt. Anderenfalls nennt man das LGS *inhomogen*. ◆

Die erweiterte Koeffizientenmatrix eines homogenen LGS wird in der Form
 $(\mathbf{A} \mid \vec{0})$
 geschrieben.

Es ergibt sich aus der Definition sofort die folgende Folgerung:

Jedes homogene lineare Gleichungssystem ist lösbar.

Diese Folgerung lässt sich sehr leicht beweisen: $(0; 0; \dots; 0)$ ist eine Lösung des LGS (1.22), denn setzt man in (1.22) $x_1 = 0, x_2 = 0, \dots, x_n = 0$ ein, so sind die linken ebenso wie die rechten Seiten aller Gleichungen Null.

Der folgende Satz enthält eine Aussage über die Struktur der Lösungsmenge eines beliebigen homogenen LGS, die in den folgenden Kapiteln noch von hoher Bedeutung sein wird.

Satz 1.1

Sind $\mathbf{x}^{(1)} = (x_1^{(1)}; x_2^{(1)}; \dots; x_n^{(1)})$ und $\mathbf{x}^{(2)} = (x_1^{(2)}; x_2^{(2)}; \dots; x_n^{(2)})$ Lösungen eines homogenen LGS und c eine beliebige reelle Zahl, so sind auch

- $\mathbf{x}^{(1)} + \mathbf{x}^{(2)} = (x_1^{(1)} + x_1^{(2)}; x_2^{(1)} + x_2^{(2)}; \dots; x_n^{(1)} + x_n^{(2)})$ sowie
- $c \cdot \mathbf{x}^{(1)} = (c \cdot x_1^{(1)}; c \cdot x_2^{(1)}; \dots; c \cdot x_n^{(1)})$

Lösungen dieses LGS.

Der Satz sagt also aus, dass die Summe zweier Lösungen eines homogenen LGS wiederum eine Lösung dieses LGS ist, und dass auch ein beliebiges reelles Vielfache einer Lösung eines homogenen LGS eine Lösung ist.

Beweis: Sind $\mathbf{x}^{(1)} = (x_1^{(1)}; x_2^{(1)}; \dots; x_n^{(1)})$ und $\mathbf{x}^{(2)} = (x_1^{(2)}; x_2^{(2)}; \dots; x_n^{(2)})$ Lösungen eines homogenen LGS (1.22) so gilt für jede Gleichung dieses LGS:

$$\begin{aligned} a_{i1}x_1^{(1)} + a_{i2}x_2^{(1)} + \dots + a_{in}x_n^{(1)} &= 0 \quad \text{und} \\ a_{i1}x_1^{(2)} + a_{i2}x_2^{(2)} + \dots + a_{in}x_n^{(2)} &= 0 \quad \text{für alle } i \in \mathbb{N} \text{ mit } 1 \leq i \leq n. \end{aligned}$$

Addition dieser beiden Gleichungen ergibt:

$$a_{i1}x_1^{(1)} + a_{i2}x_2^{(1)} + \dots + a_{in}x_n^{(1)} + a_{i1}x_1^{(2)} + a_{i2}x_2^{(2)} + \dots + a_{in}x_n^{(2)} = 0.$$

Durch Umstellen dieser Gleichung erhält man

$$a_{i1} \left(x_1^{(1)} + x_1^{(2)} \right) + a_{i2} \left(x_2^{(1)} + x_2^{(2)} \right) + \dots + a_{in} \left(x_n^{(1)} + x_n^{(2)} \right) = 0.$$

Damit ist für jede beliebige (i -te) Gleichung des gegebenen LGS auch $\mathbf{x}^{(1)} + \mathbf{x}^{(2)}$ eine Lösung, die damit eine Lösung des gesamten LGS ist.

Analog dazu folgt aus der Voraussetzung

$$c \cdot \left(a_{i1}x_1^{(1)} + a_{i2}x_2^{(1)} + \dots + a_{in}x_n^{(1)} \right) = 0 \quad \text{und somit} \\ a_{i1} \left(c \cdot x_1^{(1)} \right) + a_{i2} \left(c \cdot x_2^{(1)} \right) + \dots + a_{in} \left(c \cdot x_n^{(1)} \right) = 0.$$

Damit ist auch $c \cdot \mathbf{x}^{(1)}$ eine Lösung des gegebenen LGS. □

Der folgende Satz stellt eine Verbindung zwischen zwei Lösungen eines beliebigen (inhomogenen) linearen Gleichungssystems und einer Lösung des zugehörigen homogenen LGS her.

Satz 1.2

Sind $\mathbf{x}^{(1)} = (x_1^{(1)}; x_2^{(1)}; \dots; x_n^{(1)})$ und $\mathbf{x}^{(2)} = (x_1^{(2)}; x_2^{(2)}; \dots; x_n^{(2)})$ zwei Lösungen eines beliebigen inhomogenen LGS

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

so ist die Differenz $\mathbf{x}^{(1)} - \mathbf{x}^{(2)} = (x_1^{(1)} - x_1^{(2)}; x_2^{(1)} - x_2^{(2)}; \dots; x_n^{(1)} - x_n^{(2)})$ dieser beiden Lösungen eine Lösung des zugehörigen homogenen LGS (1.22).

Der Beweis dieses Satzes sei der Leserin bzw. dem Leser überlassen (siehe die Aufgabe 5 auf S. 39).

Satz 1.3

Ist $\mathbf{x}^{(0)}$ eine (bekannte) feste Lösung eines inhomogenen LGS, so lässt sich jede Lösung $\mathbf{x}$ dieses LGS als Summe der Lösung $\mathbf{x}^{(0)}$ und einer Lösung $\mathbf{x}^{(h)}$ des zugehörigen homogenen LGS darstellen. Umgekehrt ist für jede Lösung $\mathbf{x}^{(h)}$ des homogenen LGS die Summe $\mathbf{x}^{(0)} + \mathbf{x}^{(h)}$ eine Lösung des inhomogenen LGS.

Beweis: Die Gültigkeit des ersten Teils des Satzes 1.3 folgt direkt aus dem Satz 1.2. Die Differenz $\mathbf{x} - \mathbf{x}^{(0)}$ zweier Lösungen des gegebenen inhomogenen LGS ist stets eine Lösung des zugehörigen homogenen LGS. Es lässt sich somit zu jeder beliebigen Lösung $\mathbf{x}$ eine Lösung $\mathbf{x}^{(h)}$ des homogenen LGS mit $\mathbf{x} = \mathbf{x}^{(0)} + \mathbf{x}^{(h)}$ finden. Auf den Beweis der Umkehrung wird hier verzichtet. □

Der Satz 1.3 gibt eine Möglichkeit, Lösungen eines beliebigen linearen Gleichungssystems aus einer einzigen bekannten Lösung dieses LGS zu ermitteln, falls die allgemeine Lösung (also die gesamte Lösungsmenge) des zugehörigen homogenen LGS bekannt ist.

1.3.5 Aufgaben zu Abschnitt 1.3

1. Bestimmen Sie für die folgenden linearen Gleichungssysteme die Ränge der einfachen und der erweiterten Koeffizientenmatrix und geben Sie die Lösungsmengen an.

$$\begin{array}{ll}
 3x_1 + 6x_2 + 2x_3 + x_4 = 9 & 2x_1 + x_2 - x_3 + 5x_4 = 2 \\
 \text{a) } 2x_1 + x_2 + 4x_3 + 3x_4 = 9 & \text{b) } -2x_2 + 3x_4 = 2 \\
 x_1 + 2x_2 + 2x_4 = 9 & 2x_1 + 3x_2 - x_3 + 3x_4 = 1 \\
 4x_1 + 3x_3 + 2x_4 = 9 &
 \end{array}$$

Geben Sie außerdem jeweils die Lösungsmenge des zugehörigen homogenen LGS an.

2. Folgt aus $\text{rg } \mathbf{A} = \text{rg } (\mathbf{A} | \vec{b})$, dass das lineare Gleichungssystem $(\mathbf{A} | \vec{b})$ eindeutig lösbar ist?

Hinweis: Suchen Sie ein Gegenbeispiel.

3. Überprüfen Sie die Richtigkeit der folgenden Aussage:

Ist ein lineares Gleichungssystem mit der erweiterten Koeffizientenmatrix $(\mathbf{A} | \vec{b})$ mit n Variablen und n Gleichungen für ein $\vec{b}$ *eindeutig* lösbar, so ist es dies auch für jedes $\vec{b}$.

Hinweis: Was lässt sich über die Zeilenstufenform der erweiterten Koeffizientenmatrix aussagen?

4. a) Es seien reelle Zahlen a, b, c, d, r, s vorgegeben. Begründen Sie, dass das lineare Gleichungssystem

$$\begin{array}{l}
 ax_1 + bx_2 = r \\
 cx_1 + dx_2 = s
 \end{array}$$

im Falle $ad - bc \neq 0$ eindeutig lösbar ist und geben Sie die Lösung an.

- b) Bestimmen Sie für $m \in \mathbb{R}$ die Lösungsmenge des folgenden LGS:

$$\begin{array}{l}
 -2x_1 + 3x_2 = 2m \\
 x_1 - 5x_2 = -11
 \end{array}$$

Hinweis:

Bringen Sie die erweiterte Koeffizientenmatrix in die Zeilenstufenform.

5. Beweisen Sie den Satz 1.2 auf S. 38.

6. a) Bestimmen Sie eine (feste) Lösung des linearen Gleichungssystems

$$\begin{array}{l}
 2x_1 + 3x_2 + x_3 = 1 \\
 4x_1 - x_2 + x_3 = 2 \\
 8x_1 + 5x_2 + 3x_3 = 4
 \end{array}$$

- b) Ermitteln Sie die Lösungsmenge des zugehörigen homogenen LGS.

- c) Geben Sie nun mithilfe ihrer Ergebnisse aus den Teilen a) und b) die Lösungsmenge des gegebenen inhomogenen Systems an.

7. Ein inhomogenes LGS mit zwei Variablen besitzt die Lösungen $(1; 2)$ und $(3; 4)$. Zeigen Sie, dass dann auch alle Paare der Form $(1 + 2t; 2 + 2t)$ mit $t \in \mathbb{R}$ Lösungen sind.

1.4 Anwendungen linearer Gleichungssysteme

Bereits in Abschnitt 1.1.4 wurde auf *Mischungsaufgaben* als Anwendungen linearer Gleichungssysteme mit zwei Gleichungen und zwei Variablen eingegangen. Komplexere Aufgaben dieser Art können auch auf kompliziertere Gleichungssysteme führen, die zudem nicht immer eindeutig lösbar sein müssen.

Beispiel 1.27

Ein metallurgischer Betrieb beabsichtigt, eine neue harte und dabei relativ leichte Metalllegierung auf den Markt zu bringen. Es wird gefordert, dass diese Legierung genau 4 % Titan und 2 % Chrom enthält, während der Rest aus Aluminium bestehen soll. Reines Titan und Chrom sind auf dem Markt nur zu recht hohen Preisen erhältlich; jedoch stehen titan- und chromhaltige Legierungen zur Verfügung, die entsprechend gemischt werden müssen.

- Legierung 1: 6 % Titan, 1 % Chrom
- Legierung 2: 1 % Titan, 3 % Chrom
- Legierung 3: 4 % Titan, 0 % Chrom
- Legierung 4: 3 % Titan, 4 % Chrom

Es soll eine Tonne der neuen Metalllegierung hergestellt werden. Gibt es Kombinationen der angebotenen Legierungen, mit denen die gewünschte Zusammensetzung erreicht werden kann?

Um diese Aufgabe zu lösen, stellen wir zunächst die Mischungsverhältnisse der vier Legierungen in einer Tabelle dar:

	Legierung 1	Legierung 2	Legierung 3	Legierung 4	Ziel-Legierung
Titan	6 %	1 %	4 %	3 %	4 %
Chrom	1 %	3 %	0 %	4 %	2 %
Masse	x_1	x_2	x_3	x_4	1 t

Die Berücksichtigung der gestellten Bedingungen führt zu einem LGS, in dem die zu verwendenden Massen der Legierungen als Variablen auftreten:

$$0,06 x_1 + 0,01 x_2 + 0,04 x_3 + 0,03 x_4 = 0,04$$

$$0,01 x_1 + 0,03 x_2 \quad \quad + 0,04 x_4 = 0,02$$

$$x_1 + x_2 + x_3 + x_4 = 1$$

Dieses Gleichungssystem kann nun mithilfe des Gauß-Algorithmus gelöst werden. Dabei tritt ein Parameter auf. Setzt man $x_4 = t$, so ergibt sich:

$$x_1 = \frac{2}{3} - t, \quad x_2 = \frac{4}{9} - t, \quad x_3 = -\frac{1}{9} + t, \quad x_4 = t.$$

Nicht für alle $t \in \mathbb{R}$ entstehen für das Ausgangsproblem sinnvolle Lösungen. Es kommen dafür nur Lösungen in Frage, für die alle Massen $x_1, \dots, x_4$ größer oder gleich Null sind. Überprüft man diese Bedingung, so ergeben sich für

$$\frac{1}{9} \leq t \leq \frac{4}{9}$$

sinnvolle Lösungen. Eine mögliche Lösung (mit $t = \frac{1}{3}$) wäre die Verwendung von $\frac{1}{3} \text{ t} \approx 333 \text{ kg}$ der Legierung A, ca. 111 kg der Legierung B, ca. 222 kg von Legierung C sowie ca. 333 kg von Legierung D. ■

Eine weiteres Anwendungsfeld linearer Gleichungssysteme sind Aufgaben, die sich mit Wasser-, Menschen- oder Verkehrsströmen bzw. -kreisläufen befassen.

Beispiel 1.28

An einer weiträumigen Straßenkreuzung (siehe Abb. 1.14) wurde durch Verkehrszählungen festgestellt, wie viele Fahrzeuge je Zeiteinheit in der Hauptverkehrszeit ein- und ausfahren. Damit im Kreuzungsbereich kein Stau entsteht, müssen Zu- und Abfluss in den Knoten A, B, C und D jeweils übereinstimmen.

- Wie viele Fahrzeuge müssen die Abschnitte AB, BC, CD und DA in jeder Zeiteinheit passieren, damit diese Bedingung erfüllt ist?

Für die Beantwortung dieser Frage kann ein LGS aufgestellt werden:

$$\begin{array}{lcl}
 \text{Knoten A:} & x_4 + 140 = x_1 + 120 & -x_1 \quad \quad \quad + x_4 = -20 \\
 \text{Knoten B:} & x_1 + 90 = x_2 + 100 & \text{bzw.} \quad x_1 - x_2 = 10 \\
 \text{Knoten C:} & x_2 + 70 = x_3 + 50 & \quad \quad \quad x_2 - x_3 = -20 \\
 \text{Knoten D:} & x_3 + 80 = x_4 + 110 & \quad \quad \quad x_3 - x_4 = 30
 \end{array}$$

Es ergibt sich eine einparametrische Lösungsmenge; mit $t = x_4$ ist:

$$x_1 = t + 20, \quad x_2 = t + 10, \quad x_3 = t + 30, \quad x_4 = t.$$

Diese Lösung bedarf einer Interpretation, da nicht alle rechnerisch richtigen Lösungen des LGS sinnvolle Lösungen der ursprünglichen Sachaufgabe sind.

Natürlich kann keine der Variablen negativ werden. Für $t = 0$ ergibt sich der geringste Verkehrsdurchsatz innerhalb der Kreuzung; dies bedeutet jedoch, dass der größte Teil der Verkehrsteilnehmer an den Knoten nach rechts abbiegt und somit die Kreuzung nicht befährt – dieses Szenario erscheint wenig realistisch.

Andererseits sind auch Lösungen für große t wenig realistisch, so wären z. B. für $t = 1000$ die Streckenabschnitte AB, BC, CD und DA in jeder Zeiteinheit jeweils von 1000 bis 1030 Fahrzeugen befahren. Vergleicht man diese Zahl mit der Zahl der Fahrzeuge, welche die Kreuzung je Zeiteinheit befahren oder verlassen, so wird deutlich, dass jedes dieser Fahrzeuge durchschnittlich fast drei „Ehrenrunden drehen“ müsste. Um zu entscheiden, ob dies überhaupt möglich wäre, müssten zusätzliche Informationen über den maximalen Durchsatz, den die Fahrbahnen bewältigen können, vorliegen. ■

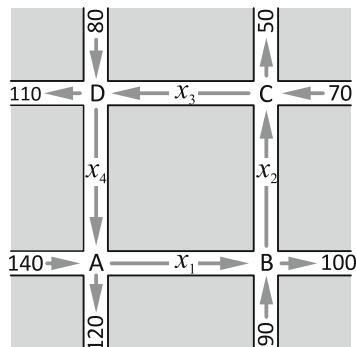


Abb. 1.14: Verkehrsströme an einer weiträumigen Straßenkreuzung

Eine gewisse Analogie zu Verkehrsströmen weisen *Stromstärken in elektrischen Schaltkreisen* auf. In einem Gleichstromkreis mit Verzweigungen (Knoten) lassen sich die Stromstärken durch die Bauteile (Widerstände) mithilfe der KIRCHHOFFschen *Regeln* berechnen:

- *Knotenregel*: In jedem Knotenpunkt eines Netzes ist die Summe der ankommenden Stromstärken gleich der Summe der abfließenden Stromstärken.
- *Maschenregel*: In jeder Masche eines Netzes ist die Summe der Spannungen gleich der Summe der Produkte aus den gerichteten Stromstärken und den Widerständen.

Beispiel 1.29

Wir betrachten den in Abb. 1.15 dargestellten Stromkreis. In beiden Knoten folgt aus dem ersten KIRCHHOFFschen Gesetz (Knotenregel) jeweils

$$I_1 = I_2 + I_3.$$

In der Masche, die nur die Widerstände R_2 und R_3 enthält, gilt nach der Maschenregel:

$$R_2 \cdot I_2 - R_3 \cdot I_3 = 0,$$

da sich innerhalb dieser Masche keine Spannungsquelle befindet. (Dieses Ergebnis lässt sich auch anhand der Tatsache überlegen, dass über den Widerständen R_2 und R_3 dieselbe Spannung U_2 anliegt und nach dem OHMSchen Gesetz $U_2 = R_2 \cdot I_2$ sowie gleichzeitig $U_2 = R_3 \cdot I_3$ sein muss.)

Weiterhin lassen sich Maschen betrachten, die aus der Spannungsquelle U_0 und den Widerständen R_1 , R_2 sowie aus U_0 , R_1 und R_3 bestehen. Hier gilt nach der Maschenregel

$$R_1 \cdot I_1 + R_2 \cdot I_2 = U_0 \quad \text{sowie}$$

$$R_1 \cdot I_1 + R_3 \cdot I_3 = U_0.$$

Wir fassen alle diese Zusammenhänge nun zu einem LGS zusammen:

$$I_1 - I_2 - I_3 = 0$$

$$R_2 \cdot I_2 - R_3 \cdot I_3 = 0$$

$$R_1 \cdot I_1 + R_2 \cdot I_2 = U_0$$

$$R_1 \cdot I_1 + R_3 \cdot I_3 = U_0$$

und lösen dieses nach den Stromstärken. Dabei lässt sich feststellen dass das LGS eine eindeutige Lösung besitzt:

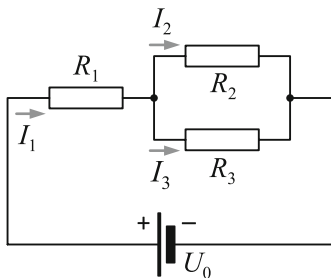


Abb. 1.15: Stromstärken in einem verzweigten Stromkreis

$$I_1 = U_0 \cdot \frac{R_2 + R_3}{R_1 R_2 + R_1 R_3 + R_2 R_3},$$

$$I_2 = U_0 \cdot \frac{R_3}{R_1 R_2 + R_1 R_3 + R_2 R_3},$$

$$I_3 = U_0 \cdot \frac{R_2}{R_1 R_2 + R_1 R_3 + R_2 R_3}.$$

Als konkretes Beispiel betrachten wir einen Stromkreis mit einer Spannungsquelle mit $U_0 = 12 \text{ V}$ und Widerständen $R_1 = 20 \Omega$, $R_2 = 10 \Omega$ und $R_3 = 20 \Omega$. Durch Einsetzen dieser Werte ergibt sich $I_1 = 0,45 \text{ A}$, $I_2 = 0,3 \text{ A}$ und $I_3 = 0,15 \text{ A}$. ■

1.4.1 Aufgaben zu Abschnitt 1.4

1. Edelstahl ist eine Legierung aus Eisen, Chrom und Nickel (und mitunter noch weiteren Stoffen). Der häufigste Legierungstyp eines nichtrostenden Stahls, der uns im Alltag begegnet, ist die Legierung X5CrNi18-10, oft als 18/10-Stahl bezeichnet. Die Bezeichnung 18/10 gibt den Anteil an Chrom/Nickel an: 18/10-Stahl besteht aus 72 % Eisen, 18 % Chrom und 10 % Nickel.

Es stehen drei Legierungen A, B und C zur Verfügung:

- Legierung A: 70 % Eisen, 22 % Chrom, 8 % Nickel
- Legierung B: 74 % Eisen, 18 % Chrom, 8 % Nickel
- Legierung C: 78 % Eisen, 15 % Chrom, 7 % Nickel

Um aus diesen Legierungen 18/10-Stahl herzustellen, muss zusätzlich Nickel (N) zugesetzt werden, der recht teuer ist. Ermitteln Sie, welche Massen der Legierungen A, B, C sowie N benötigt werden, um eine Tonne 18/10-Stahl herzustellen, wobei möglichst wenig reiner Nickel zum Einsatz kommen soll.

2. In Abb. 1.16 sind die Zu- und Abfahrten in einem Kreisverkehr je Zeiteinheit in der Hauptverkehrszeit dargestellt. Damit kein Stau entsteht, muss in jedem der Knoten A, B, C, D, E und F die Anzahl der Zufahrten gleich der Anzahl der Abfahrten sein. Wie viele Fahrzeuge müssen die Streckenabschnitte AB, BC, CD, DE, EF und FA in jeder Zeiteinheit passieren, damit diese Bedingung erfüllt ist? Welcher Streckenabschnitt wird dabei am stärksten belastet?

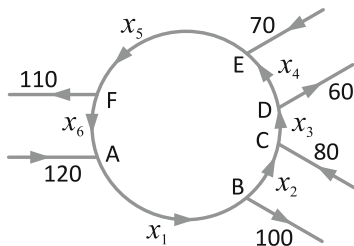


Abb. 1.16: Kreisverkehr (Aufgabe)

3. Das Schaltbild in Abb. 1.17 zeigt einen Gleichstromkreis.

- a) Gegeben sind die Spannung $U_0 = 12\text{ V}$ und die Widerstände $R_1 = 20\ \Omega$, $R_2 = 40\ \Omega$, $R_3 = 10\ \Omega$ und $R_4 = 20\ \Omega$. Stellen Sie ein geeignetes LGS auf und berechnen Sie die Stromstärken I_0 , I_1 , I_2 , I_3 und I_4 .
- b) Geben Sie nun die Stromstärken $I_0, \dots, I_4$ allgemein in Abhängigkeit von U_0 und $R_1, \dots, R_4$ an (siehe Beispiel 1.29).

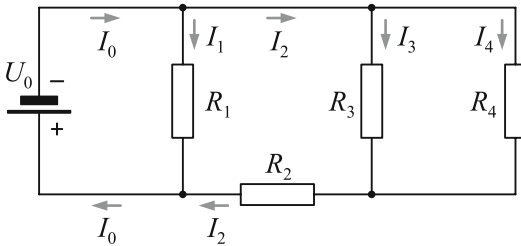


Abb. 1.17: Verzweigung von Stromstärken (Aufgabe)

1.5 Lösen linearer Gleichungssysteme mithilfe des Computers

Wie bei der Lösung der Übungsaufgaben dieses Kapitels festzustellen war, ist das „händische“ Lösen linearer Gleichungssysteme oft mühsam und fehleranfällig. Geeignete Computerprogramme (und mittlerweile auch Taschencomputer und graphische Taschenrechner) benötigen selbst für das Lösen „großer“ LGS nur Sekundenbruchteile. Besonders gut eignen sich für das Lösen linearer Gleichungssysteme und eine Vielzahl weiterer Berechnungen der Linearen Algebra sowie für Visualisierungen Computeralgebrasysteme (CAS) wie Mathematica, Maple, MuPAD und Derive. Inzwischen steht auch ein leistungsfähiges freies Computeralgebrasystem zur Verfügung: *Maxima*. Da dieses kostenlose CAS auch im Bereich der Linearen Algebra eine Vielzahl von Möglichkeiten bietet, wird es an vielen Stellen innerhalb dieses Buches zum Einsatz kommen.⁶ Laden Sie Maxima zunächst unter

<http://maxima.sourceforge.net/>

für Ihr Betriebssystem (Windows, Linux oder MacOS) herunter und installieren Sie die Software. Starten Sie dann die Benutzeroberfläche **wxMaxima** (Abb. 1.18),

⁶Andere CAS wie Maple, Mathematica und MuPAD lassen sich sehr ähnlich zu Maxima nutzen und stellen die für dieses Buch benötigten Funktionen ebenfalls zur Verfügung. Sollten Sie mit einem dieser Systeme bereits vertraut sein, so werden Sie die hier für Maxima beschriebenen Befehle in „Ihrem“ CAS auch recht leicht finden, wenngleich diese zwischen den verschiedenen CAS jeweils etwas voneinander abweichen.

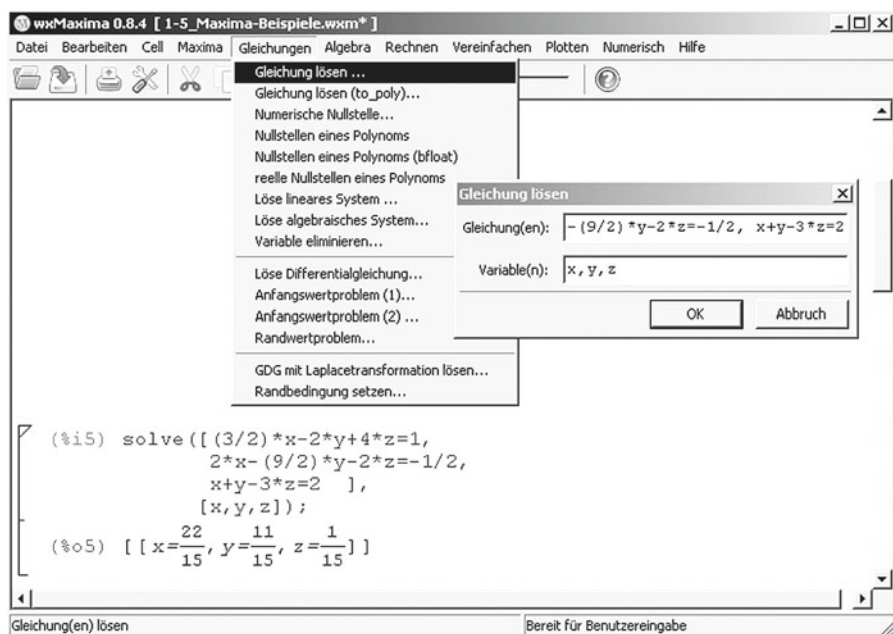


Abb. 1.18: Das Computeralgebrasystem Maxima mit der Benutzeroberfläche wxMaxima

mit der sich Maxima gut bedienen lässt. Bei den folgenden Anleitungen und Beispielen wird davon ausgegangen, dass Maxima mit wxMaxima genutzt wird.

Um das LGS aus dem Beispiel 1.17 (S. 20) zu lösen, gibt man in Maxima ein:

```
solve( [(3/2)*x-2*y+4*z=1, 2*x-(9/2)*y-2*z=-1/2, x+y-3*z=2] ,
        [x,y,z] );
```

und übergibt diese Anweisung an Maxima zur Berechnung durch gleichzeitiges Drücken der **Shift**- und **Enter**-Taste. Maxima gibt sofort die Lösungsmenge in folgender Form aus:

$$\left[\left[x = \frac{22}{15}, y = \frac{11}{15}, z = \frac{1}{15} \right] \right]$$

Alternativ zur Eingabe des Befehls kann auch das Dialogfenster „Gleichung lösen“ im Menü „Gleichungen“ aufgerufen werden, wo dann die Gleichungen und die Variablen, nach denen sie gelöst werden sollen, eingegeben werden (siehe Abb. 1.18).

Der **solve**-Befehl löst einzelne (nicht nur lineare) Gleichungen und Gleichungssysteme. Man muss Maxima die Gleichungen mitteilen und außerdem eingeben, nach welchen Variablen die Gleichungen bzw. LGS gelöst werden sollen (hier x, y, z). Die allgemeine Syntax des **solve**-Befehls ist:

```
solve( [ Gleichung 1, ..., Gleichung n ] , [ Variable 1, ..., Variable n ] );
```

Zeilenumbrüche und Leerzeichen dürfen in Maxima beliebig (aber nicht innerhalb von Namen wie **solve**) eingefügt werden; hingegen ist unbedingt auf *Groß- und Kleinschreibung* sowie die *Verwendung der richtigen Klammern* zu achten.

Bei dem betrachteten Beispiel wurde vielleicht nicht verständlich, warum eingegeben werden muss, *nach welchen Variablen* ein LGS gelöst werden soll – es traten hier nur drei Variablen auf, und nach diesen sollte das LGS auch gelöst werden. Jedoch können in einem Gleichungssystem auch durch Buchstaben bezeichnete Konstanten auftreten, und Variablen können beliebige Namen haben. Ohne die Angabe, nach welchen Variablen ein LGS gelöst werden soll, kann Maxima dann natürlich keine Lösung angeben. Dies wird an dem LGS deutlich, das in dem Beispiel 1.29 auf S. 42 gelöst wurde. Die Spannung U_0 und die Widerstände R_1, R_2, R_3 waren hier Konstanten und das Gleichungssystem sollte nach I_1, I_2 und I_3 gelöst werden. Dem entspricht in Maxima die Eingabe

```
solve([i1-i2-i3=0, r2*i2-r3*i3=0, r1*i1+r2*i2=u0, r1*i1+r3*i3=u0],
      [i1, i2, i3] );
```

Maxima „antwortet“ mit der Lösung

$$\left[\left[i_1 = \frac{r_3 u_0 + r_2 u_0}{(r_2 + r_1) r_3 + r_1 r_2}, i_2 = \frac{r_3 u_0}{(r_2 + r_1) r_3 + r_1 r_2}, i_3 = \frac{r_2 u_0}{(r_2 + r_1) r_3 + r_1 r_2} \right] \right]$$

(Die Leserin oder der Leser kann nun versuchen, das LGS etwa nach R_1, R_2, R_3 zu lösen.) Interessant ist folgende Zusatzinformation, die Maxima ausgibt:

solve: dependent equations eliminated (4)

Die Software hat also erkannt, dass eine Gleichung von den anderen abhängt und diese „eliminiert“.

Wir untersuchen im Folgenden, wie sich Maxima bei LGS verhält, die *nicht lösbar* sind oder *unendlich viele Lösungen* besitzen und ziehen dazu die beiden Gleichungssysteme aus dem Beispiel 1.18 auf S. 20 heran. Für das nicht lösbare LGS a erhalten wir nach der Eingabe

```
solve([3*x-4*y+8*z=16, 6*x-y+5*z=-4, -3*x-3*y+3*z=-8], [x, y, z] );
```

die schlichte Ausgabe

[]

Für das LGS mit einparametriger Lösungsmenge (Beispiel 1.18 b) führt

```
solve([3*x-4*y+8*z=2, 6*x-y+5*z=4, -3*x-3*y+3*z=-2], [x, y, z] );
```

zu der Ausgabe

$$\left[\left[x = -\frac{12 \%r_3 - 14}{21}, y = \frac{11 \%r_3}{7}, z = \%r_3 \right] \right]$$

Maxima hat also z als Parameter gewählt und ihm den Namen $\%r_3$ gegeben. Leicht lässt sich auch die Wahl einer anderen Variablen als Parameter erreichen, wenn bekannt ist, dass ein LGS unendlich viele Lösungen besitzt. Gibt man

```
solve([3*x-4*y+8*z=2, 6*x-y+5*z=4, -3*x-3*y+3*z=-2], [x, z] );
```

ein, so erhält man

$$\left[\left[x = -\frac{12 y - 22}{33}, z = \frac{7 y}{11} \right] \right]$$

Maxima nimmt in diesem Falle y als konstant an. Die Ausgabe lässt sich aber auch in dem Sinne interpretieren, dass y als Parameter gewählt wird, denn y kann beliebige Werte annehmen.

Oft ist es sinnvoll, ein LGS mit Symbolen für die Koeffizienten in Maxima einzugeben, diesen aber dann *konkrete Werte zuzuweisen*, wenn man Lösungen für bestimmte Koeffizienten erhalten möchte. Wir betrachten dazu erneut das Beispiel 1.29; es sollen die Stromstärken für die festgelegten Werte $U_0 = 12\text{ V}$, $R_1 = 20\ \Omega$, $R_2 = 10\ \Omega$ und $R_3 = 20\ \Omega$ berechnet werden. Dazu gibt man ein:

```
u0:12; r1:20; r2:10; r3:20;
solve([i1-i2-i3=0, r2*i2-r3*i3=0, r1*i1+r2*i2=u0, r1*i1+r3*i3=u0],
      [i1, i2, i3] );
```

Maxima gibt das bereits aus Beispiel 1.29 bekannte Ergebnis aus:

$$\left[\left[i_1 = \frac{9}{20}, i_2 = \frac{3}{10}, i_3 = \frac{3}{20} \right] \right]$$

Nun lassen sich bequem die Werte für die Spannung und die Widerstände variieren und die entsprechenden Stromstärken berechnen.

Achtung: Löscht man die Eingaben für **u0**, **r1**, **r2**, **r3** und löst das LGS erneut, so erhält man nicht mehr die allgemeine Lösung, sondern eine spezielle Lösung, die sich auf die zuletzt eingegebenen Werte bezieht. Wurden nämlich Symbole mit Werten belegt, so „merkt“ sich Maxima diese Werte und lässt z. B. **u0**, **r1**, **r2**, und **r3** nicht mehr als freie Variablen zu. Um alle Werte zu löschen und die entsprechenden Symbole zu „befreien“, gibt man **kill(all)**; ein oder wählt den Menübefehl „Maxima → Speicher löschen“.

Maxima kann nicht nur Berechnungen ausführen, sondern auch *graphische Darstellungen* erstellen. Neben Funktionsgraphen lassen sich auch durch die Eingabe von Gleichungen in zwei oder drei Variablen die entsprechenden Geraden oder Ebenen visualisieren. Dazu stehen Funktionen zur Darstellung „implizit“ gegebener Objekte zur Verfügung. Mit „implizit“ ist gemeint, dass Graphen (Geraden oder Ebenen) nicht direkt (explizit) durch Funktionsgleichungen, sondern durch Gleichungen gegeben sind und Maxima selbst Lösungen bestimmen muss, um die betreffenden Geraden oder Ebenen darzustellen.

Um die Lösungsmengen der Gleichungen des LGS aus dem Beispiel 1.3 auf S. 5 darzustellen, kann in Maxima

```
G111: 4*x - 5*y = 13;
G112: 3*x + 4*y = 3;
load("draw")$
draw2d( color = blue, implicit(G111, x,-4,4, y,-4,4),
        color = red , implicit(G112, x,-4,4, y,-4,4) );
```

eingegeben werden. Die Geraden werden dann innerhalb der für x und y angegebenen Intervalle (jeweils $[-4; 4]$) dargestellt. Dazu öffnet Maxima ein eigenes Fenster (siehe Abb. 1.19 a). Innerhalb der **draw2d**-Anweisung können zahlreiche Optionen angegeben werden, die in der Maxima-Hilfe erklärt sind. So bewirkt **user_preamble = ["set size ratio 1" , "set zeroaxis"]**

dass x - und y -Achse gleich skaliert und die Achsen gezeichnet werden.

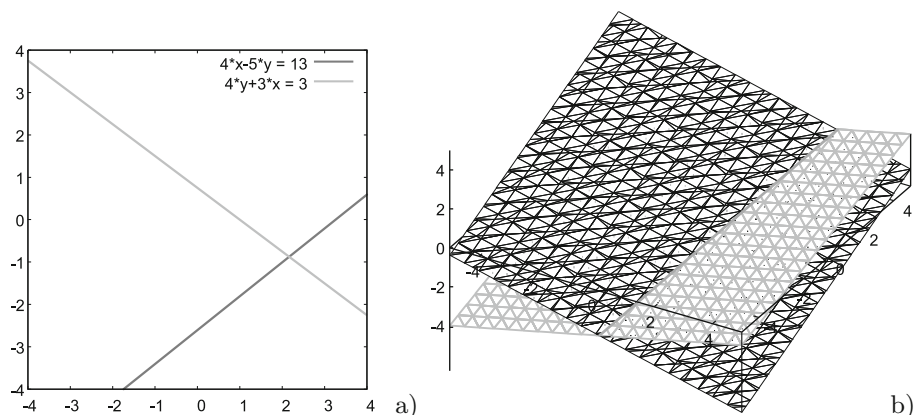


Abb. 1.19: Darstellungen der Lösungsmengen von Gleichungen in Maxima

In ähnlicher Weise lassen sich Lösungsmengen von Gleichungen mit drei Variablen als Ebenen darstellen. Durch die folgende Eingabe werden die durch die beiden Gleichungen aus Beispiel 1.16 auf S. 19 beschriebenen Ebenen dargestellt:

```
G11 : (3/2)*x - 2*y + 4*z = 1;
```

```
G12 : x + y - 3*z = 2;
```

```
load(draw)$
```

```
draw3d( user_preamble = ["set size ratio 1"], surface_hide = true,
        color = blue, implicit(G11, x,-5,5, y,-5,5, z,-5,5),
        color = red , implicit(G12, x,-5,5, y,-5,5, z,-5,5) );
```

Es öffnet sich wiederum ein separates Grafikenfenster (siehe Abb. 1.19 b), in welchem sich die Darstellung mit der Maus interaktiv drehen und dadurch aus verschiedenen Richtungen betrachten lässt. Die Darstellung erfolgt in dem Bereich, der durch die x-, y- und z-Werte -5 und 5 begrenzt wird. Der Darstellungsbereich lässt sich verändern, indem die Werte in `x,-5,5, y,-5,5, z,-5,5` (für beide Gleichungen) angepasst werden. Die Option `surface_hide = true` bewirkt, dass verdeckte Bereiche von Ebenen nicht dargestellt werden. Die Ebenen werden zunächst recht grob aufgelöst durch Dreiecke dargestellt. Für eine feinere Auflösung kann vor der mit `color = blue` beginnenden Zeile eine Zeile

```
x_voxel = 20, y_voxel=20, z_voxel=20,
```

eingefügt werden. Höhere Werte bewirken hierbei eine feiner aufgelöste Darstellung, verlängern aber die für die Berechnung der Grafiken benötigte Zeit.

Es sei hier zunächst nur erwähnt, dass sich mit Maxima Gleichungssysteme auch in der Matrixschreibweise eingeben und lösen lassen. Zudem können Ränge von Matrizen bestimmt und vielfältige weitere Operationen mit Matrizen ausgeführt werden. Darauf wird in dem Kapitel 6 noch näher eingegangen.

Eine Maxima-Datei mit allen in diesem Abschnitt besprochenen Beispielen steht auf der Internetseite www.afiller.de/linalg zur Verfügung.

2 Koordinatengeometrie

Übersicht

2.1	Geraden in der Ebene.....	50
2.2	Abstände von Punkten; Kreisgleichungen.....	55
2.3	Kugeln und Kegel.....	67
2.4	Kurven zweiter Ordnung: Ellipsen, Hyperbeln und Parabeln.....	72

Im vorangegangenen Kapitel wurden bereits mehrfach geometrische Veranschaulichungen für Lösungsmengen linearer Gleichungssysteme genutzt. Umgekehrt sind Beschreibungen geometrischer Objekte durch Koordinaten und Gleichungen von großem Nutzen für die Geometrie, auf die dadurch rechnerische und algebraische Methoden angewendet werden können.

Als Begründer der Koordinatenmethode und somit der analytischen Geometrie werden RENÉ DESCARTES (1596–1650) und PIERRE DE FERMAT (1607–1665) angesehen, wenngleich Ansätze hierfür bereits früher zu finden waren. Besonderen Einfluss erlangte das Werk „La Géométrie“ (1637) von DESCARTES. Sein Leitmotiv für die Verwendung von Koordinaten war die *Beschreibung und Klassifikation ebener Kurven durch Gleichungen*. Im Mittelpunkt des Interesses standen dabei zunächst die Kegelschnitte als algebraische Kurven zweiter Ordnung. Diese werden in diesem Kapitel nach der Behandlung von Gleichungen für Geraden und Kreise thematisiert.

Es wird sich im Folgenden zeigen, dass mithilfe von Koordinatenbeschreibungen viele Untersuchungen von Eigenschaften und Lagebeziehungen ebener geometrischer Objekte möglich sind. Gleichzeitig zeigen sich aber vor allem bei der Koordinatengeometrie des Raumes Schwierigkeiten und Grenzen, weshalb Geraden im Raum, Ebenen sowie Kugeln und Kegel zunächst nur ansatzweise behandelt werden. Die Grenzen reiner Koordinatenbeschreibungen werden dann in den folgenden Kapiteln – nach einem tieferen Eindringen in die Vektorrechnung – mithilfe vektorieller Vorgehensweisen überwunden.

2.1 Geraden in der Ebene

2.1.1 Koordinatengleichungen von Geraden

In dem Abschnitt 1.1.1 zeigte sich, dass sich Geraden in der Ebene durch Gleichungen mit zwei Variablen beschreiben lassen:

$$g: ax + by = c \quad (a, b, c \in \mathbb{R}, a \neq 0 \text{ oder } b \neq 0). \quad (2.1)$$

Durch eine derartige Gleichung wird eindeutig eine Gerade beschrieben. Es lassen sich aber zu jeder Geraden mehrere Gleichungen angeben. Z. B. beschreiben

$$2x + 5y = 12 \quad \text{und} \quad \frac{1}{6}x + \frac{5}{12}y = 1$$

dieselbe Gerade, denn beide Gleichungen sind äquivalent und besitzen daher dieselbe Lösungsmenge.

Sind zwei Punkte $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ einer Geraden g gegeben, so lassen sich deren Koordinaten in (2.1) einsetzen, um die Koeffizienten a , b und c so zu bestimmen, dass (2.1) die durch $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ verlaufende Gerade beschreibt. Dies geschieht durch Lösen des linearen Gleichungssystems

$$\begin{aligned} x_1 a + y_1 b &= c \\ x_2 a + y_2 b &= c. \end{aligned}$$

Dieses LGS mit den drei Unbekannten a , b und c ist offensichtlich nicht eindeutig lösbar, was auch daher plausibel ist, dass es zu jeder Geraden mehrere Gleichungen der Form (2.1) gibt. Es ist daher ein Koeffizient festzulegen (z. B. $c=1$) und dann das LGS zu lösen.

Beispiel 2.1

Gesucht sind Gleichungen der Geraden, die durch folgende Punkte verlaufen:

a) g durch $P_1(3; 5)$ und $P_2(6; 7)$

Wir setzen $c=1$ und lösen das LGS

$$3a + 5b = 1$$

$$6a + 7b = 1.$$

Es ergibt sich $b = \frac{1}{3}$ und $a = -\frac{2}{9}$.

Als Gleichung für die Gerade g erhalten wir:

$$-\frac{2}{9}x + \frac{1}{3}y = 1.$$

Durch Multiplikation mit 9 lässt sich diese Gleichung auch in der Form $-2x + 3y = 9$ schreiben.

b) h durch $Q_1(-4; 3)$ und $Q_2(8; -6)$

Wir setzen wieder $c=1$. Das LGS

$$-4a + 3b = 1$$

$$8a - 6b = 1.$$

ist jedoch nicht lösbar, und es wäre auch für ein beliebiges anderes c mit $c \neq 0$ unlösbar. Deshalb setzen wir nun $c=0$ und lösen das LGS

$$-4a + 3b = 0$$

$$8a - 6b = 0.$$

Alle Paare $(a; b)$ mit $b = \frac{4}{3}a$ sind Lösungen dieses LGS, also z. B. $a = 3$, $b = 4$. Somit ist $3x + 4y = 0$ eine Gleichung der Geraden h .

Die Gerade h verläuft durch den Koordinatenursprung, deshalb muss $c = 0$ sein. Für alle Geraden, die nicht durch den Ursprung verlaufen, gibt es hingegen eine Geradengleichung der Form (2.1) mit $c = 1$. ■

Für $b \neq 0$ lässt sich die Geradengleichung (2.1) in die *Normalform*

$$y = mx + n \quad (m, n \in \mathbb{R}) \quad (2.2)$$

bringen. Dabei ist m der Anstieg der beschriebenen Geraden. Mit den Koeffizienten in der Gleichung (2.1) gilt $m = -\frac{a}{b}$.

Satz 2.1

Sind $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ zwei beliebige (voneinander verschiedene) Punkte einer durch eine Gleichung der Form (2.1) gegebenen Geraden mit $b \neq 0$, so gilt

$$\frac{y_2 - y_1}{x_2 - x_1} = -\frac{a}{b} = m.$$

Der Beweis dieses Satzes erfolgt durch Einsetzen der Koordinaten beider Punkte in (2.1) und Umstellen (siehe Aufgabe 2 auf S. 54).

Der Satz 2.1 führt zu weiteren Möglichkeiten, Geraden durch Gleichungen zu beschreiben. Da dieser Satz für *jeweils* zwei beliebige Punkte einer Geraden anwendbar ist, gilt für drei paarweise verschiedene Punkte $P_1(x_1; y_1)$, $P_2(x_2; y_2)$ und $P(x; y)$ sowohl $\frac{y_2 - y_1}{x_2 - x_1} = m$ als auch $\frac{y - y_1}{x - x_1} = m$, siehe Abb. 2.1. Zusammengefasst ergibt sich daraus die *Zweipunkteform der Geradengleichung*:

$$\frac{y - y_1}{x - x_1} = \frac{y_2 - y_1}{x_2 - x_1} \quad \text{bzw.} \quad y - y_1 = \frac{y_2 - y_1}{x_2 - x_1}(x - x_1). \quad (2.3)$$

Hierbei können $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ als feste (bekannte) Punkte und $P(x; y)$ als beliebiger (variabler) Punkt einer Geraden aufgefasst werden. Durch Einsetzen der bekannten Koordinaten lässt sich also mithilfe der Zweipunkteform recht leicht eine Geradengleichung aufstellen, wenn zwei Punkte bekannt sind (siehe dazu die Aufgabe 1 auf S. 54).

Eine besonders einfache Geradengleichung lässt sich aufstellen, falls statt zweier beliebiger Punkte einer Geraden deren Schnittpunktkoordinaten $(0; n_y)$ mit der y -Achse und $(n_x; 0)$ mit der x -Achse gegeben sind; n_x und n_y werden dabei als *Achsenabschnitte* bezeichnet, siehe Abb. 2.2 (n_y ist mit n in Gleichung (2.2) identisch). In diesem Falle lässt sich aus der Zweipunkteform die *Achsenabschnittsform der Geradengleichung* herleiten (siehe Aufgabe 4 auf S. 54):

$$\frac{x}{n_x} + \frac{y}{n_y} = 1. \quad (2.4)$$

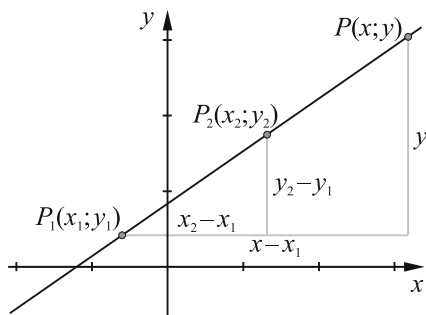


Abb. 2.1: Steigungsdreiecke

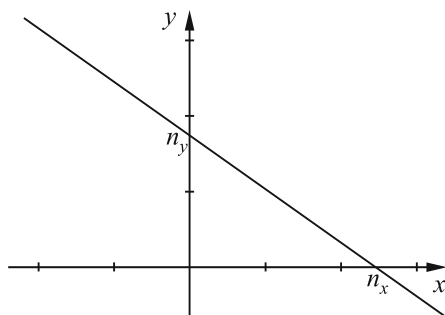


Abb. 2.2: Achsenabschnitte

Keine der hier aufgeführten Formen der Geradengleichung ist auf *Geraden im dreidimensionalen Raum* übertragbar da hierbei drei Variablen (die Koordinaten x , y und z) auftreten würden. Lösungsmengen einzelner linearer Gleichungen in drei Variablen können jedoch niemals Geraden beschreiben, siehe Abschnitt 1.2. Dort wurde auch bereits herausgearbeitet, dass Lösungsmengen linearer Gleichungssysteme mit drei Variablen und zwei Gleichungen in den meisten Fällen Geraden beschreiben (vgl. Beispiel 1.16 auf S. 19). Um eine durch zwei Punkte des Raumes verlaufende Gerade durch ein LGS zu beschreiben, sind zwei lineare Gleichungen aufzustellen, welche die Koordinaten der Punkte erfüllen, wobei diese Gleichungen keine Vielfachen voneinander sein dürfen.

Beispiel 2.2

Bestimmung eines Gleichungssystems, das die Gerade durch die beiden Punkte $P(3; -4; 7)$ und $Q(8; 5; -9)$ des Raumes beschreibt.

Gesucht ist ein lineares Gleichungssystem der Form

$$\begin{aligned} a_1 x + b_1 y + c_1 z &= d_1 \\ a_2 x + b_2 y + c_2 z &= d_2, \end{aligned}$$

welches die Lösungstriple $(3; -4; 7)$ und $(8; 5; -9)$ besitzt. Um dafür geeignete Koeffizienten a_1, b_1, c_1, d_1 sowie a_2, b_2, c_2 und d_2 zu bestimmen, setzt man die Koordinaten der gegebenen Punkte in beide Gleichungen ein und führt zunächst den ersten Schritt des Gauß-Algorithmus durch.

$$\begin{array}{rclcl} 3a_1 - 4b_1 + 7c_1 &= d_1 & | \cdot (-8) & 3a_2 - 4b_2 + 7c_2 &= d_2 & | \cdot (-8) \\ 8a_1 + 5b_1 - 9c_1 &= d_1 & | \cdot 3 & 8a_2 + 5b_2 - 9c_2 &= d_2 & | \cdot 3 \\ \hline 47b_1 - 83c_1 &= -5d_1 & & 47b_2 - 83c_2 &= -5d_2 & \end{array}$$

Jeweils zwei von den Werten $b_{1/2}, c_{1/2}, d_{1/2}$ können frei festgelegt werden, allerdings ist darauf zu achten, dass die linken Seiten des obigen Gleichungssystems keine Vielfachen voneinander sind. Wählt man zum Beispiel $c_1 = 0, d_1 = 47$ und $b_2 = 0$ und $d_2 = 83$, so ist diese Forderung erfüllt ($d_{1/2}$ wurden so ausgewählt, dass bei der folgenden Berechnung Brüche vermieden werden). Berechnet man mit diesen Werten in den beiden LGS die fehlenden Koeffizienten, so ergibt sich

$$a_1 = 9, b_1 = -5, c_1 = 0, d_1 = 47, \quad a_2 = 16, b_2 = 0, c_2 = 5, d_2 = 83.$$

Ein LGS, das die Gerade durch die beiden gegebenen Punkte beschreibt, ist also

$$\begin{aligned} 9x - 5y &= 47 \\ 16x + 5z &= 83. \end{aligned}$$

Um Eigenschaften von Geraden im Raum zu untersuchen, ist die Beschreibung durch lineare Gleichungssysteme recht „unhandlich“. Besser gelingt dies mithilfe von Parametergleichungen. Diese wurden bereits kurz in Abschnitt 1.1.1 für Geraden in der Ebene sowie in Abschnitt 1.2.1 für Ebenen und Geraden des Raumes als Darstellungen von Lösungsmengen linearer Gleichungssysteme betrachtet. Nach einer Vertiefung der Vektorrechnung werden Parameterdarstellungen von Geraden und Ebenen in dem Kapitel 4 ausführlicher behandelt.

2.1.2 Winkel zwischen Geraden in der Ebene

Der *Steigungswinkel* einer Geraden in einem ebenen Koordinatensystem lässt sich anhand eines Steigungsdreiecks berechnen, das durch zwei beliebige Punkte P_1 und P_2 der Geraden festgelegt wird, siehe Abb. 2.3. Die Längen der Katheten dieses rechtwinkligen Dreiecks entsprechen den Beträgen $|x_2 - x_1|$ und $|y_2 - y_1|$ der Koordinatendifferenzen der beiden Punkte. Für den Steigungswinkel α gilt:

$$\tan \alpha = \frac{y_2 - y_1}{x_2 - x_1} = m.$$

Dabei ist m der Anstieg der Geraden. α ist ein orientiertes Winkelmaß, d. h. für $\frac{y_2 - y_1}{x_2 - x_1} < 0$ ist $\alpha < 0$, während positiven Anstiegen positive Maße des Steigungswinkels entsprechen.

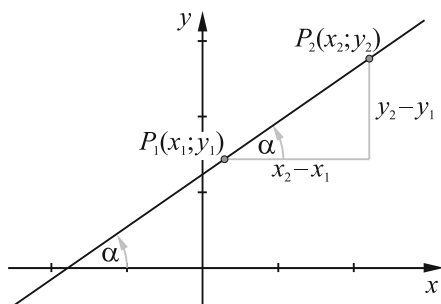


Abb. 2.3: Steigungswinkel einer Geraden

Um den *Schnittwinkel* β zweier Geraden g und h zu bestimmen, sind ihre Steigungswinkel α_g und α_h zu subtrahieren. Falls der Betrag $|\alpha_g - \alpha_h|$ der Differenz kleiner als 90° bzw. $\frac{\pi}{2}$ ist, so gilt für den Schnittwinkel der Geraden g und h :

$$\beta = |\alpha_g - \alpha_h|,$$

und zwar unabhängig davon, ob α_g und α_h gleiche oder unterschiedliche Vorzeichen haben, siehe Abb. 2.4 a), b). Da allerdings als Winkelmaß zwischen zwei Geraden g und h immer das Maß des kleineren der beiden Winkel zwischen g und h angegeben wird, setzt man für den Fall $|\alpha_g - \alpha_h| > 90^\circ$ (vgl. Abb. 2.4 c):

$$\beta = 180^\circ - |\alpha_g - \alpha_h|.$$

Mithilfe des Additionstheorems $\tan(\alpha_2 - \alpha_1) = \frac{\tan \alpha_2 - \tan \alpha_1}{1 + \tan \alpha_1 \cdot \tan \alpha_2}$ der Tangensfunktion (für $\tan \alpha_1 \cdot \tan \alpha_2 \neq -1$) lässt sich der Schnittwinkel zweier Geraden

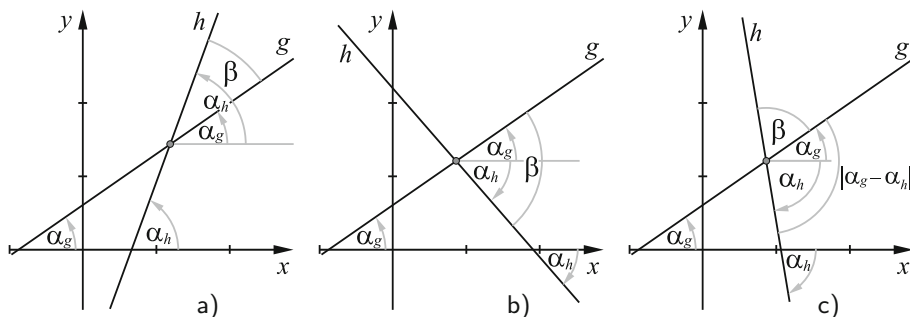


Abb. 2.4: Schnittwinkel zweier Geraden

g und h direkt aus den Anstiegen $m_g = \tan \alpha_g$ und $m_h = \tan \alpha_h$ der Geraden berechnen:

$$\tan \beta = \frac{m_g - m_h}{1 + m_g \cdot m_h} \quad \text{falls } m_g \cdot m_h \neq -1.$$

Für $m_g \cdot m_h = -1$ liegt ein besonderer Fall vor. Die obige Gleichung kann dann nicht zur Berechnung des Schnittwinkels verwendet werden. Dies hängt damit zusammen, dass der Tangens eines rechten Winkels nicht definiert ist (denn der Kosinus ist in diesem Falle Null). In der Tat lässt sich zeigen, dass zwei Geraden senkrecht sind, wenn das Produkt ihrer Anstiege -1 ist (siehe Aufgabe 10).

Satz 2.2

Zwei Geraden g und h sind genau dann orthogonal (senkrecht), wenn für ihre Anstiege m_g und m_h gilt: $m_h = -\frac{1}{m_g}$.

2.1.3 Aufgaben zu Abschnitt 2.1

- Stellen Sie Gleichungen der Geraden durch die gegebenen Punkte auf:
 - $P_1(3; -2)$ und $P_2(11; -11)$
 - $Q_1(\frac{2}{3}; \frac{3}{4})$ und $Q_2(8; 9)$
- Beweisen Sie den Satz 2.1 (siehe Seite 51).
- Stellen Sie für die in der Aufgabe 1 gegebenen Geraden Gleichungen in der Zweipunkteform (2.3) auf und wandeln Sie diese in die allgemeine Form (2.1) um. Vergleichen Sie die Ergebnisse mit Ihren Ergebnissen zu Aufgabe 1.
- Leiten Sie aus der Zweipunkteform (2.3) die Achsenabschnittsform (2.4) her. Setzen Sie dazu in (2.3) die Koordinaten der Schnittpunkte $P_1(0; n_y)$ und $P_2(n_x; 0)$ ein und formen Sie die entstehende Gleichung um.
- Bestimmen Sie ein lineares Gleichungssystem (bestehend aus zwei Gleichungen), das die Gerade durch die Punkte $P(0; 5; -2)$ und $Q(14; 3; 2)$ beschreibt.
- Berechnen Sie die Steigungswinkel folgender Geraden:
 - $g: 3x - y = 9$
 - $h: y + \frac{1}{2} = \frac{1}{2}(x - 2)$
 - PQ mit $P(3; 1)$, $Q(-1; \frac{1}{2})$
- Berechnen Sie jeweils den Schnittpunkt und den Schnittwinkel der Geraden g und h :
 - $g: y - 1 = \frac{1}{2}(x - \frac{5}{2})$, $h: y = -2x - 1$
 - $g: y = 4x - 3$, $h: y = x + 5$
- Skizzieren Sie das Dreieck $\triangle ABC$, dessen Seiten auf den Geraden $a: y = 0,5x$, $b: y = 3x$ und $c: x + y = 6x$ liegen. Berechnen Sie die Koordinaten der Eckpunkte und die Maße der Innenwinkel dieses Dreiecks.
- Gegeben sei eine Gerade g mit einem Anstieg m durch einen Punkt $P_0(x_0; y_0)$. Geben Sie eine Gleichung der Geraden g sowie eine Gleichung der zu g senkrechten Geraden durch P_0 an.
- Beweisen Sie den Satz 2.2.

2.2 Abstände von Punkten; Kreisgleichungen

2.2.1 Abstände von Punkten in der Ebene und im Raum

Für den Abstand $|P_1P_2|$ zweier Punkte $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ der Ebene gilt nach dem Satz des Pythagoras $|P_1P_2|^2 = (x_2 - x_1)^2 + (y_2 - y_1)^2$ (siehe Abb. 2.5).

Zwei Punkte $P_1(x_1; y_1; z_1)$ und $P_2(x_2; y_2; z_2)$ des Raumes bestimmen einen Quader mit zu den Koordinatenachsen parallelen Kanten, in dem sie Endpunkte einer Raumdiagonalen sind (siehe Abb. 2.6). Für die Länge $|P_1Q|$ der Diagonalen der Grundfläche dieses Quaders gilt wiederum nach dem Satz des Pythagoras $|P_1Q|^2 = (x_2 - x_1)^2 + (y_2 - y_1)^2$. Um $|P_1P_2|$ zu berechnen, wendet man in dem bei Q rechtwinkligen Dreieck $\triangle P_1QP_2$ erneut den Satz des Pythagoras an und erhält $|P_1P_2|^2 = |P_1Q|^2 + (z_2 - z_1)^2 = (x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2$.

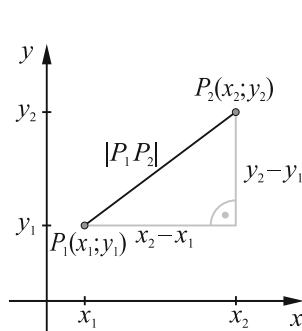


Abb. 2.5: Abstand zweier Punkte in der Ebene

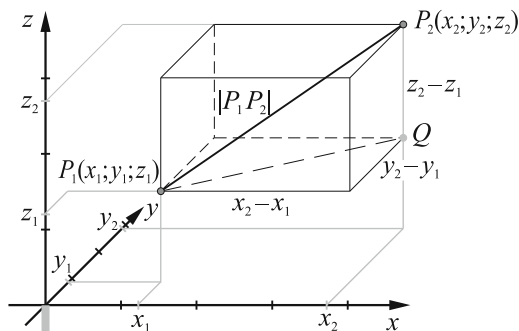


Abb. 2.6: Abstand zweier Punkte des Raumes

Für den Abstand zweier Punkte $P_1(x_1; y_1)$ und $P_2(x_2; y_2)$ der Ebene gilt:

$$|P_1P_2| = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}. \quad (2.5)$$

Der Abstand zweier Punkte $P_1(x_1; y_1; z_1)$, $P_2(x_2; y_2; z_2)$ des Raumes beträgt:

$$|P_1P_2| = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2}. \quad (2.6)$$

Mithilfe der Abstandsformel (2.5) sowie den in dem vorangegangenen Abschnitt behandelten Geradengleichungen und Schnittwinkeln lassen sich bereits verschiedene Aufgaben der ebenen Geometrie lösen und einige Sätze beweisen.

Satz 2.3

Der Mittelpunkt M einer Strecke $\overline{AB}$ mit $A(x_A; y_A)$ und $B(x_B; y_B)$ hat die Koordinaten $x_M = \frac{1}{2}(x_A + x_B)$ und $y_M = \frac{1}{2}(y_A + y_B)$.

Beweis: Es muss gezeigt werden, dass M auf der Geraden AB liegt und $|AM| = |BM|$ ist.

Um zu zeigen, dass $M \in AB$ gilt, stellen wir die Zweipunktegleichung (2.3) für diese Gerade auf, setzen $x = x_M$ und $y = y_M$ ein und formen äquivalent um:

$$\begin{aligned}\frac{y_M - y_A}{x_M - x_A} &= \frac{y_B - y_A}{x_B - x_A} \\ \frac{\frac{1}{2}(y_A + y_B) - y_A}{\frac{1}{2}(x_A + x_B) - x_A} &= \frac{y_B - y_A}{x_B - x_A} \\ \frac{\frac{1}{2}(y_B - y_A)}{\frac{1}{2}(x_B - x_A)} &= \frac{y_B - y_A}{x_B - x_A}.\end{aligned}$$

Dies ist offensichtlich eine wahre Aussage, also erfüllen x_M , y_M die für AB aufgestellte Geradengleichung, und M ist daher ein Punkt der Geraden AB .

Es bleibt zu prüfen, ob $|AM| = |BM|$ gilt:

$$\begin{aligned}|AM| &= \sqrt{\left(\frac{1}{2}(x_A + x_B) - x_A\right)^2 + \left(\frac{1}{2}(y_A + y_B) - y_A\right)^2} \\ &= \sqrt{\left(\frac{1}{2}(-x_A + x_B)\right)^2 + \left(\frac{1}{2}(-y_A + y_B)\right)^2} \\ &= \frac{1}{2}\sqrt{(x_B - x_A)^2 + (y_B - y_A)^2} = \frac{1}{2}|AB| \\ |BM| &= \sqrt{\left(\frac{1}{2}(x_A + x_B) - x_B\right)^2 + \left(\frac{1}{2}(y_A + y_B) - y_B\right)^2} \\ &= \frac{1}{2}\sqrt{(x_A - x_B)^2 + (y_A - y_B)^2} = \frac{1}{2}|AB|.\end{aligned}$$

Offensichtlich gilt $|AM| = |BM|$. Da bereits gezeigt wurde, dass M auf der Geraden AB liegt, ist M der Mittelpunkt der Strecke $\overline{AB}$. $\square$

Beispiel 2.3

Gegeben ist das Dreieck $\triangle ABC$ mit $A(-8; 0)$, $B(12; -6)$ und $C(8; 12)$. Gesucht sind die Gleichungen der Seitenhalbierenden und die Koordinaten ihres Schnittpunktes.

- Die Koordinaten der Seitenmittelpunkte des Dreiecks $\triangle ABC$ sind nach Satz 2.3: $M_{AB}(2; -3)$, $M_{BC}(10; 3)$ und $M_{AC}(0; 6)$.
- Aufstellen von Geradengleichungen für die Seitenhalbierenden AM_{BC} , BM_{AC} und CM_{AB} in der Zweipunkteform (2.3) und Umformung:

$$AM_{BC} : \frac{y - y_A}{x - x_A} = \frac{y_{M_{BC}} - y_A}{x_{M_{BC}} - x_A} \rightarrow \frac{y - 0}{x + 8} = \frac{3 - 0}{10 + 8} \rightarrow y = \frac{1}{6}x + \frac{4}{3}$$

$$BM_{AC} : \frac{y - y_B}{x - x_B} = \frac{y_{M_{AC}} - y_B}{x_{M_{AC}} - x_B} \rightarrow \frac{y + 6}{x - 12} = \frac{6 + 6}{0 - 12} \rightarrow y = -x + 6$$

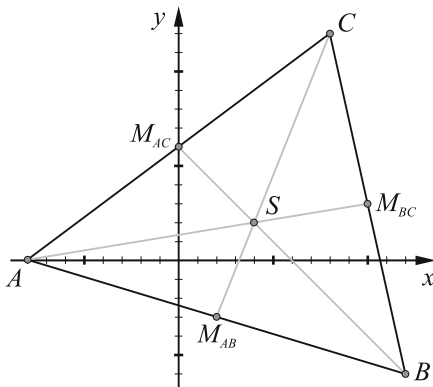


Abb. 2.7: Schnittpunkt der Seitenhalbierenden eines Dreiecks

$$CM_{AB} : \frac{y - y_C}{x - x_C} = \frac{y_{M_{AB}} - y_C}{x_{M_{AB}} - x_C} \rightarrow \frac{y - 12}{x - 8} = \frac{-3 - 12}{2 - 8} \rightarrow y = \frac{5}{2}x - 8$$

- Berechnung der Schnittpunkte jeweils zweier der drei Seitenhalbierenden durch Lösen der entsprechenden LGS:

$$\begin{array}{rcl} y = \frac{1}{6}x + \frac{4}{3} & y = \frac{1}{6}x + \frac{4}{3} & y = -x + 6 \\ y = -x + 6 & y = \frac{5}{2}x - 8 & y = \frac{5}{2}x - 8 \\ \hline x = 4, y = 2 & x = 4, y = 2 & x = 4, y = 2 \end{array}$$

Da bekannt ist, dass sich die drei Seitenhalbierenden eines Dreiecks in einem Punkt (dem Schwerpunkt) schneiden, hätte es auch gereicht, nur eines der Gleichungssysteme zu lösen. Es sei darauf hingewiesen, dass die Koordinaten des Schwerpunktes $S(4; 2)$ die arithmetischen Mittelwerte der Koordinaten der Eckpunkte des Dreiecks sind. ■

Beispiel 2.4

Ermittlung des Abstandes des Punktes $P(1; 3)$ von der Geraden $g: 4y - 3x = -2$.

Der Abstand eines Punktes P von einer Geraden g ist die Länge des Lotes von P auf g (Abb. 2.8). Um diese zu bestimmen, müssen eine Gleichung der durch P verlaufenden, zu g orthogonalen Geraden h aufgestellt, der Schnittpunkt L von h und g ermittelt sowie schließlich der Abstand $|PL|$ berechnet werden.

- Ermittlung einer Gleichung der zu g senkrechten Geraden h , der P angehört:

Der Anstieg von g ist $m_g = \frac{3}{4}$, nach Satz 2.2 muss h den Anstieg $m_h = -\frac{4}{3}$ haben. Nach Satz 2.1 gilt also für jeden Punkt $Q(x; y)$ von h :

$$\frac{y - y_P}{x - x_P} = -\frac{4}{3} \Leftrightarrow \frac{y - 3}{x - 1} = -\frac{4}{3} \Leftrightarrow y = -\frac{4}{3}x + \frac{13}{3} \quad \text{bzw.} \quad 3y + 4x = 13.$$

- Um den Schnittpunkt von g und h zu ermitteln, wird das LGS gelöst, welches aus den Gleichungen von g und h besteht:

$$\begin{array}{rcl} 4y - 3x & = & -2 \\ 3y + 4x & = & 13. \end{array}$$

Es ergibt sich $x = \frac{58}{25}, y = \frac{31}{25}$, also ist $L(\frac{58}{25}; \frac{31}{25})$ der gesuchte Schnittpunkt.

- Der Abstand von P zu g berechnet sich als Abstand der Punkte P und L :

$$d(P, g) = |PL| = \sqrt{(x_L - x_P)^2 + (y_L - y_P)^2} = \sqrt{\left(\frac{58}{25} - 1\right)^2 + \left(\frac{31}{25} - 3\right)^2} = 2,2. \quad \blacksquare$$

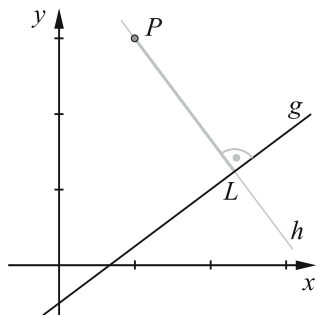


Abb. 2.8: Abstand eines Punktes von einer Geraden

2.2.2 Kreise

Ein Kreis ist die Menge aller Punkte der Ebene, die von einem festen Punkt M denselben Abstand r haben. Dabei wird M als Mittelpunkt des Kreises und r als Radius bezeichnet. Unter Verwendung der Gleichung (2.5) für den Abstand zweier Punkte der Ebene (siehe S. 55) ergeben sich aus dieser Definition Gleichungen für Kreise in Mittelpunktslage (falls M der Koordinatenursprung ist) und für Kreise mit beliebigen Mittelpunkten (siehe Abb. 2.9).

Der Kreis mit dem Mittelpunkt $O(0;0)$ und dem Radius r ist die Menge aller Punkte $P(x;y)$ der Ebene, für die folgende Gleichung erfüllt ist:

$$x^2 + y^2 = r^2. \quad (2.7)$$

Der Kreis mit dem Mittelpunkt $M(x_M; y_M)$ und dem Radius r ist die Menge aller Punkte $P(x;y)$ der Ebene, für die folgende Gleichung erfüllt ist:

$$(x-x_M)^2 + (y-y_M)^2 = r^2. \quad (2.8)$$

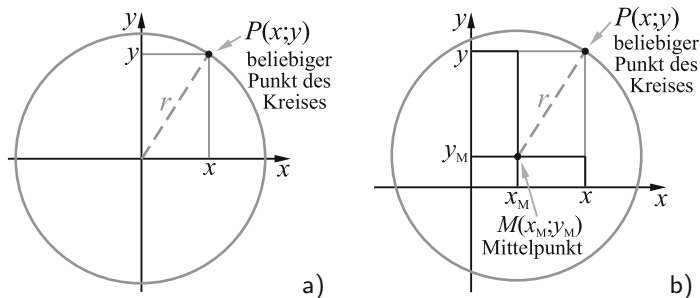


Abb. 2.9: Kreise in Mittelpunktslage und in allgemeiner Lage

Darstellung von Kreisen mithilfe des Computeralgebrasystems Maxima

Bei den folgenden Überlegungen und den zugehörigen Aufgaben ist oft die graphische Darstellung von Kreisen, mitunter zusammen mit Geraden, sinnvoll. Hierfür kann z. B. das CAS Maxima genutzt werden, welches in dem Abschnitt 1.5 vorgestellt wurde. Die Darstellung von Kreisen, die durch Gleichungen gegeben sind, erfolgt völlig analog zu Geraden mithilfe des `implicit`-Objekts. Die folgende Eingabe stellt einen Kreis `Kr1`, eine Gerade `G1` und den Mittelpunkt des Kreises (mithilfe der `points`-Eingabe) dar, siehe Abb. 2.10:

```
Kr1: (x-1.5)^2 + (y+1)^2 = 16;
G1: 3*x + 4*y = -5;
load("draw")$
draw2d( user_preamble = ["set size ratio 1"], ip_grid=[200,200],
        points ([[1.5,-1]]),
        color = blue, implicit(Kr1, x,-6,6, y,-6,6),
        color = red , implicit(G1, x,-6,6, y,-6,6) );
```

Die Anweisung `ip_grid=[200,200]` bewirkt eine feinere Darstellung der Objekte durch Berechnung von 200 Punkten in jeder Richtung (Standard sind 50 Punkte). Die weiteren Anweisungen wurden bereits auf S. 47 beschrieben.

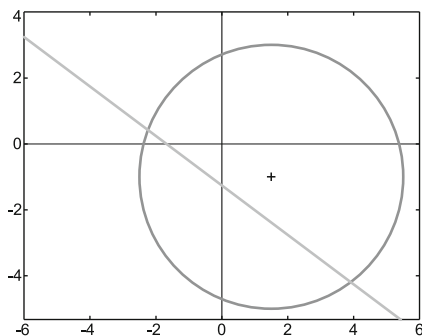


Abb. 2.10: Darstellung eines Kreises und einer Geraden mithilfe des CAS Maxima

Kreise sind keine Graphen von Funktionen, da die Zuordnung $x \mapsto y$ nicht eindeutig ist (es gibt jeweils zwei Punkte eines Kreises mit derselben x -Koordinate). Jedoch lassen sich Halbkreise als Funktionsgraphen darstellen. Dazu wird die Kreisgleichung (2.8) nach y umgestellt:

$$\begin{aligned}(y - y_M)^2 &= -(x - x_M)^2 + r^2 \\ y_{1/2} - y_M &= \pm \sqrt{r^2 - (x - x_M)^2} \\ y_{1/2} &= \pm \sqrt{r^2 - (x - x_M)^2} + y_M\end{aligned}$$

Die Halbkreise werden also durch die Graphen der Funktionen f_1 und f_2 mit

$$f_1(x) = \sqrt{r^2 - (x - x_M)^2} + y_M \quad \text{und} \quad f_2(x) = -\sqrt{r^2 - (x - x_M)^2} + y_M \quad (2.9)$$

dargestellt. Der Definitionsbereich dieser Funktionen ist jeweils $[x_M - r; x_M + r]$, die Punkte des Kreises mit den Koordinaten $(x_M - r; y_M)$ und $(x_M + r; y_M)$ gehören beiden Funktionsgraphen an.

Mittels der Beschreibung von Halbkreisen durch Funktionsgraphen ist eine schnellere Berechnung graphischer Darstellungen von Kreisen in dem CAS Maxima möglich als anhand der impliziten Beschreibung durch Kreisgleichungen. Um den zuvor erzeugten Kreis (siehe Abb. 2.10) durch zwei Halbkreise darzustellen, wird in dem auf S. 58 angegebenen Maxima-Code die erste Zeile durch

`f1: sqrt(r^2 - (x-1.5)^2) - 1;    f2: -sqrt(r^2 - (x-1.5)^2) - 1;`

sowie die `implicit`-Anweisung für den Kreis durch

`explicit(f1, x, xM-r, xM+r),    explicit(f2, x, xM-r, xM+r),`

ersetzt. Damit erhält man eine deutlich feinere Darstellung des Kreises bei gleichzeitig kürzerer Rechenzeit. Die Ursache hierfür besteht darin, dass die Software bei explizit gegebenen Funktionstermen lediglich zu (genügend vielen) x -Werten die zugehörigen y -Werte berechnen muss, während bei der `implicit`-Anweisung für jeden Stützpunkt die (quadratische) Kreisgleichung gelöst wird.

Bestimmung einer Gleichung eines Kreises, von dem drei Punkte gegeben sind

Durch drei nicht auf einer Geraden liegende Punkte A , B und C verläuft genau ein Kreis, nämlich der Umkreis des Dreiecks $\triangle ABC$. Der Mittelpunkt dieses

Kreises ist der Schnittpunkt der drei Mittelsenkrechten des Dreiecks $\triangle ABC$, siehe Abb. 2.11 (vgl. hierzu auch die Aufgabe 1 auf S. 65). Um den Mittelpunkt und den Radius eines durch drei gegebene Punkte A, B, C verlaufenden Kreises zu bestimmen, lassen sich also Gleichungen mindestens zweier Mittelsenkrechten des Dreiecks $\triangle ABC$ aufstellen, ihr Schnittpunkt bestimmen sowie dessen Abstand zu einem der Punkte A, B, C berechnen (die Abstände des Schnittpunktes der Mittelsenkrechten zu jedem der Eckpunkte sind gleich). Im Folgenden wird nun eine Möglichkeit aufgezeigt, Mittelpunkt, Radius und eine Gleichung eines durch drei Punkte gegebenen Kreises auf kürzerem Wege zu bestimmen.

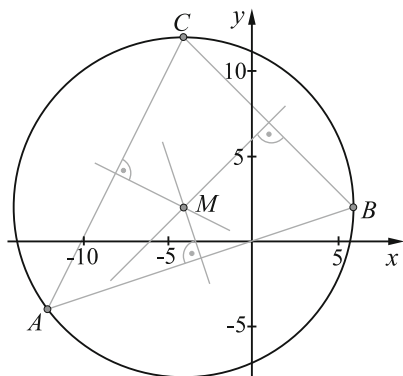


Abb. 2.11: Konstruktion eines Kreises durch drei Punkte

Beispiel 2.5

Bestimmung der Gleichung des Kreises k , der durch die Punkte $A(-12; -4)$; $B(6; 2)$ und $C(-4; 12)$ verläuft (siehe Abb. 2.11).

Zu bestimmen sind der Radius r und die Koordinaten x_M, y_M des Mittelpunktes M von k . Dazu setzen wir die Koordinaten der Punkte A, B und C in die Gleichung des Kreises ein und erhalten das quadratische Gleichungssystem

$$\begin{aligned} \text{I: } & (-12 - x_M)^2 + (-4 - y_M)^2 = r^2 \\ \text{II: } & (6 - x_M)^2 + (2 - y_M)^2 = r^2 \\ \text{III: } & (-4 - x_M)^2 + (12 - y_M)^2 = r^2 \quad \text{bzw.} \\ \text{I: } & x_M^2 + y_M^2 + 24x_M + 8y_M = r^2 - 160 \\ \text{II: } & x_M^2 + y_M^2 - 12x_M - 4y_M = r^2 - 40 \\ \text{III: } & x_M^2 + y_M^2 + 8x_M - 24y_M = r^2 - 160. \end{aligned}$$

Durch Subtraktion der Gleichungen II bzw. III von I ergibt sich ein lineares Gleichungssystem:

$$\begin{aligned} \text{IV: } & 36x_M + 12y_M = -120 & \text{bzw.} & \text{IV: } 3x_M + y_M = -10 \\ \text{V: } & 16x_M + 32y_M = 0 & & \text{V: } x_M + 2y_M = 0. \end{aligned}$$

Die Lösung dieses LGS liefert die Mittelpunktskoordinaten $x_M = -4$ und $y_M = 2$. Durch Einsetzen in I erhalten wir daraus $r^2 = 100$ bzw. $r = 10$. Eine Gleichung des Kreises k ist also

$$k: (x + 4)^2 + (y - 2)^2 = 100. \quad \blacksquare$$

Kreisgleichungen der Form $ax^2 + ay^2 + bx + cy = d$

Nicht immer liegen Kreisgleichungen in der Form $(x-x_M)^2 + (y-y_M)^2 = r^2$ vor, aus der direkt die Koordinaten des Mittelpunktes und der Radius ersichtlich sind. So ließe sich z. B. die Gleichung $(x-2)^2 + (y-3)^2 = 25$ des in dem Beispiel 2.7 betrachteten Kreises auch in der Form $x^2 + y^2 - 4x - 6y = 12$ angeben. Sind Gleichungen für Kreise in dieser oder ähnlicher Form gegeben, so lassen sich weder die Mittelpunktskoordinaten noch der Radius unmittelbar entnehmen. Es ist noch nicht einmal sicher, dass durch eine derartige Gleichung tatsächlich ein Kreis beschrieben wird. Oft ist es daher notwendig, Gleichungen der Gestalt $ax^2 + ay^2 + bx + cy = d$ in die Form $(x-x_M)^2 + (y-y_M)^2 = r^2$ (2.8) zu bringen.

Beispiel 2.6

Die folgenden Gleichungen werden daraufhin untersucht, ob sie Kreise beschreiben; ist dies der Fall, so werden deren Mittelpunkte und Radien ermittelt.

$$(a) \quad x^2 + y^2 - 14x + 10y = -38$$

$$(b) \quad x^2 + y^2 + 8x - 6y = -29$$

Es wird jeweils eine quadratische Ergänzung der Terme durchgeführt, die x bzw. y enthalten:

$$x^2 - 14x = (x-7)^2 - 49$$

$$y^2 + 10y = (y+5)^2 - 25$$

Durch Einsetzen in (a) ergibt sich:

$$(x-7)^2 - 49 + (y+5)^2 - 25 = -38.$$

bzw.

$$(x-7)^2 + (y+5)^2 = 36.$$

Durch die Gleichung (a) wird somit ein Kreis mit dem Mittelpunkt $M(7; -5)$ und dem Radius 6 beschrieben.

$$x^2 + 8x = (x+4)^2 - 16$$

$$y^2 - 6y = (y-3)^2 - 9$$

Durch Einsetzen in (b) ergibt sich:

$$(x+4)^2 - 16 + (y-3)^2 - 9 = -29.$$

bzw.

$$(x+4)^2 + (y-3)^2 = -4.$$

Diese Gleichung besitzt keine Lösung, da die Summe zweier Quadrate nicht negativ sein kann. Durch (b) wird daher kein Kreis beschrieben. ■

2.2.3 Lagebeziehungen von Kreisen und Geraden

Aus der geometrischen Anschauung heraus ist klar, dass ein Kreis und eine Gerade keinen, genau einen oder zwei Schnittpunkte besitzen können. Eine Gerade, die einen Kreis k nicht schneidet, heißt *Passante*; eine Gerade, die mit k genau einen Punkt gemeinsam hat, *Tangente* sowie eine Gerade, die mit k zwei Schnittpunkte besitzt, *Sekante* des Kreises k (siehe Abb. 2.12).

Beispiel 2.7

Ermittlung der gegenseitigen Lage und Bestimmung der Schnittpunkte des Kreises $k: (x-2)^2 + (y-3)^2 = 25$ und der Geraden $g_3: 4y - 3x = 12$.

Falls der Kreis k und die Gerade g_3 einen oder mehrere Schnittpunkte haben, so müssen deren Koordinaten sowohl die Kreis- als auch die Geradengleichung

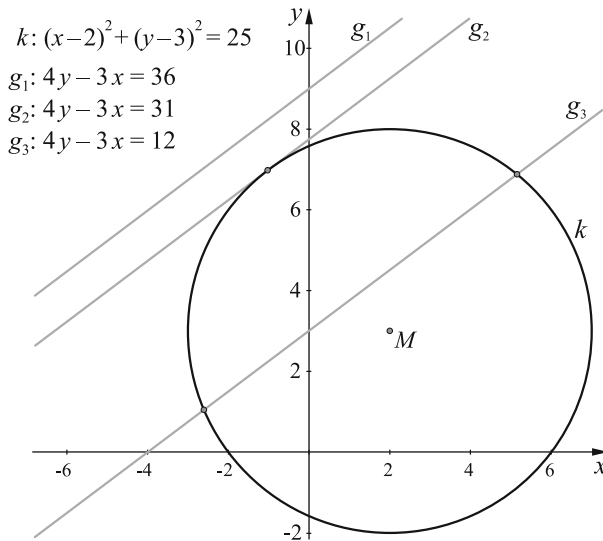


Abb. 2.12: Passante, Tangente und Sekante eines Kreises

erfüllen, also Lösungen des aus einer quadratischen und einer linearen Gleichung bestehenden Gleichungssystems

$$\begin{aligned}(x-2)^2 + (y-3)^2 &= 25 \\ 4y - 3x &= 12\end{aligned}$$

sein. Um dieses Gleichungssystem zu lösen, lässt sich die Geradengleichung nach y oder x umstellen, z. B. $y = \frac{3}{4}x + 3$. Einsetzen in die Kreisgleichung ergibt dann

$$\begin{aligned}(x-2)^2 + \left(\frac{3}{4}x + 3 - 3\right)^2 &= 25 \\ x^2 + \frac{9}{16}x^2 - 4x + 4 &= 25 \\ 25x^2 - 64x &= 336 \\ x^2 - \frac{64}{25}x - \frac{336}{25} &= 0 \\ x_{1/2} &= \frac{32}{25} \pm \sqrt{\frac{32^2}{25^2} + \frac{336}{25}} \approx 1,28 \pm 3,88 \\ x_1 &\approx 5,16, \quad x_2 \approx -2,60.\end{aligned}$$

Durch Einsetzen in die ursprüngliche Geraden- oder Kreisgleichung erhält man die y -Koordinaten der beiden Schnittpunkte. Der Kreis k und die Gerade g_3 schneiden sich in den Punkten mit den näherungsweisen Koordinaten $(5,16; 6,87)$ und $(-2,60; 1,05)$, somit ist g_3 eine Sekante von k .

Zur Ermittlung der gegenseitigen Lage des Kreises k und der Geraden g_1 (siehe Abb. 2.12) ist das Gleichungssystem

$$\begin{aligned}(x-2)^2 + (y-3)^2 &= 25 \\ 4y - 3x &= 36\end{aligned}$$

zu lösen. Mit denselben Umformungsschritten wie oben ergibt sich hierbei

$$\begin{aligned}25x^2 + 80x &= -240 \\ x^2 + \frac{80}{25}x + \frac{240}{25} &= 0 \\ x_{1/2} &= -\frac{40}{25} \pm \sqrt{\frac{40^2}{25^2} - \frac{240}{25}} = -\frac{40}{25} \pm \sqrt{\frac{1600}{25^2} - \frac{6000}{25^2}}.\end{aligned}$$

Offensichtlich existiert aufgrund des negativen Ausdrucks unter der Wurzel keine Lösung; k und g_1 besitzen also keinen gemeinsamen Punkt.

Auf völlig analoge Weise wird schließlich die gegenseitige Lage des Kreises k und der Geraden g_2 ermittelt. Dabei ergibt sich

$$\begin{aligned} 25x^2 + 50x &= -25 \\ x^2 + 2x + 1 &= 0 \\ x_{1/2} &= -1 \pm \sqrt{1-1} = -1. \end{aligned}$$

Es existiert nur eine Lösung; durch Einsetzen erhält man $y_1 = 7$. Somit besitzen k und g_2 genau einen Schnittpunkt $S(-1; 7)$; g_2 ist also Tangente an k . ■

Berechnung von Schnittpunkten mithilfe eines Computeralgebrasystems

Da die für die Ermittlung von Lagebeziehungen zwischen Kreisen und Geraden sowie für Schnittpunktbestimmungen durchzuführenden Rechnungen zeit-aufwändig sind, empfiehlt es sich, hierfür ein Computeralgebrasystem zu verwenden. Die in dem Abschnitt 1.5 beschriebene `solve`-Anweisung des CAS Maxima lässt sich nicht nur für das Lösen linearer Gleichungssysteme, sondern u. a. auch für Systeme nutzen, die aus einer quadratischen und einer linearen Gleichung bestehen. Dabei sind zwei Eigenschaften von CAS zu beachten.

- CAS ermitteln standardmäßig exakte Lösungen. Sind die Lösungen einer Gleichung oder eines Gleichungssystems irrational (wie in dem Beispiel 2.7 für die Schnittpunkte des Kreises k und der Geraden g_3), so werden als Lösungen oft recht komplizierte Wurzelausdrücke angegeben. Oft möchte man jedoch statt dessen Näherungswerte erhalten. Dies wird in Maxima durch die Option `numer` der `solve`-Anweisung erreicht. Durch die Eingabe

```
solve([(x-2)^2+(y-3)^2=25, 4*y-3*x=12],[x,y]), numer;
```

erhält man Näherungswerte der Koordinaten der beiden Schnittpunkte:

```
[x=-2.60309155, y=1.04768133], [x=5.16309155, y=6.87231866]]
```

- Quadratische Gleichungen, die in den reellen Zahlen unlösbar sind (wie in dem Beispiel 2.7 für k und g_1) besitzen Lösungen in den komplexen Zahlen. Diese Lösungen, welche die imaginäre Einheit i (in Maxima `%i`) enthalten, werden von CAS standardmäßig angegeben. So führt die Eingabe

```
solve([(x-2)^2+(y-3)^2=25, 4*y-3*x=36],[x,y]), numer;
```

zu den Lösungen

```
[x=-0.2*(13.266499161*%i+8), y=-0.2*(9.94987437106*%i-39)],  
[x=0.2*(13.266499161*%i-8), y=0.2*(9.94987437106*%i+39)]]
```

Wenn (wie hier bei Schnittpunktbestimmungen) nur reelle Lösungen von Interesse sind, so lässt sich mithilfe der Option `realonly:true` erreichen, dass die Ausgabe komplexer Lösungen mit von Null verschiedenem Imaginärteil unterdrückt wird. Als Ergebnis der Eingabe

```
solve([(x-2)^2+(y-3)^2=25, 4*y-3*x=36],[x,y]), realonly:true;
```

gibt Maxima

```
[ ]
```

aus, es existieren also keine reellen Lösungen.

Tangenten an Kreise

Es gibt verschiedene Zugänge, um eine Gleichung der Tangente t an einen Kreis k in einem gegebenen Punkt P_0 des Kreises zu bestimmen:

- Eine Tangente an einen Kreis ist stets senkrecht zu der Geraden MP_0 durch dessen Mittelpunkt und den Berührungspunkt (siehe Abb. 2.13). Damit kann der Tangentenanstieg bestimmt werden.
- Eine Tangente an einen Kreis hat mit diesem Kreis stets genau einen Punkt gemeinsam. Dies ist genau dann der Fall, wenn das aus der Kreis- und der Geradengleichung bestehende Gleichungssystem genau eine Lösung besitzt.
- Halbkreise lassen sich als Funktionsgraphen auffassen. Durch Ableiten der Funktionsgleichungen können Tangentenanstiege ermittelt werden.

Beispiel 2.8

Ermittlung einer Gleichung der Tangente an den Kreis $k: (x-5)^2 + (y-1)^2 = 3,25^2$ in dem Punkt $P_0(3,75; 4)$.

Der Anstieg m_0 der Geraden durch M und P_0 ist (nach Satz 2.1 auf S. 51):

$$m_0 = \frac{y_0 - y_M}{x_0 - x_M} = \frac{4 - 1}{3,75 - 5} = \frac{3}{-1,25} = -\frac{12}{5}.$$

Für den Anstieg m der Tangente t an k in P_0 gilt, da diese senkrecht auf der Geraden MP_0 stehen muss, nach Satz 2.2 auf S. 54:

$$m = -\frac{1}{m_0} = \frac{5}{12}.$$

Eine Gleichung der Tangente ist damit $y = \frac{5}{12}x + n$, wobei der Achsenabschnitt n noch unbekannt ist. Er lässt sich durch Einsetzen der Werte $x_0 = 3,75 = \frac{15}{4}$ und $y_0 = 4$ ermitteln. Aus $4 = \frac{5}{12} \cdot \frac{15}{4} + n$ folgt $n = \frac{39}{16}$. Eine Gleichung der Tangente an k in P_0 ist somit

$$t: y = \frac{5}{12}x + \frac{39}{16}$$

Es lässt sich überprüfen, dass das Gleichungssystem, welches aus der Kreisgleichung und dieser Tangentengleichung besteht, genau eine Lösung besitzt (nämlich die Koordinaten des Punktes P_0), siehe Aufgabe 13 auf S. 66. ■

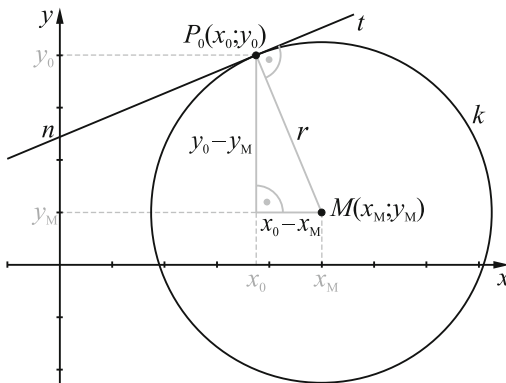


Abb. 2.13: Tangente an einen Kreis

Das in dem Beispiel 2.8 praktizierte Vorgehen lässt sich verallgemeinern, um die Gleichung der Tangente an einen beliebigen Kreis $k: (x-x_M)^2 + (y-y_M)^2 = r^2$ in einem beliebigen Punkt $P_0(x_0; y_0)$ dieses Kreises herzuleiten (mit Ausnahme der zu den Koordinatenachsen parallelen Tangenten mit $x_0 = x_M$ bzw. $y_0 = y_M$). Für den Anstieg m_0 der Geraden MP_0 gilt nach dem Satz 2.1 auf S. 51:

$$m_0 = \frac{y_0 - y_M}{x_0 - x_M}.$$

Für den Anstieg m der Tangente t an k in P_0 folgt nach dem Satz 2.2 auf S. 54:

$$m = -\frac{x_0 - x_M}{y_0 - y_M}.$$

Setzt man diesen Anstieg sowie die Koordinaten x_0 für x und y_0 für y in die Form $y = mx + n$ der Geradengleichung ein, so lässt sich n ermitteln:

$$n = y_0 + \frac{x_0 - x_M}{y_0 - y_M} x_0.$$

Eine Gleichung der Tangente an k in P_0 ist also

$$t: y = -\frac{x_0 - x_M}{y_0 - y_M} x + y_0 + \frac{x_0 - x_M}{y_0 - y_M} x_0 = -\frac{x_0 - x_M}{y_0 - y_M} (x - x_0) + y_0. \quad (2.10)$$

Bestimmung von Kreistangenten mithilfe der Differentialrechnung

Mittels der Auffassung von Halbkreisen als Funktionsgraphen (siehe (2.9) auf S. 59) lässt sich die Differentialrechnung nutzen, um den Anstieg von Tangenten an Kreise zu bestimmen. Leitet man die Funktionsgleichung

$$f_1(x) = \sqrt{r^2 - (x - x_M)^2} + y_M$$

des „oberen“ Halbkreises nach x ab, so erhält man

$$f'_1(x) = \frac{1}{2\sqrt{r^2 - (x - x_M)^2}} \cdot (-2x + 2x_M) = \frac{-(x - x_M)}{\sqrt{r^2 - (x - x_M)^2}} = -\frac{x - x_M}{f_1(x) - y_M}.$$

Setzt man hierin $x = x_0$ und $f(x_0) = y_0$, so ergibt sich wiederum der Tangentenanstieg

$$m = -\frac{x_0 - x_M}{y_0 - y_M}.$$

Die Tangentengleichung (2.10) lässt sich daraus wie oben gezeigt ableiten.

2.2.4 Aufgaben zu Abschnitt 2.2

- Bestimmen Sie die Koordinaten des Schnittpunktes S der Mittelsenkrechten des Dreiecks $\triangle ABC$ mit $A(-11; 1)$, $B(6; 8)$ und $C(13; -9)$. Beachten Sie zum Aufstellen der Gleichungen der Mittelsenkrechten die Aufgabe 9 auf S. 54.
- Weisen Sie nach, dass die Menge M aller Punkte $P(x; y)$, die von zwei gegebenen Punkten $A(x_A; y_A)$ und $B(x_B; y_B)$ jeweils gleich weit entfernt sind, eine Gerade bilden, die durch den Mittelpunkt der Strecke $\overline{AB}$ verläuft und zu der Geraden AB orthogonal ist (die Mittelsenkrechte der Strecke $\overline{AB}$).
- Bestimmen Sie jeweils den Abstand des Punktes P von der Geraden g .
 - $P(1; 2)$, $g: 2y + 1, 5x = -11$
 - $P(0,6; -1,5)$, $g: 3y - 2x = 4$

4. Geben Sie jeweils eine Gleichung des Kreises mit dem Mittelpunkt M an, der durch den Punkt P_0 verläuft.
 - a) $M(0; 0)$, $P_0(2; 1,5)$
 - b) $M(-2, 1)$, $P_0(4; -3)$
5. Ermitteln Sie, welche der folgenden Punkte innerhalb, außerhalb oder auf dem Kreis $k: (x-3)^2 + (y+4)^2 = 36$ liegen: $P(6; 1)$, $Q(9; -4)$, $R(5; 2)$, $S(8; 8)$.
6. Der Kreis k berührt die x -Achse im Punkt $P(4; 0)$ und die y -Achse im Punkt $Q(0; -4)$. Bestimmen Sie den Mittelpunkt und den Radius des Kreises k . Fertigen Sie dazu zunächst eine Skizze an.
7. Wie viele Kreise mit dem Radius $r = \sqrt{50}$ existieren, die durch die Punkte $P(3; 5)$ und $Q(-1; -7)$ verlaufen? Geben Sie Gleichungen dieser Kreise an.
8. Bestimmen Sie den Mittelpunkt und den Radius des Kreises, der durch die Punkte A , B und C verläuft. Geben Sie eine Gleichung dieses Kreises an.
 - a) $A(1; 1)$, $B(7; -7)$, $C(8; 0)$
 - b) $A(10; -9)$, $B(-4; 5)$, $C(-8; -3)$
 - c) $A(0; 2)$, $B(4; 8)$, $C(-6; 6)$
 - d) $A(0; -1)$, $B(2; 2)$, $C(-3; 1)$
9. Untersuchen Sie, ob die folgenden Gleichungen Kreise beschreiben. Ermitteln Sie in den Fällen, in denen dies zutrifft, deren Mittelpunkte und Radien.
 - a) $x^2 + y^2 - 8x + 6y = -21$
 - b) $x^2 + y^2 - 22x + 2y = -138$
 - c) $x^2 + y^2 + 10x + 6y = -9$
 - d) $x^2 + y^2 - 4x - 14y = -54$
10. Bestimmen Sie eine Gleichung des zu $k: 5x^2 + 5y^2 + 24x - 32y = 9$ konzentrischen Kreises, der
 - a) die x -Achse berührt,
 - b) durch den Punkt $P(1; 2)$ verläuft.
11. Untersuchen Sie jeweils die gegenseitige Lage der Geraden g und des Kreises k . Bestimmen Sie gemeinsame Punkte von g und k , falls diese existieren.
 - a) $g: y = 2x - 5$ $k: x^2 + y^2 = 25$
 - b) $g: -7x + 4y = 15$ $k: (x-2)^2 + (y-5)^2 = 100$
 - c) $g: y = 4$ $k: (x-3)^2 + (y+1)^2 = 25$
 - d) $g: -9x + 4y = -12$ $k: (x-7)^2 + (y+6)^2 = 49$
12. Bestimmen Sie, falls vorhanden, die Schnittpunkte der Kreise k_1 und k_2 :
 - a) $k_1: x^2 + y^2 = 25$ $k_2: (x+1)^2 + (y-3)^2 = 9$
 - b) $k_1: x^2 + y^2 + 2x - 4y - 4 = 0$ $k_2: 4x^2 + 4y^2 + 12x - 12y - 3 = 0$
13. Verifizieren Sie, dass die in dem Beispiel 2.8 auf S. 64 ermittelte Tangente t und der gegebene Kreis k tatsächlich genau einen gemeinsamen Punkt besitzen. Zeigen Sie dazu, dass das Gleichungssystem, das aus der Kreisgleichung $k: (x-5)^2 + (y-1)^2 = 3,25^2$ und der Geradengleichung $t: y = \frac{5}{12}x + 2,4375$ besteht, eindeutig lösbar ist.
14. Leiten Sie mithilfe der Differentialrechnung eine Tangentengleichung für den „unteren Halbkreis“ (siehe S. 59) her.
15. Bestimmen Sie Gleichungen der Tangenten in den Punkten $P_1(5; 7)$ und $P_2(6; 0)$ an den Kreis k mit dem Mittelpunkt $M(2; 3)$ und dem Radius $r = 5$.

2.3 Kugeln und Kegel

2.3.1 Kugelgleichungen

Analog zu einem Kreis in der Ebene ist eine Kugel die Menge aller Punkte des Raumes, die von einem vorgegebenen Punkt M einen festen Abstand r haben. M wird als Mittelpunkt der Kugel und r als ihr Radius bezeichnet. Aufgrund der Gleichung (2.6) für den Abstand zweier Punkte des Raumes (vgl. S. 55) ergeben sich aus dieser Definition Gleichungen für Kugeln in Mittelpunktslage (falls M der Koordinatenursprung ist, siehe Abb. 2.14) und für Kugeln mit beliebigen Mittelpunkten (siehe Abb. 2.15).

Die Kugel mit dem Mittelpunkt $O(0; 0; 0)$ und dem Radius r ist die Menge aller Punkte $P(x; y; z)$ des Raumes, für die folgende Gleichung erfüllt ist:

$$x^2 + y^2 + z^2 = r^2. \quad (2.11)$$

Die Kugel mit dem Mittelpunkt $M(x_M; y_M; z_M)$ und dem Radius r ist die Menge aller Punkte $P(x; y; z)$, für die folgende Gleichung erfüllt ist:

$$(x - x_M)^2 + (y - y_M)^2 + (z - z_M)^2 = r^2. \quad (2.12)$$

Bestimmung einer Gleichung einer Kugel, von der vier Punkte gegeben sind

Um eine Kugel im Raum eindeutig zu bestimmen, werden vier Punkte benötigt, die nicht in einer Ebene liegen. Ähnlich wie in der Ebene jedes Dreieck einen eindeutig bestimmten Umkreis besitzt, gibt es für jeden Körper mit vier Ecken (d. h. für jedes Tetraeder) genau eine „Umkugel“ (siehe Abb. 2.16).

Beispiel 2.9

Aufstellen einer Gleichung der Kugel k durch die Punkte $P(4; 4; 1)$, $Q(3; 2; 2)$, $R(0; 0; -3)$ und $S(6; 4; -1)$.

Zu bestimmen sind der Radius r sowie die Koordinaten x_M , y_M und z_M des Mittelpunktes M von k . Dazu setzen wir die Koordinaten eines jeden der vier

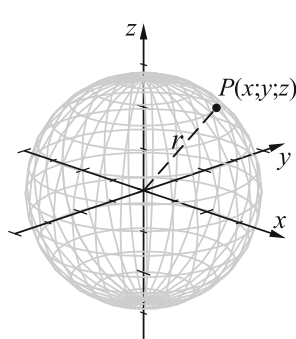


Abb. 2.14: Kugel in Mittelpunktslage

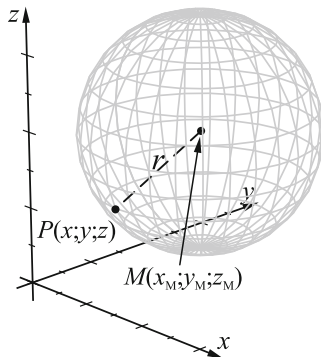


Abb. 2.15: Kugel in allgemeiner Lage

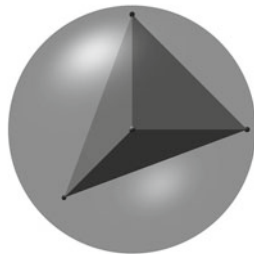


Abb. 2.16: Umkugel eines Tetraeders

gegebenen Punkte in die Gleichung (2.12) ein und erhalten folgendes System quadratischer Gleichungen:

$$\begin{aligned} \text{I: } & (4-x_M)^2 + (4-y_M)^2 + (1-z_M)^2 = r^2 \\ \text{II: } & (3-x_M)^2 + (2-y_M)^2 + (2-z_M)^2 = r^2 \\ \text{III: } & x_M^2 + y_M^2 + (-3-z_M)^2 = r^2 \\ \text{IV: } & (6-x_M)^2 + (4-y_M)^2 + (-1-z_M)^2 = r^2. \end{aligned}$$

Durch Anwendung der binomischen Formeln und Sortieren ergibt sich daraus:

$$\begin{aligned} \text{I: } & x_M^2 + y_M^2 + z_M^2 - 8x_M - 8y_M - 2z_M = r^2 - 33 \\ \text{II: } & x_M^2 + y_M^2 + z_M^2 - 6x_M - 4y_M - 4z_M = r^2 - 17 \\ \text{III: } & x_M^2 + y_M^2 + z_M^2 + 6z_M = r^2 - 9 \\ \text{IV: } & x_M^2 + y_M^2 + z_M^2 - 12x_M - 8y_M + 2z_M = r^2 - 53. \end{aligned}$$

Um die quadratischen Terme zu eliminieren, subtrahieren wir II, III und IV jeweils von der Gleichung I und erhalten das lineare Gleichungssystem:

$$\begin{aligned} \text{V: } & -2x_M - 4y_M + 2z_M = -16 & \text{V: } & x_M + 2y_M - z_M = 8 \\ \text{VI: } & -8x_M - 8y_M - 8z_M = -24 & \text{bzw. VI: } & x_M + y_M + z_M = 3 \\ \text{VII: } & 4x_M - 4z_M = 20 & \text{VII: } & x_M - z_M = 5 \end{aligned}$$

sowie schließlich durch Subtraktion der Gleichungen VI und VII von V:

$$\begin{aligned} \text{V: } & x_M + 2y_M - z_M = 8 \\ \text{VIII: } & y_M - 2z_M = 5 \\ \text{IX: } & 2y_M = 3. \end{aligned}$$

Die Lösung dieses linearen Gleichungssystems liefert die Mittelpunktskoordinaten $x_M = 3,25$, $y_M = 1,5$ und $z_M = -1,75$. Durch Einsetzen in I erhalten wir daraus $r^2 = 14,375$ bzw. $r \approx 3,791$. Eine Gleichung der Kugel k lautet also

$$k: (x-3,25)^2 + (y-1,5)^2 + (z+1,75)^2 = 14,375. \quad \blacksquare$$

Der rechnerische Aufwand für die Bestimmung einer Gleichung der Kugel durch vier vorgegebene Punkte ist, wie das Beispiel 2.9 verdeutlicht, enorm. Die *Verwendung eines Computeralgebrasystems* führt zu einer beachtlichen Zeiterparnis. Das ursprüngliche Gleichungssystem aus vier quadratischen Gleichungen lässt sich mithilfe des CAS Maxima problemlos lösen. Die Eingabe

```
G11: (4-xM )^2 + (4-yM )^2 + (1-zM )^2 = r^2;
G12: (3-xM )^2 + (2-yM )^2 + (2-zM )^2 = r^2;
G13: xM^2 + yM^2 + (-3-zM )^2 = r^2;
G14: (6-xM )^2 + (4-yM )^2 + (-1-zM )^2 = r^2;
solve([G11, G12, G13, G14], [xM, yM, zM, r]), numer;
```

führt unmittelbar zu den Werten für x_M , y_M , z_M und r :

```
[ [xM=3.25,yM=1.5,zM=-1.75,r=-3.791437722025774],
  [xM=3.25,yM=1.5,zM=-1.75,r=3.791437722025774]]
```

Davon ist natürlich nur die Lösung mit positivem r sinnvoll.

Auch die graphische Darstellung von Kugeln ist mithilfe von Maxima möglich, siehe eine entsprechende Datei auf der Internetseite zu diesem Buch. Leider erfolgt die Darstellung jedoch meist verzerrt; eine gleiche Skalierung aller Achsen ist bei dreidimensionalen Grafiken gegenwärtig (2010) in Maxima nicht möglich.

2.3.2 Kegel

Im allgemeinen Sprachgebrauch wird ein Kreiskegel oft als Körper aufgefasst, der von einer kreisförmigen Grundfläche, einer Spitze S und allen Verbindungsstrecken zwischen der Spitze und Punkten des Grundkreises begrenzt wird. In der Mathematik ist es allerdings üblich, einen Kreiskegel K als unendliches Gebilde aufzufassen, nämlich als die Menge aller Punkte derjenigen Geraden, die einen Punkt S des Raumes mit den Punkten eines Kreises k verbinden, siehe Abb. 2.17. Diese Geraden heißen *Mantellinien*, der Punkt S *Spitze* und der Kreis k *Grundkreis* des Kreiskegels K . Die Gerade durch die Spitze S und den Mittelpunkt M des Grundkreises wird als *Achse* des Kreiskegels K bezeichnet. Steht die Achse eines Kreiskegels K senkrecht auf der Grundkreisebene ε , so ist K ein gerader Kreiskegel. Ist α der Winkel zwischen der Kegelachse und den Mantellinien, so heißt 2α *Öffnungswinkel* von K . Kreiskegel bestehen nach dieser Begriffsauffassung aus zwei Kegelästen, es handelt sich also um „Doppelkegel“.

Der Grundkreis eines Kreiskegels ist nicht eindeutig bestimmt. So könnte jeder der in Abb. 2.18 dargestellten Kreise als Grundkreis ein und desselben Kegels verwendet werden. Hingegen sind jedem geraden Kreiskegel eindeutig eine Achse, eine Spitze und ein Öffnungswinkel zugeordnet, und durch diese wird umgekehrt ein gerader Kreiskegel eindeutig bestimmt.

Es seien a eine Gerade und S ein Punkt von a . Der *gerade Kreiskegel mit der Spitze S , der Achse a und dem Öffnungswinkel 2α* ist die Menge aller Punkte derjenigen Geraden, die durch S verlaufen und mit a den Winkel α einschließen.

Um eine Gleichung für gerade Kreiskegel herzuleiten, betrachten wir einen geraden Kreiskegel, dessen Spitze im Koordinatenursprung liegt und dessen Achse die z -Achse ist (siehe Abb. 2.18). Dieser Kegel besteht aus den Punkten P von Kreisen, die in zur x - y -Ebene parallelen Ebenen liegen und deren Mittelpunkte jeweils der z -Achse angehören. Jeder dieser Kreise ist durch die z -Koordinate seines Mittelpunktes (welche mit der z -Koordinate jedes Punktes des Kreises

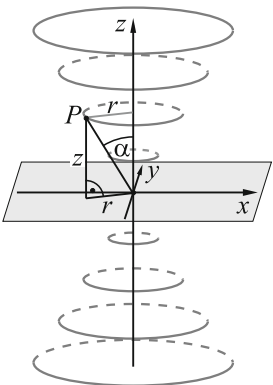
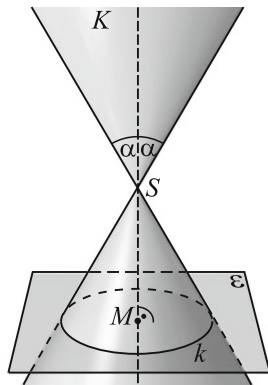


Abb. 2.17: Kegel

Abb. 2.18: Herleitung einer Kegelgleichung

übereinstimmt) eindeutig bestimmt. Der Radius eines solchen Kreises hängt mit seiner z -Koordinate zusammen: Da der Winkel bei P in dem in Abb. 2.18 dargestellten rechtwinkligen Dreieck gleich dem halben Öffnungswinkel α des Kreiskegels ist, gilt

$$\tan \alpha = \frac{r}{|z|},$$

und somit

$$r = |z| \tan \alpha.$$

Da jeder Punkt des Kegels auf einem der beschriebenen Kreise liegt, gilt also für die Koordinaten eines beliebigen Punktes $P(x; y; z)$ des Kreiskegels:

$$x^2 + y^2 = r^2 = z^2 \cdot (\tan \alpha)^2.$$

Da für α beliebige spitze Winkel in Frage kommen, kann $\tan \alpha$ beliebige positive Werte annehmen.

Jeder gerade Kreiskegel, dessen Spitze im Koordinatenursprung liegt und dessen Achse die z -Achse ist, wird durch eine Gleichung der Form

$$x^2 + y^2 = A^2 z^2 \quad (2.13)$$

beschrieben. Dabei ist $\arctan A$ der halbe Öffnungswinkel des Kreiskegels.

Schnittfiguren einer Ebene und eines Kreiskegels

Für einige einfache Fälle können wir nun Gleichungen der Schnittkurven einer Ebene und eines Kreiskegels ermitteln.¹ Wir betrachten dazu den Kegel mit der Gleichung $K: x^2 + y^2 = z^2$ (also mit $A=1$, d. h. $\alpha=45^\circ$) und die durch folgende Gleichungen beschriebenen Ebenen (siehe Abb. 2.19²):

$$\varepsilon_1: z = 1,5$$

$$\varepsilon_2: y = -1$$

$$\varepsilon_3: z = y + 1$$

Durch Einsetzen dieser Gleichungen in die Gleichung von K ergibt sich:

$$x^2 + y^2 = 2,25$$

$$z^2 - x^2 = 1$$

$$x^2 = 2y + 1 \text{ bzw. } y = \frac{1}{2}x^2 - \frac{1}{2}$$

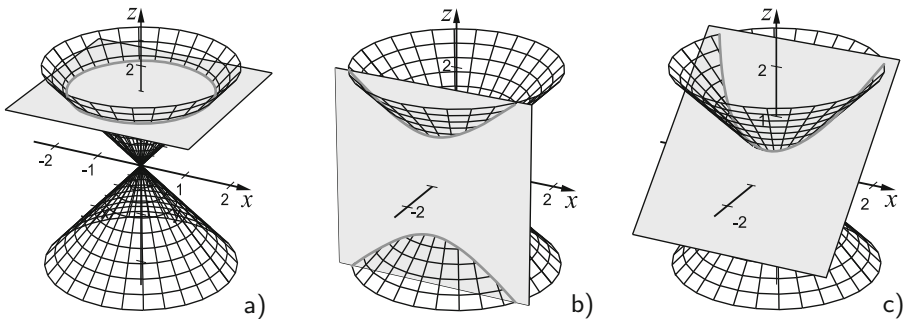


Abb. 2.19: Schnittkurven eines Kreiskegels mit ausgewählten Ebenen

¹Eine allgemeine Herleitung der Gleichungen von Schnittfiguren beliebiger Kreiskegel mit beliebigen Ebenen ist mit Mitteln der Vektorrechnung möglich. Ein entsprechender ergänzender Abschnitt steht auf der Internetseite zu diesem Buch zur Verfügung.

²Interaktive Versionen dieser Abbildungen, die aus allen Richtungen betrachtet werden können, sowie ein Video, welches verschiedene Kegelschnitte zeigt, enthält die Internetseite zu diesem Buch.

Diese Gleichungen beschreiben die Projektionen der Schnittkurven von K mit ε_1 , ε_2 und ε_3 in die x - y - bzw. in die x - z -Ebene. Bei den Schnittkurven handelt es sich um einen Kreis, eine Hyperbel und eine Parabel. Weiterhin können als Schnittkurven eines Kreiskegels und einer Ebene, die nicht durch die Kegelspitze verläuft, auch Ellipsen (die keine Kreise sind) auftreten. Alle diese Kurven werden durch quadratische Gleichungen (in zwei Variablen) beschrieben und heißen daher *Kurven zweiter Ordnung*. Eine genauere Untersuchung derartiger Kurven erfolgt im folgenden Abschnitt.

2.3.3 Aufgaben zu Abschnitt 2.3

1. Geben Sie jeweils den Radius und eine Gleichung der Kugel mit dem Mittelpunkt M an, die durch den Punkt P_0 verläuft.

- a) $M(0; 0; 0)$; $P_0(-1; -1; -1)$ b) $M(1; 2; 3)$; $P_0(0; 0; 0)$
 c) $M(4; -5; 6)$; $P_0(8; 3; -2)$ d) $M(-\frac{1}{3}; \frac{1}{2}; -\frac{1}{6})$; $P_0(\frac{2}{3}; \frac{1}{6}; -\frac{1}{2})$

2. Überprüfen Sie, ob der Punkt P auf, innerhalb oder außerhalb der Kugel k mit dem Mittelpunkt M und dem Radius r liegt.

- a) $M(0; 0; 0)$, $r=6$, $P(-7; -4; 6)$ b) $M(-3; 4; -5)$, $r=3$, $P(-1; 6; -6)$
 c) $M(11; 32; -17)$, $r=10$, $P(14; 29; -9)$ d) $M(-7; 8; 2)$, $r=15$, $P(3; 3; -8)$

3. Überprüfen Sie, ob durch die folgenden Gleichungen Kugeln beschrieben werden, und bestimmen Sie (falls möglich) Mittelpunkt und Radius.

Hinweis: Versuchen Sie, die Gleichungen mittels quadratischer Ergänzung in die Form der Kugelgleichung (2.12) zu bringen (siehe Beispiel 2.6 auf S. 61).

- a) $k: x^2 + y^2 + z^2 + 12x + 10y + 14z = -11$
 b) $k: x^2 + y^2 + z^2 - 8x + 22y + 4z = -200$
 c) $k: x^2 + y^2 + z^2 + 0,5x + 2y - z = -1,25$

4. Bestimmen Sie die Koordinaten des Mittelpunktes, den Radius sowie eine Gleichung der Kugel k , die durch die Punkte A , B , C und D verläuft.

- a) $A(3; 0; 0)$, $B(0; -3; 0)$, $C(-2; 1; 2)$, $D(2; 2; -1)$
 b) $A(2; 14; 5)$, $B(4; 10; 7)$, $C(4; 15; 4)$, $D(7; 15; 3)$
 c) $A(-5; -7; -11)$, $B(-1; -1; -7)$, $C(-2; -4; -3)$, $D(-8; 0; -7)$

5. Bestimmen Sie die Koordinaten des Mittelpunktes sowie eine Gleichung der Kugel k mit dem Radius $r = 10$, die durch die Punkte A , B und C verläuft.

- a) $A(9; 8; 3)$, $B(1; 8; 11)$, $C(-5; 2; -5)$
 b) $A(11; -10; 0)$, $B(19; -4; -14)$, $C(17; -12; -8)$
 c) $A(5; -18; 16)$, $B(-3; -10; 4)$, $C(13; -4; 10)$

6. Leiten Sie je eine Gleichung der Form $z = f(x, y)$ für den oberen und den unteren Ast eines geraden Kreiskegels mit dem Öffnungswinkel 2α her, dessen Spitze im Koordinatenursprung liegt und dessen Achse mit der z -Achse identisch ist (siehe Abb. 2.18). Drücken Sie dazu für jeden Punkt des Kegels die z -Koordinate in Abhängigkeit von den zugehörigen x - und y -Werten aus.

2.4 Ellipsen, Hyperbeln und Parabeln

Bei der Untersuchung der Schnittfiguren eines Kreiskegels mit einigen Ebenen im vorangegangenen Abschnitt traten einige Kurven zweiter Ordnung (Kurven, die durch quadratische Gleichungen in zwei Variablen beschrieben werden) auf. Derartige Kurven – auch Kegelschnitte genannt – sind in vielfältigen Kontexten in der Mathematik sowie in Anwendungen von Bedeutung. Im Folgenden werden Kurven zweiter Ordnung (Ellipsen, Hyperbeln und Parabeln) zunächst unter geometrisch-konstruktiven Aspekten betrachtet. Davon ausgehend erfolgen dann Herleitungen von Gleichungen für diese Klassen von Kurven, und schließlich werden dann ihre Tangenten untersucht.

2.4.1 Ortsdefinitionen von Ellipsen, Hyperbeln und Parabeln

In früheren Jahrhunderten wurde häufig großer Wert auf elliptisch geformte Blumenbeete und Hecken gelegt. Aus dieser Zeit stammt die „Gärtnerkonstruktion“, eine Ellipsenkonstruktion, die der Gärtner mit zwei Pfählen und einer Schnur ausführt. Die Schnur ist durch die beiden Pfähle und einen Zeichenstock jederzeit fest gespannt; der Gärtner läuft dabei um die Pfähle herum, wobei er eine Ellipse zeichnet, siehe Abb. 2.20. Ausgehend von der Gärtnerkonstruktion lassen sich Ellipsen folgendermaßen definieren:³

Ortsdefinition der Ellipse

Eine Ellipse ist die Menge aller Punkte P der Ebene, für welche die Summe der Abstände zu zwei vorgegebenen Punkten F_1 und F_2 konstant und größer als der Abstand von F_1 und F_2 ist (siehe Abb. 2.21):

$$|F_1P| + |F_2P| = \text{const} , \quad |F_1P| + |F_2P| > |F_1F_2| .$$

Die Punkte F_1 und F_2 werden als *Brennpunkte* der Ellipse bezeichnet.

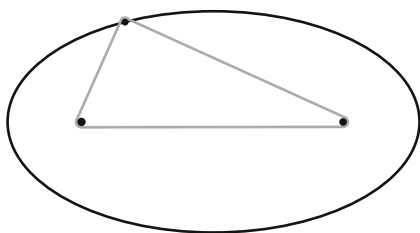


Abb. 2.20: Gärtnerkonstruktion der Ellipse

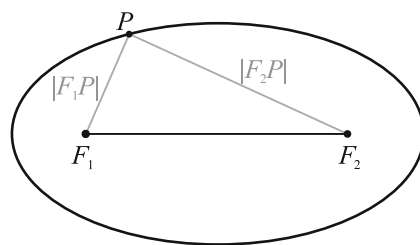


Abb. 2.21: Ortsdefinition der Ellipse

³Die Bezeichnung „Ortsdefinition“ rührt daher, dass früher statt von Punktmengen von geometrischen Orten gesprochen wurde. Die obige Definition wurde also begonnen mit „Eine Ellipse ist der geometrische Ort aller Punkte ...“.

Beispiel 2.10

Punktweise Konstruktion einer Ellipse (siehe Abb. 2.22)

Es soll eine Ellipse mit dem Abstand der Brennpunkte $|F_1 F_2| = 5 \text{ cm}$ gezeichnet werden, für welche die Abstandssumme $|F_1 P| + |F_2 P|$ ihrer Punkte P zu den Brennpunkten 7 cm beträgt. Hierfür werden Paare von Kreisen mit den Mittelpunkten F_1 und F_2 konstruiert, für welche die Summe der Radien 7 cm ist. Für jedes derartige Paar von Kreisen werden die Schnittpunkte markiert. Diese Schnittpunkte sind Punkte der zu konstruierenden Ellipse. Nach der Konstruktion genügend vieler Punkte der Ellipse werden diese verbunden.

Gut eignet sich für die Durchführung einer derartigen Konstruktion eine dynamische Geometriesoftware, die Spurkurven von Schnittpunkten darstellen kann (z. B. Geogebra). Die Internetseite zu diesem Buch enthält entsprechende dynamische Konstruktionen von Ellipsen sowie von Hyperbeln und Parabeln. ■

Bestimmungstücke und -größen von Ellipsen

Die Hälfte des Abstandes $|F_1 F_2|$ zwischen den Brennpunkten einer Ellipse heißt *lineare Exzentrizität* e . Der Mittelpunkt der Strecke $\overline{F_1 F_2}$ wird *Mittelpunkt* M der Ellipse genannt. Die Gerade, die durch die Brennpunkte einer Ellipse verläuft, heißt *Hauptachse* (mit der Länge $2a$) die dazu senkrechte Achse durch den Mittelpunkt *Nebenachse* (mit der Länge $2b$).

Die beiden Punkte S_1 und S_2 einer Ellipse, die auf der Hauptachse liegen, heißen *Hauptscheitel*. Die Schnittpunkte S_3 und S_4 einer Ellipse mit ihrer Nebenachse heißen *Nebenscheitel* (siehe Abb. 2.23).

Satz 2.4

- a) Für jeden Punkt P einer Ellipse mit den beiden Brennpunkten F_1 und F_2 sowie der Länge $2a$ der Hauptachse gilt:

$$|F_1 P| + |F_2 P| = 2a.$$

- b) Für jede Ellipse mit der linearen Exzentrizität e , der Hauptachsenlänge $2a$ sowie der Länge $2b$ der Nebenachse gilt:

$$a^2 = b^2 + e^2.$$

Zum Beweis des Satzes 2.4 siehe Aufgabe 2 auf S. 83.

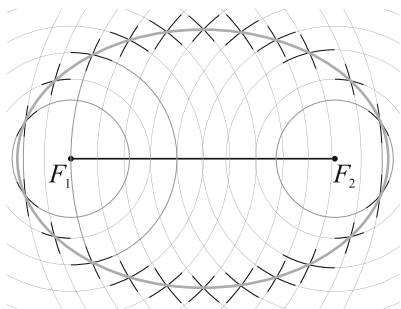


Abb. 2.22: Punktweise Konstruktion einer Ellipse

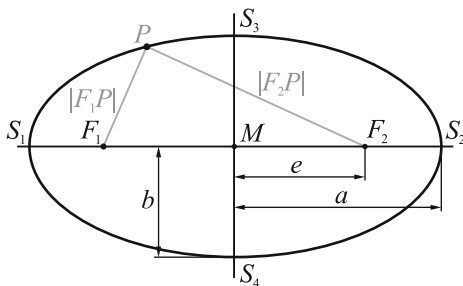


Abb. 2.23: Bestimmungstücke und -größen von Ellipsen

Ortsdefinition der Hyperbel

Eine Hyperbel ist die Menge aller Punkte P der Ebene, für welche der Absolutbetrag der Differenz der Abstände zu zwei vorgegebenen Punkten F_1 und F_2 konstant und positiv ist (siehe Abb. 2.24):

$$||F_1P| - |F_2P|| = \text{const}, \quad ||F_1P| - |F_2P|| > 0.$$

Die Punkte F_1 und F_2 werden als *Brennpunkte* der Hyperbel bezeichnet.

Konstruktionen von Hyperbeln lassen sich anhand dieser Definition analog zu der in dem Beispiel 2.10 beschriebenen Ellipsenkonstruktion führen. Eine entsprechende dynamische Konstruktion enthält die Internetseite zu diesem Buch.

Wie bei der Ellipse werden die Hälfte des Abstandes $|F_1F_2|$ zwischen den Brennpunkten einer Hyperbel als *lineare Exzentrizität* e und der Mittelpunkt der Strecke $\overline{F_1F_2}$ als *Mittelpunkt* M der Hyperbel bezeichnet. Die Schnittpunkte S_1 und S_2 einer Hyperbel mit ihrer Achse F_1F_2 heißen *Scheitel* der Hyperbel.

Satz 2.5

Für jeden Punkt P einer Hyperbel mit den beiden Brennpunkten F_1 und F_2 sowie den Scheitelpunkten S_1 und S_2 gilt:

$$||F_1P| - |F_2P|| = |S_1S_2| = 2a.$$

Beispiel 2.11

Der Graph der Funktion f mit $f(x) = \frac{1}{x}$ ist eine Hyperbel mit den Brennpunkten $F_1(\sqrt{2}; \sqrt{2})$ und $F_2(-\sqrt{2}; -\sqrt{2})$, siehe Abb. 2.25.

Es ist zu zeigen, dass $||F_1P| - |F_2P||$ für alle Punkte $P(x; \frac{1}{x})$ mit $x \in \mathbb{R}, x \neq 0$ gleich, also unabhängig von x ist. Wir berechnen dazu $(|F_1P| - |F_2P|)^2$:

$$\begin{aligned} (|F_1P| - |F_2P|)^2 &= \left(\sqrt{(\sqrt{2}-x)^2 + (\sqrt{2}-\frac{1}{x})^2} - \sqrt{(-\sqrt{2}-x)^2 + (-\sqrt{2}-\frac{1}{x})^2} \right)^2 \\ &= (\sqrt{2}-x)^2 + (\sqrt{2}-\frac{1}{x})^2 + (\sqrt{2}+x)^2 + (\sqrt{2}+\frac{1}{x})^2 \\ &\quad - 2\sqrt{\left((\sqrt{2}-x)^2 + (\sqrt{2}-\frac{1}{x})^2\right) \cdot \left((\sqrt{2}+x)^2 + (\sqrt{2}+\frac{1}{x})^2\right)} \end{aligned}$$

Der Ausdruck unter der letzten Wurzel lässt sich zu $x^4 + \frac{1}{x^4} + 2$ vereinfachen. Damit (und nach Vereinfachung der darüber stehenden Zeile) ergibt sich

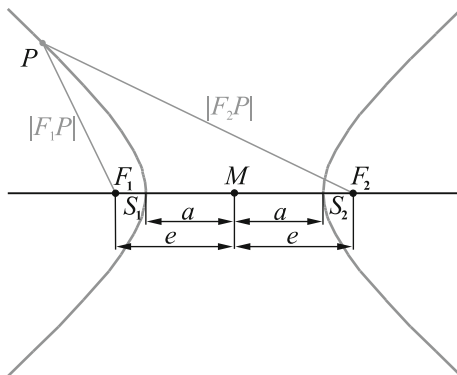


Abb. 2.24: Hyperbel

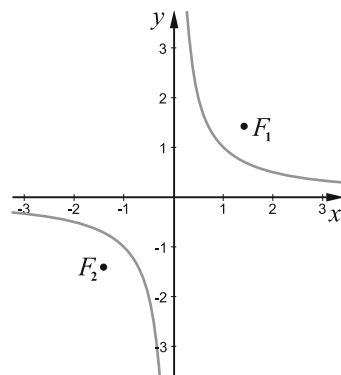


Abb. 2.25: Graph der Funktion f mit $f(x) = \frac{1}{x}$

$$\begin{aligned}
 (|F_1P| - |F_2P|)^2 &= 8 + 2x^2 + 2\frac{1}{x^2} - 2\sqrt{x^4 + \frac{1}{x^4} + 2} \\
 &= 8 + 2x^2 + 2\frac{1}{x^2} - 2\left(x^2 + \frac{1}{x^2}\right) = 8.
 \end{aligned}$$

Es gilt somit $||F_1P| - |F_2P|| = \sqrt{8}$ für alle Punkte P des Graphen der Funktion f mit $f(x) = \frac{1}{x}$. Alle Punkte des Funktionsgraphen gehören somit zu der Hyperbel mit den Brennpunkten F_1 und F_2 sowie dem Abstand der Scheitelpunkte $2a = \sqrt{8} = 2\sqrt{2}$. Ohne Beweis sei angemerkt, dass umgekehrt alle Punkte dieser Hyperbel auch Punkte des Funktionsgraphen sind. ■

Ortsdefinition der Parabel

Eine Parabel ist die Menge aller Punkte P der Ebene, für die der Abstand von einem festen Punkt F gleich dem Abstand von einer festen Geraden l ist (siehe Abb. 2.26):

$$|FP| = d(P, l).$$

Der Punkt F wird als *Brennpunkt*, die Gerade l als *Leitlinie* der Parabel bezeichnet.

Der Abstand p des Brennpunktes von der Leitlinie heißt *Halbparameter*, die Größe $2p$ *Parameter* der Parabel. Die zur Leitlinie l senkrechte Gerade durch den Brennpunkt wird *Achse* und der Schnittpunkt S einer Parabel mit ihrer Achse *Scheitelpunkt* der Parabel genannt.

Beispiel 2.12

Der Graph der Funktion $f: f(x) = x^2$ ist eine Parabel mit dem Brennpunkt $F(0; \frac{1}{4})$ und der Leitlinie $l: y = -\frac{1}{4}$, siehe Abb. 2.27.

Für die Abstände eines beliebigen Punktes $P(x; x^2)$ vom Brennpunkt und von der Leitlinie gilt

$$\begin{aligned}
 |FP|^2 &= x^2 + \left(x^2 - \frac{1}{4}\right)^2 = x^2 + x^4 - \frac{1}{2}x^2 + \frac{1}{16} = x^4 + \frac{1}{2}x^2 + \frac{1}{16} \quad \text{und} \\
 (d(P, l))^2 &= \left(x^2 + \frac{1}{4}\right)^2 = x^4 + \frac{1}{2}x^2 + \frac{1}{16}.
 \end{aligned}$$

Beide Abstände sind also für beliebige Punkte des Funktionsgraphen gleich. Umgekehrt liegen alle Punkte, für die $|FP| = d(P, l)$ gilt, auf dem Graphen der Funktion f mit $f(x) = x^2$. Dieser Funktionsgraph und die Parabel mit dem Brennpunkt $F(0; \frac{1}{4})$ und der Leitlinie $l: y = -\frac{1}{4}$ sind somit identisch. ■

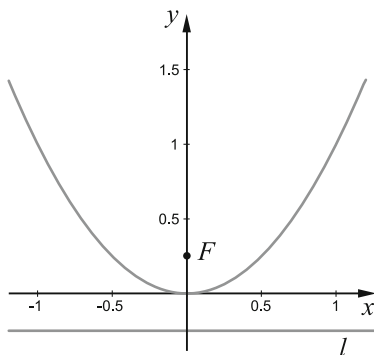
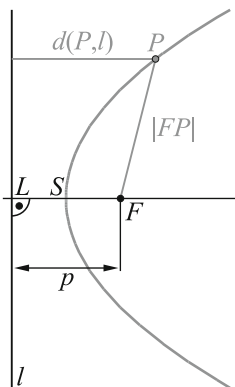


Abb. 2.26: Parabel (links)

Abb. 2.27: Graph der Funktion f mit $f(x) = x^2$ (rechts)

2.4.2 Gleichungen von Ellipsen, Hyperbeln und Parabeln

Die Gleichung der Ellipse in Hauptachsenlage

Um möglichst einfach eine Gleichung für eine Ellipse herleiten zu können, sollte ihre Hauptachse mit der x -Achse und die Nebenachse mit der y -Achse übereinstimmen (eine solche Lage heißt *Hauptachsenlage*, siehe Abb. 2.28). Durch eine geeignete Wahl des Koordinatensystems lässt sich dies für jede Ellipse erreichen.

Ist P ein von den beiden Hauptscheiteln verschiedener Punkt der Ellipse und P' der Fußpunkt des Lotes von P auf die x -Achse, so sind $\triangle F_1PP'$ und $\triangle F_2PP'$ rechtwinklige Dreiecke. Dabei ist $|PP'| = y$ sowie

$$|F_1P'| = |e + x| \quad \text{und} \quad |F_2P'| = |e - x|.$$

Nach dem Satz des Pythagoras gilt:

$$|F_1P| = \sqrt{y^2 + (e + x)^2} \quad \text{und} \quad |F_2P| = \sqrt{y^2 + (e - x)^2}.$$

Daraus ergibt sich:

$$|F_1P| + |F_2P| = \sqrt{y^2 + (e + x)^2} + \sqrt{y^2 + (e - x)^2}.$$

Wegen $|F_1P| + |F_2P| = 2a$ (siehe Satz 2.4) folgt daraus:

$$\sqrt{y^2 + (e + x)^2} + \sqrt{y^2 + (e - x)^2} = 2a.$$

Diese Gleichung wird nun vereinfacht. Zunächst wird auf beiden Seiten $\sqrt{y^2 + (e - x)^2}$ subtrahiert, und es werden beide Seiten der entstehenden Gleichung quadriert:

$$\begin{aligned} \sqrt{y^2 + (e + x)^2} &= 2a - \sqrt{y^2 + (e - x)^2} \\ y^2 + (e + x)^2 &= 4a^2 + y^2 + (e - x)^2 - 4a\sqrt{y^2 + (e - x)^2}. \end{aligned}$$

Durch Auflösen der Klammern und Vereinfachen ergibt sich daraus

$$a\sqrt{y^2 + (e - x)^2} = a^2 - ex,$$

erneutes Quadrieren beider Seiten führt zu

$$a^2 [y^2 + (e - x)^2] = a^4 + e^2x^2 - 2a^2ex$$

bzw.

$$(a^2 - e^2)x^2 + a^2y^2 = a^2(a^2 - e^2).$$

Nach Ersetzen von $a^2 - e^2$ durch b^2 (siehe Satz 2.4 auf S. 73) folgt daraus

$$b^2x^2 + a^2y^2 = a^2b^2.$$

Division dieser Gleichung durch a^2b^2 auf beiden Seiten ergibt schließlich:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1.$$

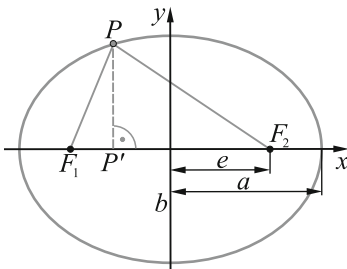


Abb. 2.28: Herleitung der Ellipsengleichung in Hauptachsenlage

Bislang wurde gezeigt, dass alle Punkte der gegebenen Ellipse die Gleichung $(*) \frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$ erfüllen. Um nachzuweisen, dass diese Gleichung tatsächlich eine Gleichung der Ellipse ist, muss noch gezeigt werden, dass *nur* Punkte der Ellipse die Gleichung erfüllen bzw. dass alle Punkte, deren Koordinaten die Gleichung $(*)$ erfüllen, der Ellipse angehören. Dazu berechnen wir zunächst die Abstände eines beliebigen Punktes $Q(x; y)$ zu den Brennpunkten F_1 und F_2 :

$$(**) \quad |QF_1| = \sqrt{(x+e)^2 + y^2}, \quad |QF_2| = \sqrt{(x-e)^2 + y^2}.$$

Ist Q ein Punkt, der die Gleichung $(*)$ erfüllt, so gilt $y^2 = b^2 \left(1 - \frac{x^2}{a^2}\right)$. Wegen $b^2 = a^2 - e^2$ ergibt sich daraus:

$$y^2 = a^2 - e^2 - x^2 + \frac{e^2}{a^2} x^2.$$

Für $(**)$ kann deshalb geschrieben werden:

$$|QF_1| = \sqrt{2ex + a^2 + \frac{e^2}{a^2} x^2}, \quad |QF_2| = \sqrt{-2ex + a^2 + \frac{e^2}{a^2} x^2}$$

beziehungsweise

$$(***) \quad |QF_1| = \pm \left(a + \frac{e}{a} x\right), \quad |QF_2| = \pm \left(a - \frac{e}{a} x\right).$$

Da in diesen beiden Gleichungen auf den linken Seiten Abstände, also positive Zahlen, stehen, müssen auf den rechten Seiten die Vorzeichen so gewählt werden, dass sich ebenfalls positive Werte ergeben. Da der Punkt $Q(x; y)$ der Gleichung $(*)$ genügen soll, muss $\frac{x^2}{a^2} \leq 1$ und somit $|x| \leq a$ sein. Wegen $a^2 = e^2 + b^2$ muss weiterhin $e < a$ sein. Somit sind in $(***)$ die beiden Ausdrücke in den Klammern positiv, und es ergibt sich daraus

$$|QF_1| + |QF_2| = a + \frac{e}{a} x + a - \frac{e}{a} x = 2a.$$

Der Punkt Q erfüllt daher die Ortsdefinition (siehe S. 72) und ist ein Punkt der betrachteten Ellipse. Damit ist nachgewiesen, dass die Gleichung $(*)$ alle Punkte der Ellipse und nur diese beschreibt.

Mittelpunktsgleichung der Ellipse

Eine Ellipse, deren Mittelpunkt der Koordinatenursprung ist, deren Hauptachse auf der x -Achse und deren Nebenachse auf der y -Achse liegt, hat die Gleichung

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1. \quad (2.14)$$

Dabei ist a die halbe Länge der Haupt-, b die halbe Länge der Nebenachse.

Kreise sind spezielle Ellipsen. Falls die Brennpunkte einer Ellipse identisch sind, also $e = |F_1 F_2| = 0$ gilt, so folgt daraus nach Satz 2.4 $a = b$. Die Gleichung (2.14) nimmt damit die Gestalt $\frac{x^2}{a^2} + \frac{y^2}{a^2} = 1$ bzw. $x^2 + y^2 = a^2$ an und entspricht der Kreisgleichung in Mittelpunktslage (2.7), siehe S. 58, mit dem Radius $r = a$.

Beispiel 2.13

Aufstellen einer Gleichung der Ellipse mit den Brennpunkten $F_1(-4; 0)$ und $F_2(4; 0)$ sowie dem Ellipsenpunkt $P(4; 6)$.

Die Länge der Hauptachse der Ellipse erhält man durch

$$2a = |F_1 P| + |F_2 P| = \sqrt{8^2 + 6^2} + 6 = 16.$$

Für die lineare Exzentrizität gilt $2e = |F_1 F_2| = 8$, also $e = 4$. Die Länge der Nebenachse ist nach Teil b) des Satzes 2.4 auf S. 73:

$$b = \sqrt{a^2 - e^2} = \sqrt{48}.$$

Somit ergibt sich wegen $a = 8$ und $b = \sqrt{48}$ die folgende Gleichung der Ellipse:

$$\frac{x^2}{64} + \frac{y^2}{48} = 1 \quad \text{bzw.} \quad 3x^2 + 4y^2 = 192. \quad \blacksquare$$

Graphische Darstellung von Ellipsen mithilfe eines Computeralgebrasystems

Ellipsen lassen sich in dem CAS Maxima nach den Vorgehensweisen darstellen, die auf S. 58 und S. 59 für Kreise beschrieben wurden. Bei Verwendung der **implicit**-Anweisung ist in dem auf S. 58 angegebenen Maxima-Code die Kreisgleichung durch die Gleichung einer Ellipse zu ersetzen. Feinere Darstellungen bei gleichzeitig kürzerer Berechnungszeit sind mittels der **explicit**-Anweisung durch die Darstellung von zwei Halbellipsen möglich, die als Graphen von Funktionen f_1 und f_2 gegeben sind. Dazu wird die Gleichung (2.14) nach y umgestellt: $y = \pm b\sqrt{1 - \frac{x^2}{a^2}}$. Funktionsgleichungen der beiden Halbellipsen sind somit:

$$f_1(x) = b\sqrt{1 - \frac{x^2}{a^2}} \quad \text{und} \quad f_2(x) = -b\sqrt{1 - \frac{x^2}{a^2}} \quad (2.15)$$

Eine Beispieldatei für die Darstellung von Ellipsen, Hyperbeln und Parabeln nach beiden Vorgehensweisen enthält die Internetseite zu diesem Buch.

Die Gleichung der Hyperbel in Hauptachsenlage

Auf analoge Weise wie für Ellipsen lässt sich eine Gleichung für Hyperbeln in Hauptachsenlage herleiten (siehe Aufgabe 12 auf S. 83).

Mittelpunktsgleichung der Hyperbel

Eine Hyperbel, deren Mittelpunkt der Koordinatenursprung ist und deren Brennpunkte auf der x -Achse liegen, hat die Gleichung

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1. \quad (2.16)$$

Dabei ist a die Hälfte des Abstandes zwischen den beiden Scheiteln, und es ist $b = \sqrt{e^2 - a^2}$ mit der linearen Exzentrizität e , siehe Abb. 2.30.

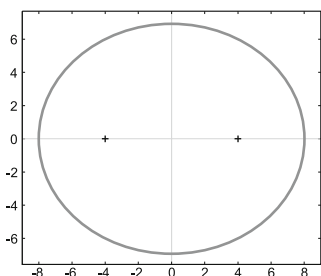


Abb. 2.29: Ellipse in dem CAS Maxima

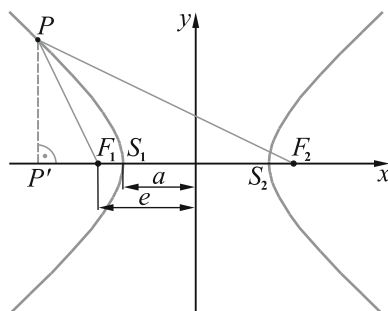


Abb. 2.30: Herleitung der Hyperbelgleichung

Beispiel 2.14

In dem Beispiel 2.11 auf S. 74 wurde festgestellt, dass der Graph der Funktion f mit $f(x) = \frac{1}{x}$ eine Hyperbel ist. Durch eine *Koordinatentransformation* wird diese Hyperbel nun in ein Koordinatensystem überführt, bezüglich dessen sie sich in Hauptachsenlage befindet, und eine Gleichung der Hyperbel ermittelt.

Ein Koordinatensystem (mit den Achsen x' und y'), bezüglich dessen sich die Hyperbel in Hauptachsenlage befindet, geht aus dem ursprünglichen Koordinatensystem durch eine Drehung um 45° hervor, siehe Abb. 2.31. Die „alten“ können durch die „neuen“ Koordinaten folgendermaßen ausgedrückt werden:

$$x = \frac{1}{2}\sqrt{2} x' - \frac{1}{2}\sqrt{2} y', \quad y = \frac{1}{2}\sqrt{2} x' + \frac{1}{2}\sqrt{2} y'.$$

Diese Transformationsgleichungen können hier noch nicht fundiert begründet werden; dies wird erst nach allgemeineren Betrachtungen zu Koordinatensystemen in dem Abschnitt 5.5.2 möglich sein (siehe Beispiel 5.35 auf S. 209).

Durch Ersetzen von x und y durch x' und y' in der Gleichung $y = \frac{1}{x}$ mittels der obigen Transformationsgleichungen ergibt sich die Gleichung

$$\frac{1}{2}\sqrt{2} x' + \frac{1}{2}\sqrt{2} y' = \frac{1}{\frac{1}{2}\sqrt{2} x' - \frac{1}{2}\sqrt{2} y'},$$

die sich leicht in die Form

$$\frac{x'^2}{2} - \frac{y'^2}{2} = 1$$

bringen lässt. Für die lineare Exzentrizität erhält man $e = \sqrt{a^2 + b^2} = 2$. ■

Die Scheitelpunktsgleichung der Parabel

Um eine Parabelgleichung herzuleiten, ist – wie bereits bei der Ellipse und der Hyperbel – ein geeignetes Koordinatensystem zu wählen. Dazu legt man die x -Achse auf die Achse der Parabel und den Koordinatenursprung auf ihren Scheitelpunkt. Die y -Achse verläuft dann parallel zur Leitlinie, siehe Abb. 2.32.

Für einen beliebigen Punkt $P(x;y)$ ist $|PF|^2 = (x - \frac{p}{2})^2 + y^2$ und $d(P, l) = x + \frac{p}{2}$ (siehe Abb. 2.32). Falls P ein Punkt der Parabel ist, so folgt aus der Ortsdefinition (vgl. S. 75) $|PF|^2 = (d(P, l))^2$; durch Einsetzen ergibt sich daraus

$$\left(x - \frac{p}{2}\right)^2 + y^2 = \left(x + \frac{p}{2}\right)^2.$$

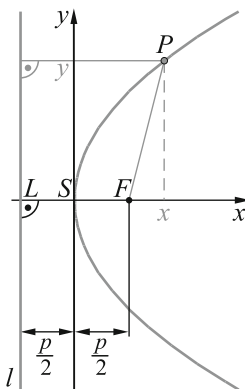
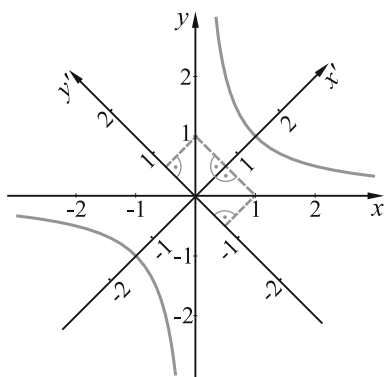


Abb. 2.31:
Überführung einer Hyperbel in Hauptachsenlage

Abb. 2.32:
Herleitung einer Parabelgleichung

Diese Gleichung lässt sich leicht in die Form $y^2 = 2px$ bringen.

Scheitelpunktsgleichung der Parabel

Eine Parabel, deren Scheitelpunkt der Koordinatenursprung ist, und deren Leitlinie parallel zur y -Achse verläuft, hat die Gleichung

$$y^2 = 2px. \quad (2.17)$$

Dabei ist p der Halbparameter der Parabel (Abstand des Brennpunktes von der Leitlinie, siehe Abb. 2.32).

Gleichungen von Ellipsen, Hyperbeln und Parabeln in achsenparalleler Lage

Die bisher behandelten Mittelpunktsgleichungen von Ellipsen und Hyperbeln sowie die Scheitelfgleichung der Parabel setzen ein Koordinatensystem voraus, dessen Ursprung mit dem Mittelpunkt der betrachteten Ellipse oder Hyperbel (bzw. dem Scheitelpunkt der Parabel) zusammenfällt. Im Folgenden werden nun etwas allgemeinere Gleichungen hergeleitet, für deren Anwendung nur noch vorausgesetzt werden muss, dass die Koordinatenachsen zu den Achsen der beschriebenen Ellipsen, Hyperbeln bzw. Parabeln parallel sind.

Um eine Gleichung einer Ellipse mit dem Mittelpunkt $M(x_M; y_M)$ bezüglich eines Koordinatensystems KS herzuleiten, dessen x -Achse parallel zur Haupt- und dessen y -Achse parallel zur Nebenachse der Ellipse ist, betrachten wir ein Hilfskoordinatensystem KS', dessen Achsen parallel zu denen von KS sind und dessen Ursprung der Mittelpunkt der Ellipse ist, siehe Abb. 2.33. In diesem Hilfskoordinatensystem hat die betrachtete Ellipse die Mittelpunktsgleichung

$$\frac{x'^2}{a^2} + \frac{y'^2}{b^2} = 1.$$

Zwischen den Koordinaten eines beliebigen Punktes P bezüglich des Koordinatensystems KS und denen bezüglich KS' bestehen die Beziehungen

$$x = x' + x_M, \quad y = y' + y_M \quad \text{bzw.} \quad x' = x - x_M, \quad y' = y - y_M.$$

Durch Einsetzen in die darüber stehende Ellipsengleichung ergibt sich die *Gleichung der Ellipse in achsenparalleler Lage*:

$$\frac{(x-x_M)^2}{a^2} + \frac{(y-y_M)^2}{b^2} = 1.$$

Analog lassen sich Gleichungen für Hyperbeln und Parabeln in achsenparalleler Lage herleiten:

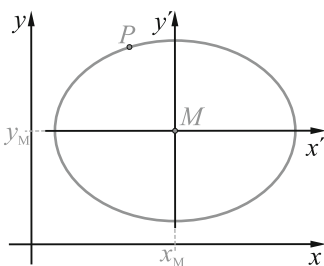


Abb. 2.33: Ellipse in achsenparalleler Lage

Hyperbel: $\frac{(x-x_M)^2}{a^2} - \frac{(y-y_M)^2}{b^2} = 1$, Parabel: $(y-y_M)^2 = 2p(x-x_M)$.

Eine allgemeinere Gleichung, die Ellipsen, Parabeln und Hyperbeln beschreibt, kann mit Mitteln der Vektorrechnung hergeleitet werden, siehe dazu einen ergänzenden Abschnitt auf der Internetseite zu diesem Buch.

2.4.3 Tangenten an Ellipsen, Hyperbeln und Parabeln

Für die Herleitung einer Tangentengleichung für Ellipsen betrachten wir Halbellipsen als Funktionsgraphen, siehe (2.15) auf S. 78. Die Ableitung der Funktionsgleichung der „oberen“ Halbellipse

$$f_1(x) = b\sqrt{1 - \frac{x^2}{a^2}}$$

ist nach der Kettenregel

$$f_1'(x_0) = b \cdot \frac{1}{2\sqrt{1 - \frac{x^2}{a^2}}} \cdot \left(-\frac{2}{a^2}x\right) = \frac{-b^2x}{a^2b\sqrt{1 - \frac{x^2}{a^2}}} = -\frac{b^2}{a^2} \frac{x}{y}.$$

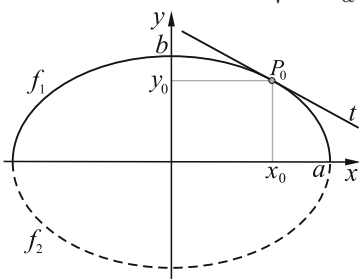


Abb. 2.34: Tangente an eine Ellipse

Somit ist der Anstieg der Tangente in einem beliebigen Punkt P_0 der oberen Halbellipse (mit Ausnahme der beiden Hauptscheitel):

$$m = -\frac{b^2}{a^2} \frac{x_0}{y_0}.$$

Damit lässt sich für die Tangente in P_0 die folgende Gleichung aufstellen:

$$t: y = t(x) = mx + n = -\frac{b^2}{a^2} \frac{x_0}{y_0} x + n.$$

Durch Einsetzen von x_0 für x und von y_0 für y lässt sich n bestimmen:

$$n = y_0 + \frac{b^2}{a^2} \frac{x_0^2}{y_0}.$$

Einsetzen dieses Ausdrucks für n in die Gleichung von t führt zu:

$$t: y = -\frac{b^2}{a^2} \frac{x_0}{y_0} x + y_0 + \frac{b^2}{a^2} \frac{x_0^2}{y_0}.$$

Durch Multiplikation dieser Gleichung mit y_0 , Division durch b^2 und Anwendung der Ellipsengleichung ergibt sich:

$$t: \frac{yy_0}{b^2} + \frac{xx_0}{a^2} = \frac{y_0^2}{b^2} + \frac{x_0^2}{a^2} = 1.$$

Es lässt sich zeigen, dass dieselbe Tangentengleichung für Punkte der unteren Halbellipse gilt, siehe Aufgabe 22 auf S. 84. Auf analogem Wege lässt sich eine Tangentengleichung für Hyperbeln in Hauptachsenlage herleiten (Aufgabe 23).

Tangentengleichungen für Ellipsen und Hyperbeln in Hauptachsenlage

Die Tangente an eine Ellipse mit der Gleichung $\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$ in einem Punkt $P_0(x_0; y_0)$ hat die Gleichung $\frac{xx_0}{a^2} + \frac{yy_0}{b^2} = 1$.

Die Tangente an eine Hyperbel mit der Gleichung $\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$ in einem Punkt $P_0(x_0; y_0)$ hat die Gleichung $\frac{xx_0}{a^2} - \frac{yy_0}{b^2} = 1$.

Um eine *Gleichung für Tangenten an Parabeln* herzuleiten, kann die Parabelgleichung (2.17) $y^2 = 2px$ nach x umgestellt und dann nach y abgeleitet werden:

$$x = f(y) = \frac{1}{2p} y^2, \quad f'(y) = \frac{1}{p} y.$$

Die Ableitung in einem Punkt $P_0(x_0; y_0)$ gibt in diesem Falle den Anstieg der Tangente in diesem Punkt bezüglich vertauschter Koordinatenachsen an, also

$$f'(y_0) = \frac{x - x_0}{y - y_0}$$

für beliebige Punkte $P(x; y)$ der Tangente. Der Anstieg der Tangente ist reziprok zu diesem Wert (und somit gleich der Ableitung der Umkehrfunktion von f):

$$m = \frac{1}{f'(y_0)} = \frac{p}{y_0}.$$

Durch Einsetzen in die Normalform der Geradengleichung ergibt sich

$$t: y = mx + n = \frac{p}{y_0} x + n,$$

und durch Einsetzen von x_0 für x und von y_0 für y lässt sich n bestimmen:

$$n = y_0 - \frac{p}{y_0} x_0.$$

Die Tangentengleichung erhält damit die Form

$$t: y = \frac{p}{y_0} x + y_0 - \frac{p}{y_0} x_0 \quad \text{bzw.} \quad t: y y_0 = p x + y_0^2 - p x_0,$$

Durch Einsetzen von $2px_0$ für y_0^2 ergibt sich daraus $t: y y_0 = p(x + x_0)$.

Tangentengleichungen für Parabeln in Scheitelpunktslage

Die Tangente an eine Parabel mit der Gleichung $y^2 = 2px$ in einem Punkt $P_0(x_0; y_0)$ hat die Gleichung $y y_0 = p(x + x_0)$.

Ellipsen, Parabeln und Hyperbeln besitzen zahlreiche weitere interessante Fakten und Anwendungen, die im Rahmen dieses Buches nicht berücksichtigt werden konnten. Empfohlen sei hierzu das Buch von Schupp (2000).

2.4.4 Aufgaben zu Abschnitt 2.4

1. Begründen Sie, dass jeder Kreis eine Ellipse ist, also der Definition auf S. 72 entspricht. Überlegen Sie dazu, welche lineare Exzentrizität ein Kreis besitzt und welcher Zusammenhang zwischen dem Mittelpunkt des Kreises und den Brennpunkten der Ellipse besteht.

2. Beweisen Sie den Satz 2.4 auf S. 73.
3. Warum muss in der Ortsdefinition der Hyperbel $||F_1P| - |F_2P|| > 0$ gefordert werden? Was für eine Figur erhält man für $||F_1P| - |F_2P|| = 0$?
4. Begründen Sie den Satz 2.5 auf S. 74.
5. Weisen Sie nach, dass die Länge der Sehne $\overline{P_1P_2}$ einer Parabel, die durch den Brennpunkt F dieser Parabel verläuft und senkrecht auf ihrer Achse steht, gleich dem Parameter $2p$ dieser Parabel ist (siehe dazu Abb. 2.26 auf S. 75).
6. Zeigen Sie, dass der Graph der Funktion $f: f(x) = ax^2$ eine Parabel mit dem Brennpunkt $F(0; \frac{1}{4a})$ und der Leitlinie $l: y = -\frac{1}{4a}$ ist.
7. Es seien S der Scheitelpunkt und L der Schnittpunkt von Achse und Leitlinie einer Parabel (Abb. 2.26). Weisen Sie nach, dass dann gilt: $|LS| = |SF| = \frac{p}{2}$.
8. Stellen Sie für die Ellipsen mit den folgenden Stücken Gleichungen auf:
 - a) $a = 5, b = 4$
 - b) $2e = 8, 2a = 10$
 - c) $b = 2, 2e = 6$
 - d) $2a = 20, F_1(-8; 0), F_2(8; 0)$
 - e) $F_1(-2; 0), F_2(2; 0), P(2; 3)$
 - (P liegt auf der Ellipse)
9. Berechnen Sie für die Ellipse mit der Gleichung $16x^2 + 25y^2 = 1600$ die Längen der Achsen, die lineare Exzentrizität und die Brennpunktkoordinaten.
10. Beweisen Sie, dass für Punkte innerhalb bzw. außerhalb einer Ellipse die Ungleichungen $\frac{x^2}{a^2} + \frac{y^2}{b^2} < 1$ bzw. $\frac{x^2}{a^2} + \frac{y^2}{b^2} > 1$ gelten.
Hinweis: Beachten Sie, dass ein Punkt P innerhalb (außerhalb) der Ellipse liegt, falls ein Ellipsenpunkt P_0 existiert, der mit P auf einer Geraden durch den Mittelpunkt M der Ellipse liegt und von M einen größeren (kleineren) Abstand als P hat.
11. a) Konstruieren Sie die Ellipse mit der Gleichung $\frac{x^2}{16} + \frac{y^2}{4} = 1$.
 b) Verschieben Sie die Ellipse um 3 Einheiten in x -Richtung und um 2 Einheiten in y -Richtung. Wie lautet die Gleichung der verschobenen Ellipse?
12. Leiten Sie die Mittelpunktsleichung (2.16) der Hyperbel her, siehe S. 78. Nutzen Sie dazu Abb. 2.30 und orientieren Sie sich an der Herleitung der Ellipsengleichung (S. 76).
13. Geben Sie eine Gleichung der Hyperbel in Hauptachsenlage an, von der folgende Größen bekannt sind.
 - a) $2a = 6, 2b = 8$
 - b) $2a = 24, 2e = 26$
 - c) $a = 8, F_1(-6; 0), F_2(6; 0)$
14. Von einer Hyperbel in Hauptachsenlage sind zwei Punkte P_1 und P_2 gegeben. Stellen Sie die Mittelpunktsleichung der Hyperbel auf.
 - a) $P_1(5; \frac{9}{4}), P_2(2\sqrt{6}; \frac{3}{2}\sqrt{2})$
 - b) $P_1(3; 2), P_2(4; 3)$
15. Gegeben ist die Ellipse mit der Gleichung $\frac{x^2}{8} + \frac{y^2}{5} = 1$. Stellen Sie eine Mittelpunktsleichung der Hyperbel auf, deren Scheitel sich in den Brennpunkten und deren Brennpunkte sich in den Scheiteln der gegebenen Ellipse befinden.

16. Gesucht ist die Seitenlänge eines Quadrates, dessen Eckpunkte auf der Hyperbel mit der Gleichung $\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1$ liegen. Für welche Hyperbeln lässt sich ein solches Quadrat finden?
17. Stellen Sie die Scheitelgleichung der Parabel auf, deren Scheitel im Koordinatenursprung liegt, deren Leitlinie parallel zur y -Achse ist und für die folgende zusätzliche Angabe vorhanden ist:
 a) $F(2;0)$ b) $l: x = -3$ c) $F(-1;0)$ d) $l: x = 10$
18. Wie viele Parabeln mit dem Halbparameter $p = 0,5$ gibt es, deren Scheitel im Ursprung liegen und deren Leitlinien parallel zur y -Achse verlaufen? Geben Sie die Gleichungen dieser Parabeln an.
19. Geben Sie die Koordinaten der Brennpunkte und Gleichungen der Leitlinien für die durch folgende Gleichungen gegebenen Parabeln an.
 a) $y^2 = 20x$ b) $y^2 = -x$ c) $y = x^2$ d) $x^2 = -0,5y$
20. Wie groß ist der Halbparameter der Parabel mit der Gleichung $y^2 = 2px$, wenn diese durch den Punkt $P(5;6)$ verläuft?
21. Berechnen Sie die Koordinaten des Mittelpunktes und die Längen a und b der Halbachsen der Hyperbel H . Bringen Sie dazu die Gleichungen mittels quadratischer Ergänzung in die Form der Hyperbelgleichung in achsenparalleler Lage (siehe S. 81).
 a) $H: x^2 + 4x - 4y^2 - 24y = 36$ b) $H: 2x^2 - 2x - y^2 + y = 0$
22. Leiten Sie die Tangentengleichung $\frac{yy_0}{b^2} + \frac{xx_0}{a^2} = 1$ für Punkte der unteren Halbellipse her (siehe dazu die Herleitung auf S. 81).
23. Leiten Sie eine Tangentengleichung für Hyperbeln in Hauptachsenlage her. Orientieren Sie sich dabei an der Herleitung der Tangentengleichung für Ellipsen auf S. 81.
24. Stellen Sie Gleichungen für die beiden Tangenten an die Ellipse E an der Stelle x_0 auf. Berechnen Sie dazu zunächst die y -Koordinaten der beiden Punkte $P_{01}(x_0; y_{01})$ und $P_{02}(x_0; y_{02})$ der Ellipse, welche die x -Koordinate x_0 haben. Stellen Sie die Ellipsen und ihre Tangenten graphisch dar.
 a) $E: \frac{x^2}{100} + \frac{y^2}{25} = 1, x_0 = -6$ b) $E: \frac{x^2}{400} + \frac{y^2}{225} = 1, x_0 = 16$
25. Stellen Sie eine Gleichung für die Tangente an die Hyperbel H im Punkt P_0 auf und vereinfachen Sie diese so weit wie möglich.
 a) $H: \frac{x^2}{64} - \frac{y^2}{16} = 1; P_0(10; -3)$ b) $H: \frac{(x-4)^2}{16} - \frac{(y-2)^2}{9} = 1; P_0(9; -\frac{1}{4})$
26. Stellen Sie eine Gleichung der Tangente t an die Parabel $y^2 = 2x$ auf, die
 a) parallel zu der Geraden $g: 2x - y = -5$ verläuft;
 b) senkrecht auf der Geraden $g: x + y = 7$ steht.

3 Vektoren

Übersicht

3.1	Vektoren in physikalischen und geometrischen Kontexten	86
3.2	Vektoren in arithmetischen Kontexten	100
3.3	Zusammenhänge zwischen Pfeilklassen und n -Tupeln – Vektoren . .	106
3.4	Linearkombinationen von Vektoren	113
3.5	Das Skalarprodukt zweier Vektoren	122
3.6	Vektor- und Spatprodukt	133
3.7	Vektorrechnung und -darstellung mithilfe des Computers	137

Vektoren gehören zu den zentralen Strukturelementen der linearen Algebra. Sie ermöglichen es, in eleganter Weise geometrische Aufgaben auf rechnerischem Wege zu lösen und einige Beschränkungen einer reinen Koordinatengeometrie, die sich im vorangegangenen Kapitel gezeigt haben, zu überwinden. Eine hohe Bedeutung haben Vektoren auch für die Physik, z. B. bei der Addition und Zerlegung von Kräften sowie bei der Berechnung der mechanischen Arbeit.

Die Vektorrechnung entstand in einem langen historischen Prozess vor allem aufgrund des Bedürfnisses nach einem „geometrischen Kalkül“ sowie aus den Anforderungen der Physik heraus. Im weiteren Verlauf löste sich der Vektorbegriff von den konkreten Kontexten, aus denen heraus er entstand. Damit wurde einerseits eine Abstraktion vollzogen – es entstand der Begriff des Vektorraumes, der heute zu den wichtigsten algebraischen Strukturbegriffen gehört. Andererseits führte die Mächtigkeit dieser Struktur zur Erschließung neuer Anwendungsfelder von Vektoren, die heute z. B. weder aus der Informatik noch aus den Wirtschaftswissenschaften wegzudenken sind.

Das vorliegende Buch orientiert sich insofern an der historischen Entwicklung des Vektorbegriffs, als in diesem Kapitel zunächst konkrete Vektormodelle im Mittelpunkt der Betrachtungen stehen, wobei bereits Gemeinsamkeiten verschiedener Zugänge zu Tage treten. Mit den so eingeführten Vektoren wird dann Vektorrechnung betrieben, und es werden einige Anwendungen betrachtet. Der allgemeinere Begriff des Vektorraumes ist dann Gegenstand des Kapitels 5, wobei auf die hier zunächst eingeführten „konkreten“ Vektoren zurückgeblickt und eine Abstraktion vollzogen wird.

3.1 Vektoren in physikalischen und geometrischen Kontexten

3.1.1 Kräfte, Geschwindigkeiten und Verschiebungen

In der Physik gibt es *gerichtete Größen* wie Kräfte, Geschwindigkeiten und Beschleunigungen. Diese sind durch eine Maßzahl und die Maßeinheit nicht vollständig bestimmt, vielmehr sind Richtung und Richtungssinn von Bedeutung.

Den vektoriellen Charakter von *Geschwindigkeiten* verdeutlicht folgendes Beispiel: Ein Boot fährt auf einem Fluss mit einer Geschwindigkeit von $12 \frac{\text{km}}{\text{h}}$ in Richtung auf das andere Flussufer. Die Strömungsgeschwindigkeit des Flusses beträgt $6 \frac{\text{km}}{\text{h}}$. Die Maßangaben beschreiben den Sachverhalt offensichtlich nicht vollständig, vielmehr sind die Richtungen der Geschwindigkeiten von erheblicher Bedeutung. Durch maßstabsgetreue Pfeile lassen sich die Beträge und die Richtungen der Geschwindigkeiten $\vec{v}_1$ und $\vec{v}_2$ darstellen, siehe Abb. 3.1. Die resultierende Geschwindigkeit $\vec{v}$ lässt sich daraus graphisch ermitteln. Indem an einen Pfeil, der die Geschwindigkeit $\vec{v}_1$ kennzeichnet, ein Pfeil angetragen wird, der $\vec{v}_2$ beschreibt, und der Anfangspunkt des ersten Pfeils mit der Spitze des zweiten Pfeils verbunden wird, entsteht ein Pfeil, der die resultierende Geschwindigkeit $\vec{v}$ repräsentiert. Deren Betrag lässt sich aus der Länge des Pfeils ermitteln. In dem konkreten Fall ist, da $\vec{v}_1$ und $\vec{v}_2$ senkrecht zueinander sind, auch eine einfache Berechnung des Betrags der resultierenden Geschwindigkeit mithilfe des Satzes des Pythagoras möglich: $|\vec{v}| = \sqrt{\left(12 \frac{\text{km}}{\text{h}}\right)^2 + \left(6 \frac{\text{km}}{\text{h}}\right)^2} \approx 13,4 \frac{\text{km}}{\text{h}}$.

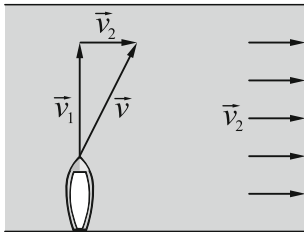


Abb. 3.1: Darstellung von Geschwindigkeiten durch Pfeile

Zu beachten ist, dass die Angabe „fährt mit einer Geschwindigkeit von $12 \frac{\text{km}}{\text{h}}$ “ zwar gebräuchlich, aber nicht ganz korrekt ist, denn sie bezeichnet nicht die Geschwindigkeit des Bootes, sondern nur ihren Betrag. Durch ihre Beträge sind Geschwindigkeiten nicht hinreichend beschrieben. Insbesondere für die Bestimmung resultierender Geschwindigkeiten ist die Kenntnis der Beträge nicht ausreichend: die Addition der Beträge würde in dem obigen Beispiel $18 \frac{\text{km}}{\text{h}}$ ergeben. Dieser Wert ist irrelevant und beschreibt den Vorgang in keiner Weise.

Es sei angemerkt, dass eine Geschwindigkeit nicht durch einen „festen“ Pfeil gekennzeichnet wird, dessen Position eindeutig festgelegt ist. So zeigt Abb. 3.1, dass die Strömungsgeschwindigkeit durch mehrere Pfeile an verschiedenen Stellen dargestellt werden kann. Von Bedeutung ist, dass alle Pfeile, welche dieselbe Geschwindigkeit kennzeichnen, gleich lang, parallel und gleich gerichtet sind.

Ebenso wie die Geschwindigkeit ist auch die *Kraft* eine gerichtete (vektorielle) Größe. Wird beispielsweise ein Schiff durch zwei Boote gezogen, so hängt die resultierende Kraft $\vec{F}$ davon ab, in welche Richtungen die Teilkräfte $\vec{F}_1$ und $\vec{F}_2$ wirken (siehe Abb. 3.2). Die resultierende Kraft kann (wie in dem vorherigen Beispiel die resultierende Geschwindigkeit) durch Antragen eines Pfeils, der $\vec{F}_2$ darstellt an einen Pfeil, der $\vec{F}_1$ beschreibt, graphisch ermittelt werden. Eine andere Möglichkeit besteht darin, einen Pfeil, der die resultierende Kraft ausdrückt, als (gerichtete) Diagonale eines Parallelogramms zu ermitteln, das von zwei die beiden Teilkräfte $\vec{F}_1$ und $\vec{F}_2$ repräsentierenden Pfeilen aufgespannt wird. Beide Möglichkeiten führen, wie aus Abb. 3.2 hervorgeht, zu demselben Ergebnis.

Außer der Kräfte„addition“ (Ermittlung resultierender Kräfte) ist die Zerlegung von Kräften in Komponenten, die entlang vorgegebener Richtungen wirken, von Interesse. So wird die *Gewichtskraft* $\vec{F}_G$ einer Kugel oder eines Wagens, der eine geneigte Ebene hinunterrollt, in zwei Komponenten aufgeteilt, siehe Abb. 3.3. Die *Hangabtriebskraft* $\vec{F}_H$ ist die für den Bewegungsvorgang relevante Kraft, sie bestimmt die Beschleunigung des herunterrollenden Objekts. Die *Normalkraft* $\vec{F}_N$ ist die Kraft, mit der das Objekt die Fahrbahn belastet. Zwischen den drei Kräften besteht der Zusammenhang $\vec{F}_G = \vec{F}_H + \vec{F}_N$.¹ Da meist die Gewichtskraft und der Neigungswinkel einer geneigten Ebene bekannt sind, sind hier häufiger die Beträge der Kräfte $\vec{F}_H$ (und daraus die Beschleunigung) sowie $\vec{F}_N$ zu bestimmen. Dies ist mithilfe der beiden rechtwinkligen Teildreiecke des in Abb. 3.3 dargestellten Kräfteparallelogramms (bei dem es sich in diesem Falle um ein Rechteck handelt) möglich. Es ist nämlich leicht zu überlegen, dass der Winkel zwischen den Kräften $\vec{F}_G$ und $\vec{F}_N$ dem Neigungswinkel α der geneigten Ebene entspricht. Mithilfe trigonometrischer Beziehungen lassen sich damit $|\vec{F}_H|$ und $|\vec{F}_N|$ aus α und $|\vec{F}_G|$ ermitteln, siehe die Aufgabe 2 auf S. 99.

Wie schon bei Geschwindigkeiten bemerkt wurde, sind auch die Positionen der Pfeile, die Kräfte beschreiben, nicht eindeutig festgelegt. So könnte z. B. in Abb. 3.3 die Gewichtskraft durch beliebige Pfeile dargestellt werden, die zu dem Pfeil, der $\vec{F}_G$ repräsentiert, parallel, gleich gerichtet und gleich lang sind.

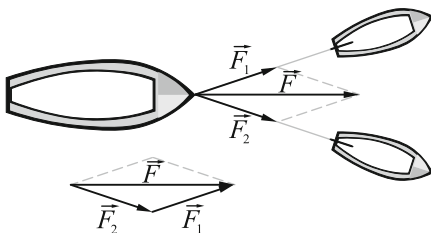


Abb. 3.2: Addition von Kräften

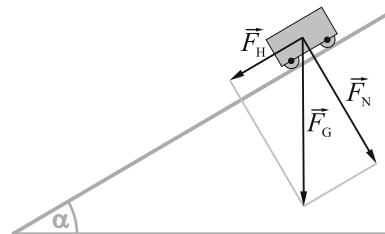


Abb. 3.3: Zerlegung einer Kraft

¹Das Operationszeichen „+“ steht hierbei für die Zusammensetzung von Kräften bzw. die später noch behandelte Vektoraddition. Es lässt sich nicht durch eine Addition der Beträge der Kräfte interpretieren. Offensichtlich gilt nicht $|\vec{F}_G| = |\vec{F}_H| + |\vec{F}_N|$, sondern (da in dem Beispiel die Teilkräfte $\vec{F}_H$ und $\vec{F}_N$ senkrecht zueinander wirken) $|\vec{F}_G| = \sqrt{|\vec{F}_H|^2 + |\vec{F}_N|^2}$.

Wir betrachten im Folgenden *Verschiebungen*. Eine Abbildung der Ebene oder des Raumes auf sich selbst ist eine Verschiebung, falls für zwei beliebige Punkte P, Q (der Ebene bzw. des Raumes) und ihre Bildpunkte P', Q' gilt:

- Die Abstände zwischen Original- und Bildpunkten sind gleich:
 $|PP'| = |QQ'|$;
- Die durch Original- und Bildpunkte verlaufenden Geraden sind parallel zueinander: $PP' \parallel QQ'$;²
- Die Pfeile $\overrightarrow{PP'}$ und $\overrightarrow{QQ'}$ sind gleich orientiert. Falls $\overrightarrow{PP'}$ und $\overrightarrow{QQ'}$ auf verschiedenen Geraden liegen, ist dies unter den Voraussetzungen $PP' \parallel QQ'$ und $|PP'| = |QQ'|$ genau dann der Fall, wenn PQ und $P'Q'$ parallel sind. Gehören P, P', Q und Q' einer Geraden an, so sind die Pfeile $\overrightarrow{PP'}$ und $\overrightarrow{QQ'}$ gleich orientiert, falls sie identisch sind oder falls P' und Q zwischen P und Q' oder P und Q' zwischen P' und Q liegen, siehe Abb. 3.4.

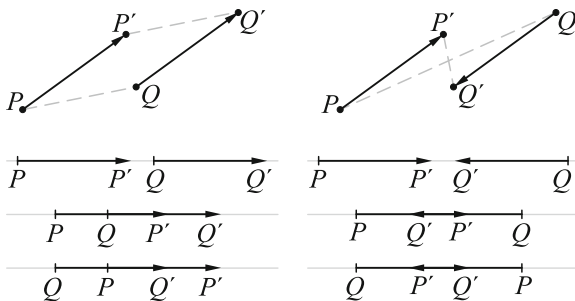


Abb. 3.4: Gleich und entgegengesetzt orientierte Pfeile

Da diese Bedingungen für beliebige Punkte und ihre Bildpunkte gelten, ist eine Verschiebung durch die Vorgabe eines einzigen Punktes und seines Bildpunktes oder eines „Verschiebungspfeils“ eindeutig bestimmt. Damit lassen sich für beliebige Punkte der Ebene oder des Raumes die Bildpunkte konstruieren. Bilder geradlinig begrenzter geometrischer Figuren und Körper entstehen, indem die Bilder der Eckpunkte konstruiert werden. Dabei wird die Tatsache genutzt, dass Verbindungsstrecken von Punkten auf die Verbindungsstrecken der entsprechenden Bildpunkte abgebildet werden (siehe Abb. 3.5 und 3.6).

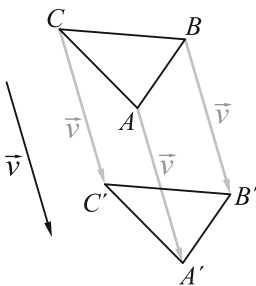


Abb. 3.5: Verschiebung in der Ebene

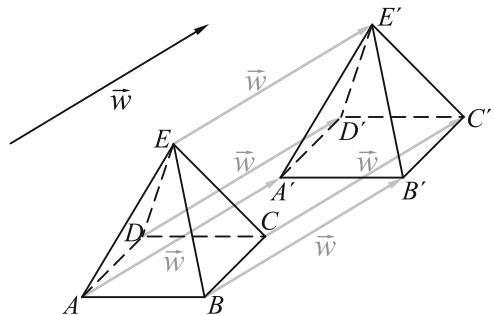


Abb. 3.6: Verschiebung im Raum

²Parallelität wird hier so aufgefasst, dass sie die Identität mit einschließt, also jede Gerade zu sich selbst parallel ist.

3.1.2 Pfeilklassen

Die im vorangegangenen Abschnitt betrachteten Beispiele weisen wesentliche Gemeinsamkeiten auf: Sowohl Geschwindigkeiten und Kräfte als auch Verschiebungen lassen sich durch Pfeile (gerichtete Strecken) beschreiben. Dabei ist es unwesentlich, an welchen Positionen die kennzeichnenden Pfeile gezeichnet werden, entscheidend für die Beschreibung einer Geschwindigkeit, Kraft oder Verschiebung sind lediglich Länge, Richtung und Richtungssinn eines Pfeils. Die Abstraktion von der Lage führt zu dem Konzept der Pfeilklassse, welches Pfeile gleicher Länge und Richtung sowie gleichen Richtungssinns identifiziert. Dazu führen wir eine Relation „parallelgleich“ zwischen Pfeilen ein und fassen dann Pfeile, die zueinander in dieser Relation stehen, zu Klassen zusammen.

Definition 3.1

Zwei Pfeile $\overrightarrow{AB}$ und $\overrightarrow{CD}$ heißen *parallelgleich*, falls sie gleich lang und gleichsinnig parallel sind, d. h. wenn gilt:

- Die Abstände zwischen Anfangs- und Endpunkt sind bei beiden Pfeilen gleich: $|AB| = |CD|$;
- Die Geraden, denen die beiden Pfeile angehören, sind parallel: $AB \parallel CD$;
- Die Pfeile $\overrightarrow{AB}$ und $\overrightarrow{CD}$ sind gleich orientiert (siehe S. 88). ◆

Satz 3.1

Die Relation „parallelgleich“ ist eine Äquivalenzrelation auf der Menge aller Pfeile der Ebene bzw. des Raumes, d. h.

- Jeder Pfeil ist zu sich selbst parallelgleich (Reflexivität).
- Ist ein Pfeil $\overrightarrow{AB}$ parallelgleich zu einem Pfeil $\overrightarrow{CD}$, so ist auch $\overrightarrow{CD}$ parallelgleich zu $\overrightarrow{AB}$ (Symmetrie).
- Ist ein Pfeil $\overrightarrow{AB}$ parallelgleich zu einem Pfeil $\overrightarrow{CD}$ und $\overrightarrow{CD}$ parallelgleich zu einem Pfeil $\overrightarrow{EF}$, so ist $\overrightarrow{AB}$ parallelgleich zu $\overrightarrow{EF}$ (Transitivität).

Beweis: Die *Reflexivität* gilt unmittelbar nach der Definition 3.1, denn jeder Pfeil ist zu sich selbst gleich lang und gleich orientiert; die Forderung ii (Parallelität) von Definition 3.1 ist wegen der Fußnote 2 auf S. 88 ebenfalls erfüllt.

Die *Symmetrie* der Relation „parallelgleich“ liegt ebenfalls unmittelbar auf der Hand, denn sowohl die Gleichheit (bezüglich der Länge) als auch die Parallelität und die Gleichheit der Orientierung sind jeweils symmetrische Relationen.

Die *Transitivität* der Parallelgleichheit gilt ebenfalls: Aus $|AB| = |CD|$ und $|CD| = |EF|$ folgt unmittelbar $|AB| = |EF|$, aus $AB \parallel CD$ und $CD \parallel EF$ ergibt sich ebenfalls direkt $AB \parallel EF$. Auf eine exakte Begründung der Transitivität der Gleichorientierung wird verzichtet, anschaulich ist sie recht einleuchtend. □

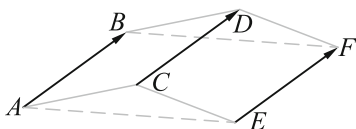


Abb. 3.7: Transitivität der Relation „parallelgleich“

Eine Äquivalenzrelation zerlegt die Menge, auf der sie definiert ist, in nicht leere, disjunkte Teilmengen, so genannte *Äquivalenzklassen*. Für Leserinnen und Leser, die mit diesem Begriff noch nicht vertraut sind, sei hier kurz auf ein sehr bekanntes Beispiel für Äquivalenzklassen eingegangen: die gebrochenen Zahlen. Unter einem Bruch versteht man ein konkretes Paar natürlicher Zahlen $(m; n)$ (das in der Form $\frac{m}{n}$ geschrieben wird). In diesem Sinne sind z. B. $\frac{1}{2}$ und $\frac{2}{4}$ verschiedene Brüche. Jedoch sind diese beiden Brüche *gleichwertig* (äquivalent), die zugehörige Äquivalenzrelation ist die Quotientengleichheit. Diese Äquivalenzrelation zerlegt die Menge aller Brüche in Äquivalenzklassen quotientengleicher Brüche. Die Äquivalenzklassen sind die gebrochenen Zahlen. Die Klassen werden ebenso bezeichnet wie die Brüche. Jedoch gehören zu der gebrochenen Zahl $\frac{1}{2}$ unendlich viele konkrete Brüche ($\frac{1}{2}$, $\frac{2}{4}$, $\frac{3}{6}$, $\frac{4}{8}$ usw.). Jeder dieser konkreten Brüche ist ein Repräsentant der gebrochenen Zahl (also der Äquivalenzklasse) $\frac{1}{2}$. Ebenso wie Brüche durch die Äquivalenzrelation „quotientengleich“ als gleichwertig zusammengefasst werden, erfolgt durch die Äquivalenzrelation „parallelgleich“ eine Zusammenfassung „gleichwertiger“ Pfeile zu *Pfeilklassen*. Gleichwertig bzw. äquivalent sind die so zusammengefassten Pfeile z. B. in der Hinsicht, dass sie dieselben Verschiebungen, Geschwindigkeiten oder Kräfte beschreiben.

Definition 3.2

Eine Pfeilkasse ist eine Äquivalenzklasse bezüglich der Äquivalenzrelation „parallelgleich“, d. h. als Pfeilkasse $\vec{u}$ bezeichnet man die Menge aller zu dem Pfeil $\vec{u}$ parallelgleichen Pfeile der Ebene bzw. des Raumes:

$$\vec{u} = \{\vec{x} \mid \vec{x} \text{ ist parallelgleich zu } \vec{u}\}.$$

Jeder Pfeil $\vec{x} \in \vec{u}$ heißt *Repräsentant* der Pfeilkasse $\vec{u}$. ◆

Bemerkungen:

- Für Pfeile und Pfeilklassen werden dieselben Bezeichnungen (wie $\vec{u}$ und $\overrightarrow{AB}$) genutzt. So bezeichnet $\vec{u}$ in Definition 3.2 innerhalb der geschweiften Klammer einen Pfeil, außerhalb der Klammer jedoch eine Pfeilkasse. Die Definition scheint zunächst unkorrekt zu sein, da ein Objekt $\vec{u}$ nicht durch sich selbst definiert werden kann. Da „ $\vec{u}$ “ jedoch verschiedene Bedeutungen hat, wird eine Pfeilkasse $\vec{u}$ anhand eines konkreten Pfeils $\vec{u}$ definiert. Die Bezeichnung unterschiedlicher Objekte mit denselben Symbolen ist vergleichbar mit Brüchen und gebrochenen Zahlen, die ebenfalls jeweils mit $\frac{m}{n}$ bezeichnet werden.
- In den Abbildungen 3.1 auf S. 86 sowie 3.5 und 3.6 auf S. 88 wurden bereits jeweils mehrere verschiedene Pfeile mit denselben Bezeichnungen $\vec{v}_2$, $\vec{v}$ bzw. $\vec{w}$ versehen, obwohl es sich trotz gleicher Bezeichnung um verschiedene Pfeile handelt. $\vec{v}_2$, $\vec{v}$ und $\vec{w}$ bezeichnen in den dortigen Beispielen somit Pfeilklassen.
- Im Folgenden werden vorrangig *Pfeilklassen* mit Symbolen wie $\vec{a}$ und $\vec{u}$ bezeichnet (nur wenn dies ausdrücklich erwähnt wird oder klar aus dem Kontext hervorgeht, ist damit ein konkreter Pfeil gemeint).

- Ist $\vec{u}$ eine Pfeilklassse in der Ebene oder im Raum, so existiert für jeden Punkt A der Ebene bzw. des Raumes ein Repräsentant (Pfeil) von $\vec{u}$ mit dem Anfangspunkt A (siehe Abb. 3.8).

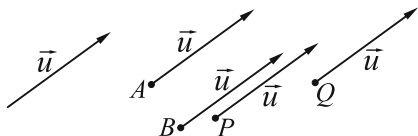


Abb. 3.8: Verschiedene Repräsentanten einer Pfeilklassse

3.1.3 Addition und skalare Multiplikation von Pfeilklassen

Pfeilklassenaddition

Definition 3.3

Es seien $\vec{u}$ und $\vec{v}$ Pfeilklassen sowie $\overrightarrow{AB} \in \vec{u}$ und $\overrightarrow{BC} \in \vec{v}$ Repräsentanten dieser Pfeilklassen (welche so gewählt sind, dass der Anfangspunkt des Repräsentanten von $\vec{v}$ gleich dem Endpunkt des Repräsentanten von $\vec{u}$ ist). Als *Summe der Pfeilklassen* $\vec{u}$ und $\vec{v}$ wird die Pfeilklassse bezeichnet, welche den Pfeil $\overrightarrow{AC}$ als einen Repräsentanten besitzt: $\vec{u} + \vec{v} = \{\vec{x} \mid \vec{x} \text{ ist parallelgleich zu } \overrightarrow{AC}\}$. ♦

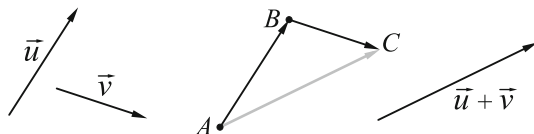


Abb. 3.9: Addition von Pfeilklassen

Bemerkungen:

- Damit die Definition 3.3 tatsächlich eine Operation beschreibt, muss gesichert sein, dass sich für beliebige Pfeilklassen die Summe bilden lässt. Dies ist der Fall, wenn für beliebige $\vec{u}$, $\vec{v}$ Repräsentanten mit der in Def. 3.3 verlangten Eigenschaft „der Anfangspunkt des Repräsentanten von $\vec{v}$ ist gleich dem Endpunkt des Repräsentanten von $\vec{u}$ “ existieren. Die vorherige Bemerkung sichert diese Tatsache, denn ist der Pfeil $\overrightarrow{AB}$ ein beliebiger Repräsentant von $\vec{u}$, so lässt sich ein Repräsentant $\overrightarrow{BC}$ von $\vec{v}$ mit dem Anfangspunkt B finden.
- In der Definition 3.3 werden der Begriff „Summe“ und das Operationszeichen „+“ verwendet, die vom Rechnen mit Zahlen her bekannt sind. Dennoch haben der Begriff und das Symbol zunächst keinen Bezug zum Rechnen mit Zahlen und sind rein geometrisch definiert. Gemeinsame Eigenschaften der Addition von Pfeilklassen und der Addition von Zahlen werden sich allerdings im Folgenden bei der Behandlung von Rechenregeln herausstellen.

Wird eine Operation von Äquivalenzklassen anhand von Repräsentanten definiert – wie in Definition 3.3 die Addition von Pfeilklassen mithilfe des Antragens repräsentierender Pfeile – so muss abgesichert werden, dass eine solche Definition *repräsentantenunabhängig* ist. Dies bedeutet, dass bei Wahl verschiedener Pfeile für $\vec{u}$ und $\vec{v}$ im Ergebnis Pfeile entstehen, welche dieselbe Pfeilklassse repräsentieren. Nur wenn dies der Fall ist, wird zwei Pfeilklassen $\vec{u}$ und $\vec{v}$ eindeutig

eine Summe $\vec{u} + \vec{v}$ zugeordnet. Würden sich je nach Wahl der Repräsentanten verschiedene Pfeilklassen als Summen ergeben, so beschriebe die Definition 3.3 keine sinnvolle Operation von Pfeilklassen mit eindeutigen Ergebnissen. Um die Bedeutung der Repräsentantenunabhängigkeit (mitunter auch „Wohldefiniertheit“ genannt) zu verdeutlichen, sei nochmals auf die Bruchrechnung verwiesen. Die Addition gebrochener Zahlen wird durch die Addition konkreter Brüche definiert. Die beiden gebrochenen Zahlen $\frac{1}{2}$ und $\frac{2}{3}$ lassen sich unter Verwendung verschiedener Repräsentanten (Brüche) addieren, z. B. $\frac{1}{2} + \frac{2}{3} = \frac{7}{6}$ und $\frac{2}{4} + \frac{6}{9} = \frac{42}{36}$. Die beiden als Ergebnisse entstehenden Brüche sind gleichwertig (äquivalent), sie repräsentieren also dieselbe gebrochene Zahl. Nur dadurch, dass dies für beliebige Repräsentanten beliebiger gebrochener Zahlen zutrifft, ist die Addition gebrochener Zahlen „wohldefiniert“. Wir zeigen mit dem folgenden Satz, dass die in der Definition 3.3 definierte Addition von Pfeilklassen ebenfalls repräsentantenunabhängig ist, zwei Pfeilklassen $\vec{u}$ und $\vec{v}$ durch diese Definition also eindeutig eine Summe $\vec{u} + \vec{v}$ zugeordnet wird.

Satz 3.2

Die Definition 3.3 ist repräsentantenunabhängig, d. h.: Sind $\overrightarrow{AB}$ und $\overrightarrow{PQ}$ Repräsentanten einer Pfeilklassse $\vec{u}$ sowie $\overrightarrow{BC}$ und $\overrightarrow{QR}$ Repräsentanten einer Pfeilklassse $\vec{v}$, so sind $\overrightarrow{AC}$ und $\overrightarrow{PR}$ Repräsentanten ein und derselben Pfeilklassse.

Beweis: Da nach Voraussetzung $\overrightarrow{AB}$ und $\overrightarrow{PQ}$ dieselbe Pfeilklassse repräsentieren, sind diese beiden Pfeile parallelgleich, es gilt also $|AB| = |PQ|$, $AB \parallel PQ$ und $\overrightarrow{AB}$ und $\overrightarrow{PQ}$ sind gleich orientiert (siehe die Definition 3.1 auf S. 89). Ebenso gilt $|BC| = |QR|$, $BC \parallel QR$, $\overrightarrow{BC}$ und $\overrightarrow{QR}$ sind gleich orientiert.

- i. Wegen $AB \parallel PQ$ gilt nach dem Stufenwinkelsatz an geschnittenen Parallelen $\beta = \beta'$ (siehe Abb. 3.10 a); wegen $BC \parallel QR$ gilt $\beta' = \beta''$, also $\beta = \beta''$. Aufgrund der Voraussetzungen $|AB| = |PQ|$ und $|BC| = |QR|$ sind die Dreiecke $\triangle ABC$ und $\triangle PQR$ nach dem Kongruenzsatz „sws“ kongruent, somit ist $|AC| = |PR|$.
- ii. Wegen der Kongruenz der Dreiecke $\triangle ABC$ und $\triangle PQR$ ist $\gamma = \rho$. Andererseits gilt wegen $BC \parallel QR$ nach dem Stufenwinkelsatz $\rho = \rho'$. Damit ist $\gamma = \rho'$, die Geraden PR und AC bilden mit BC also gleiche Stufenwinkel. Nach der Umkehrung des Stufenwinkelsatzes folgt daraus $AC \parallel PR$.

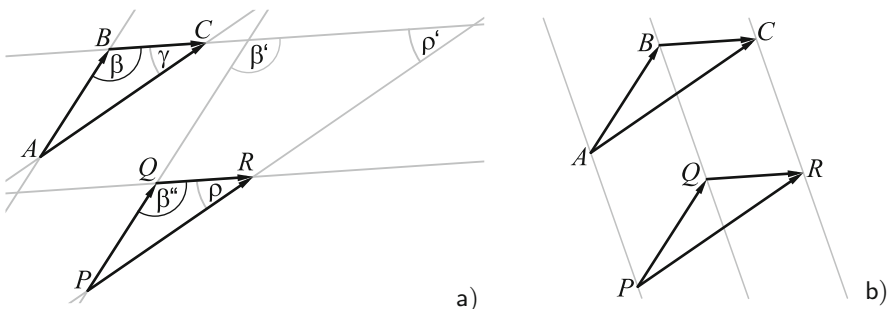


Abb. 3.10: Repräsentantenunabhängigkeit der Addition von Pfeilklassen

iii. Wegen der Gleichorientierung von $\overrightarrow{AB}$ und $\overrightarrow{PQ}$ sowie von $\overrightarrow{BC}$ und $\overrightarrow{QR}$ gilt $AP \parallel BQ$ und $BQ \parallel CR$ (siehe Abb. 3.10 b). Daraus folgt $AP \parallel CR$, also sind auch $\overrightarrow{AC}$ und $\overrightarrow{PR}$ gleich orientiert.

Damit wurde gezeigt, dass die Pfeile $\overrightarrow{AC}$ und $\overrightarrow{PR}$ nach Definition 3.1 parallelgleich sind. Sie repräsentieren also dieselbe Pfeilklassse. Unabhängig von der Wahl der Repräsentanten von $\vec{u}$ und $\vec{v}$ in der Definition 3.3 ergibt sich dieselbe Summe $\vec{u} + \vec{v}$. $\square$

Bemerkung: In dem Beweis von Satz 3.2 wurden nicht die Fälle betrachtet, dass $\overrightarrow{AB}$ und $\overrightarrow{PQ}$, $\overrightarrow{BC}$ und $\overrightarrow{QR}$ oder $\overrightarrow{AC}$ und $\overrightarrow{AR}$ auf einer Geraden liegen. In diesen Fällen kann analog vorgegangen werden, wobei sich einige Beweisschritte verkürzen. Aufgrund der einfacheren Darstellung wurde der Beweis zudem nur für Pfeile in der Ebene geführt, eine Übertragung auf Pfeile im Raum ist möglich.

Rechenregeln der Addition von Pfeilklassen

Die Pfeilklassenaddition hat grundlegende Eigenschaften, welche die Verwandtschaft zur Addition von Zahlen erkennen lassen. Die Aussagen des folgenden Satzes gelten sowohl auf der Menge aller Pfeilklassen der Ebene als auch auf der Menge aller Pfeilklassen des Raumes.

Satz 3.3

- A1. Für beliebige Pfeilklassen $\vec{u}$ und $\vec{v}$ gilt $\vec{u} + \vec{v} = \vec{v} + \vec{u}$ (Kommutativität).
- A2. Für beliebige Pfeilklassen $\vec{u}, \vec{v}, \vec{w}$ gilt $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$ (Assoziativität).
- A3. Es existiert eine Pfeilklassse $\vec{o}$, so dass für jede Pfeilklassse $\vec{u}$ gilt: $\vec{u} + \vec{o} = \vec{u}$ (Existenz einer „Nullpfeilklassse“).
- A4. Zu jeder Pfeilklassse $\vec{u}$ existiert eine Pfeilklassse $-\vec{u}$ mit $\vec{u} + (-\vec{u}) = \vec{o}$ (Existenz einer „entgegengesetzten Pfeilklassse“ zu jeder Pfeilklassse).

Beweis:

A1. Es seien der Pfeil $\overrightarrow{AB}$ ein Repräsentant von $\vec{u}$ sowie $\overrightarrow{BC}$ und $\overrightarrow{AD}$ Repräsentanten von $\vec{v}$. Es gilt offensichtlich $\overrightarrow{AB} + \overrightarrow{BC} = \overrightarrow{AC} = \overrightarrow{AD} + \overrightarrow{DC}$ (siehe Abb. 3.11 a). Es bleibt zu zeigen, dass $\overrightarrow{DC} \in \vec{u}$ ist (der Pfeil $\overrightarrow{DC}$ also der Pfeilklassse $\vec{u}$ angehört). Da wegen der Parallelgleichheit von $\overrightarrow{BC}$ und $\overrightarrow{AD}$ gilt $|BC| = |AD|$ und $BC \parallel AD$, ist das Viereck $ADCB$ ein Parallelogramm. Somit gilt auch $|AB| = |DC|$ und $AB \parallel DC$. Die Pfeile $\overrightarrow{AB}$ und $\overrightarrow{DC}$ sind deshalb parallelgleich und wegen $\overrightarrow{AB} \in \vec{u}$ ist auch $\overrightarrow{DC} \in \vec{u}$. Aus $\overrightarrow{AB} + \overrightarrow{BC} = \overrightarrow{AD} + \overrightarrow{DC}$ folgt damit $\vec{u} + \vec{v} = \vec{v} + \vec{u}$.

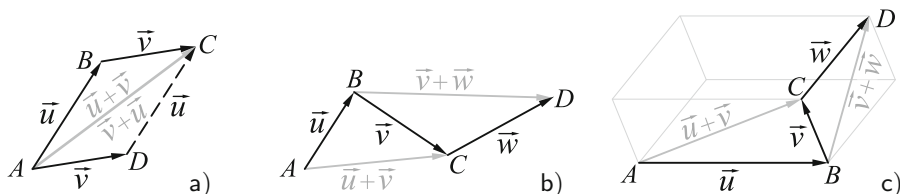


Abb. 3.11: Kommutativität und Assoziativität der Addition von Pfeilklassen

- A2. Es sei $\overrightarrow{AB}$ ein Repräsentant der Pfeilkategorie $\vec{u}$, $\overrightarrow{BC} \in \vec{v}$ und $\overrightarrow{CD} \in \vec{w}$. Dann ist nach Definition 3.3 $\overrightarrow{AC} \in \vec{u} + \vec{v}$ und $\overrightarrow{BD} \in \vec{v} + \vec{w}$. Die Summe $(\vec{u} + \vec{v}) + \vec{w}$ wird durch $\overrightarrow{AC} + \overrightarrow{CD} = \overrightarrow{AD}$ und $\vec{u} + (\vec{v} + \vec{w})$ durch $\overrightarrow{AB} + \overrightarrow{BD} = \overrightarrow{AD}$ repräsentiert, also gilt $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$, siehe Abb. 3.11 b, c.
- A3. Setzt man $\vec{o}$ als diejenige Pfeilkategorie, die alle Pfeile enthält, für die Anfangs- und Endpunkt identisch sind, so gilt für beliebige Pfeile $\overrightarrow{AB}$ einer beliebigen Pfeilkategorie $\vec{u}$: $\overrightarrow{AB} + \overrightarrow{BA} = \overrightarrow{AA}$ und somit $\vec{u} + \vec{o} = \vec{u}$.
- A4. Ist $\vec{u}$ eine beliebige Pfeilkategorie und $\overrightarrow{AB} \in \vec{u}$, so erfüllt die Pfeilkategorie $-\vec{u}$ die durch $\overrightarrow{BA}$ erzeugt wird, die Behauptung $\vec{u} + (-\vec{u}) = \vec{o}$. $\square$

Bemerkungen:

- Die Teile A3 und A4 des Satzes 3.3 beinhalten Existenzaussagen völlig unterschiedlicher Natur. Während A3 die Existenz einer einzigen „universellen“ Nullpfeilkategorie beinhaltet (welche die Eigenschaft $\vec{u} + \vec{o} = \vec{u}$ für alle Pfeilkategorien $\vec{u}$ besitzt), besagt A4, dass es für jede Pfeilkategorie $\vec{u}$ eine gewissermaßen „individuelle“ entgegengesetzte Pfeilkategorie gibt.
- Die Teile A2, A3 und A4 des Satzes beinhalten zusammengefasst, dass die Menge aller Pfeilkategorien der Ebene und die Menge aller Pfeilkategorien des Raumes mit der Operation „Addition von Pfeilkategorien“ Gruppen sind; wegen A1 handelt es sich sogar um Abelsche (kommutative) Gruppen. Für Leserinnen und Leser, die mit dem Begriff der Gruppe noch nicht vertraut sind, wird dieser im Folgenden definiert.

Exkurs: Gruppen

Definition 3.4

Eine nicht leere Menge G mit einer Operation $\circ$ heißt *Gruppe*, falls folgende Bedingungen erfüllt sind:

0. *Abgeschlossenheit*: Für zwei Elemente $a, b \in G$ ist auch $a \circ b \in G$.
1. *Assoziativität*: Für alle $a, b, c \in G$ gilt: $(a \circ b) \circ c = a \circ (b \circ c)$.
2. Es existiert ein *neutrales Element* e in G bezüglich $\circ$, d. h. es existiert $e \in G$ mit $a \circ e = e \circ a = a$ für alle $a \in G$.
3. Jedes Element a von G besitzt bezüglich der Operation $\circ$ ein *inverses Element*, d. h. ein Element b mit $a \circ b = b \circ a = e$. $\blacklozenge$

Die Forderung nach der Abgeschlossenheit müsste in der Definition des Begriffs Gruppe nicht gestellt werden, da sie bereits in dem Begriff der Operation enthalten ist. Von einer Operation $\circ$ in einer Menge wird verlangt, dass sie zwei beliebigen Elementen a, b ein Element $a \circ b$ der Menge zuordnet. Da diese Abgeschlossenheit ein wichtiges Merkmal von Gruppen ist, wurde sie in die Definition als 0. Eigenschaft aufgenommen.

Gruppen zählen zu den grundlegenden algebraischen Strukturen. Eine sehr bekannte Gruppe ist die Menge der ganzen Zahlen mit der Addition als Operation, man verwendet hierfür die Schreibweise $(\mathbb{Z}, +)$. Auch $(\mathbb{Q}, +)$ und $(\mathbb{R}, +)$ sind Gruppen. Neutrales Element ist hierbei jeweils die Null. Bezüglich der Multiplikation bilden $\mathbb{Q} \setminus \{0\}$ sowie $\mathbb{R} \setminus \{0\}$ Gruppen. Neutrales Element der Gruppen

$(\mathbb{Q} \setminus \{0\}, \cdot)$ und $(\mathbb{R} \setminus \{0\}, \cdot)$ ist jeweils die Zahl Eins. Die Null muss aus den Mengen der rationalen bzw. reellen Zahlen herausgenommen werden, da sie bezüglich der Multiplikation kein Inverses besitzt: es gibt keine Zahl b für die $0 \cdot b = 1$ gilt.

Gruppen treten auch in der Geometrie auf. So bilden die Menge aller Kongruenzabbildungen sowie die Menge aller Ähnlichkeitsabbildungen der Ebene bzw. des Raumes Gruppen bezüglich der Operation „Nacheinanderausführung von Abbildungen“.

Die genannten Beispiele für Gruppen verdeutlichen, dass die Gruppenstruktur Gemeinsamkeiten zwischen sehr unterschiedlichen Objekten der Arithmetik und der Geometrie aufdeckt. Insbesondere werden durch die Gruppeneigenschaften auch Analogien zwischen der Gruppe der reellen Zahlen mit der Operation „+“ und der Gruppe der Pfeilklassen der Ebene bzw. des Raumes bezüglich der Addition von Pfeilklassen strukturell zum Ausdruck gebracht. Das neutrale Element in der Gruppe der Pfeilklassen ist die „Nullpfeilklass“ und das inverse Element einer Pfeilklass die dazu entgegengesetzte Pfeilklass, siehe Satz 3.3.

Multiplikation von Pfeilklassen mit reellen Zahlen

Definition 3.5

Es seien $\vec{u}$ eine Pfeilklass der Ebene oder des Raumes mit einem Repräsentanten $\overrightarrow{AB}$ sowie λ eine reelle Zahl. Als *Produkt* $\lambda \cdot \vec{u}$ der Pfeilklass $\vec{u}$ mit der Zahl λ bezeichnet man die Pfeilklass, die durch den Pfeil $\overrightarrow{AC}$ mit folgenden Eigenschaften repräsentiert wird:

- $|AC| = |\lambda| \cdot |AB|$,
- falls $\lambda > 0$ ist, so gehört C dem Strahl AB und für $\lambda < 0$ dem zu AB entgegengesetzten Strahl an.

(Für $\lambda = 0$ ist $\lambda \cdot \vec{u}$ wegen $|AC| = |\lambda| \cdot |AB| = 0$ die Nullpfeilklass.)

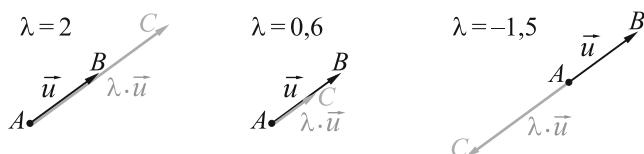


Abb. 3.12: Multiplikation von Pfeilklassen mit reellen Zahlen

Bemerkungen:

- Die Definition 3.5 weist Ähnlichkeiten mit Definitionen zentrischer Streckungen auf. Tatsächlich ist der durch die Definition festgelegte Punkt C Bild des Punktes B bei der zentrischen Streckung mit dem Streckungszentrum A und dem Streckungsfaktor λ .
- Die Multiplikation von Pfeilklassen mit reellen Zahlen ist eine *äußere Verknüpfung*, da Pfeilklassen mit Elementen einer anderen Menge (nämlich mit reellen Zahlen) verknüpft werden. Hingegen ist die Addition eine *innere Verknüpfung* (Operation), da jeweils zwei Pfeilklassen eine Pfeilklass (die Summe) zugeordnet wird.
- Statt von der „Multiplikation mit reellen Zahlen“ spricht man auch von der *skalaren Multiplikation*. Reelle Zahlen werden dabei als Skalare bezeichnet.

Da in der Definition 3.5 Pfeilklassen $\lambda \cdot \vec{u}$ anhand von Repräsentanten definiert werden, ist es – wie schon zu Definition 3.3 – notwendig, die Repräsentanten-unabhängigkeit dieser Definition nachzuweisen (siehe S. 91-92).

Satz 3.4

Sind $\overrightarrow{AB}$ und $\overrightarrow{A'B'}$ Repräsentanten einer Pfeilklass $\vec{u}$, so sind die in der Definition 3.5 festgelegten Pfeile $\overrightarrow{AC}$ und (analog dazu) $\overrightarrow{A'C'}$ Repräsentanten ein und derselben Pfeilklass.

Beweis: Es ist zu zeigen, dass die Pfeile $\overrightarrow{AC}$ und $\overrightarrow{A'C'}$ parallelgleich sind. Da es sich hierbei für $\lambda=0$ um Nullpfeile handelt (von denen klar ist, dass sie der Nullpfeilklass angehören) muss der Nachweis nur für $\lambda \neq 0$ geführt werden.

- Nach Voraussetzung ist $|AB|=|A'B'|$; da nach Definition 3.5 $|AC|=|\lambda| \cdot |AB|$ und $|A'C'|=|\lambda| \cdot |A'B'|$ gilt, folgt daraus $|AC|=|A'C'|$.
- Nach Voraussetzung sind die Geraden AB und $A'B'$ parallel zueinander. Da C auf der Geraden AB und C' auf $A'B'$ liegt, gilt auch $AC \parallel A'C'$.
- Es ist zu zeigen, dass $\overrightarrow{AC}$ und $\overrightarrow{A'C'}$ gleich orientiert sind. Dies ist (falls diese Pfeile nicht ein und derselben Geraden angehören) der Fall, wenn $AA' \parallel CC'$ gilt (siehe S. 88). Wegen i und ii sind A, C, A' und C' Eckpunkte eines Parallelogramms, es bleibt also zu zeigen, dass hierbei C und C' benachbarte Eckpunkte sind. Für $\lambda > 0$ gehört C dem Strahl AB und C' dem Strahl $A'B'$ an. Damit liegen C und C' in derselben Halbebene der durch die Pfeile $\overrightarrow{AC}$ und $\overrightarrow{A'C'}$ aufgespannten Ebene mit der Randgeraden AA' wie B und B' . Somit können C und C' nicht auf verschiedenen Seiten von AA' liegen. $\overrightarrow{CC'}$ ist somit keine Diagonale des Parallelogramms $AA'C'C$, damit sind C und C' benachbarte Punkte. $\overrightarrow{AA'}$ und $\overrightarrow{CC'}$ sind daher gegenüberliegende Seiten, woraus $AA' \parallel CC'$ folgt. Für $\lambda < 0$ lässt sich der Nachweis analog führen; auf den Nachweis für den Spezialfall, dass $\overrightarrow{AC}$ und $\overrightarrow{A'C'}$ auf einer Geraden liegen, verzichten wir hier. In jedem Falle sind $\overrightarrow{AC}$ und $\overrightarrow{A'C'}$ gleich orientiert. $\square$

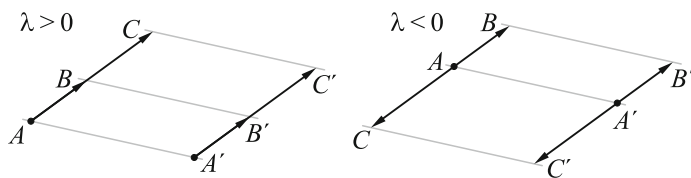


Abb. 3.13:
Repräsentanten-unabhängigkeit der Multiplikation von Pfeilklassen mit reellen Zahlen

Rechenregeln der Multiplikation von Pfeilklassen mit reellen Zahlen

Satz 3.5

- Für beliebige Pfeilklassen $\vec{u}$ gilt $1 \cdot \vec{u} = \vec{u}$.
- Für beliebige Pfeilklassen $\vec{u}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda \cdot \mu) \cdot \vec{u} = \lambda \cdot (\mu \cdot \vec{u})$ (Assoziativität).
- Für beliebige Pfeilklassen $\vec{u}, \vec{v}$ und beliebige $\lambda \in \mathbb{R}$ gilt $\lambda \cdot (\vec{u} + \vec{v}) = \lambda \cdot \vec{u} + \lambda \cdot \vec{v}$ (1. Distributivgesetz).

S4. Für beliebige Pfeilklassen $\vec{u}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda + \mu) \cdot \vec{u} = \lambda \cdot \vec{u} + \mu \cdot \vec{u}$ (2. Distributivgesetz).

Bemerkung: In der Eigenschaft S2 tritt das Zeichen „ $\cdot$ “ in zwei verschiedenen Bedeutungen auf: in $\lambda \cdot \mu$ steht es für das Produkt zweier reeller Zahlen, in $(\mu \cdot \vec{u})$ hingegen für die Multiplikation der Pfeilklassse $\vec{u}$ mit der Zahl μ . Der erste Punkt in $\lambda \cdot (\mu \cdot \vec{u})$ bezeichnet die Multiplikation der Pfeilklassse $(\mu \cdot \vec{u})$ mit der Zahl λ und der zweite Punkt in $(\lambda \cdot \mu) \cdot \vec{u}$ die Multiplikation von $\vec{u}$ mit der Zahl $\lambda \mu$. Ebenso wird das Operationszeichen „ $+$ “ in S4 in zwei Bedeutungen verwendet: $\lambda + \mu$ ist die Summe zweier reeller Zahlen, $\lambda \cdot \vec{u} + \mu \cdot \vec{u}$ die Summe zweier Pfeilklassen.

Beweis:

S1. Ist $\overrightarrow{AB}$ ein Repräsentant der Pfeilklassse $\vec{u}$, so ist nach Definition 3.5 der Pfeil $\overrightarrow{AC}$ mit $|AC| = 1 \cdot |AB|$ und der Eigenschaft, dass C auf dem Strahl AB liegt, ein Repräsentant von $1 \cdot \vec{u}$. Offensichtlich ist $C = B$ und somit $1 \cdot \vec{u} = \vec{u}$.

S2. Falls $\lambda > 0$ und $\mu > 0$ gilt und $\overrightarrow{AB}$ ein Repräsentant von $\vec{u}$ ist, so gilt nach Definition 3.5 für einen Repräsentanten $\overrightarrow{AC}$ von $\mu \cdot \vec{u}$, dass $|AC| = \mu \cdot |AB|$ ist und C auf dem Strahl AB liegt. Für einen Repräsentanten $\overrightarrow{AD}$ von $\lambda \cdot (\mu \cdot \vec{u})$ muss $|AD| = \lambda \cdot |AC| = \lambda \cdot (\mu \cdot |AB|)$ sein und D auf dem Strahl AC liegen. Dieser ist gleich dem Strahl AB . Ist $\overrightarrow{AE}$ ein Repräsentant von $(\lambda \cdot \mu) \cdot \vec{u}$, so ist $|AE| = \lambda \cdot \mu \cdot |AB|$ und (wegen $\lambda \cdot \mu > 0$) E ein Punkt des Strahls AB . Somit sind D und E identisch, es gilt also $\overrightarrow{AE} = \overrightarrow{AD}$ und somit die Behauptung.

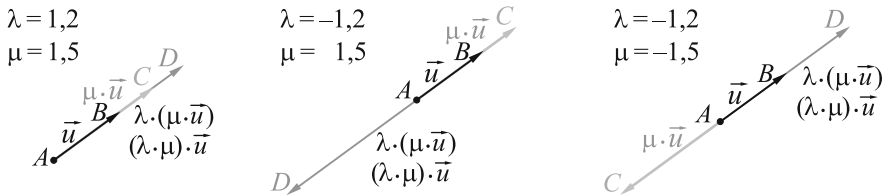


Abb. 3.14: Assoziativität der Multiplikation von Pfeilklassen mit reellen Zahlen

Für $\lambda < 0$ und $\mu > 0$ gilt für einen Repräsentanten $\overrightarrow{AC}$ von $\mu \cdot \vec{u}$ erneut, dass $|AC| = \mu \cdot |AB|$ ist und C auf dem Strahl AB liegt. Für einen Repräsentanten $\overrightarrow{AD}$ von $\lambda \cdot (\mu \cdot \vec{u})$ muss $|AD| = |\lambda| \cdot |AC| = |\lambda| \cdot (\mu \cdot |AB|)$ sein und D auf dem zu AC entgegengesetzten Strahl liegen. Dieser ist auch der zu AB entgegengesetzte Strahl. Ist $\overrightarrow{AE}$ ein Repräsentant von $(\lambda \cdot \mu) \cdot \vec{u}$, so ist $|AE| = |\lambda \cdot \mu| \cdot |AB|$ und (wegen $\lambda \cdot \mu < 0$) E ein Punkt des zu AB entgegengesetzten Strahls. Auch in diesem Falle sind somit D und E identisch, und es gilt die Behauptung.

Auf Nachweise für die Fälle $\lambda > 0, \mu < 0$ und $\lambda < 0, \mu < 0$ sowie für den Fall, dass $\lambda = 0$ oder $\mu = 0$ ist, wird hier verzichtet, siehe Aufgabe 4 auf S. 99.

S3. Um einen Repräsentanten für $\lambda \cdot (\vec{u} + \vec{v})$ zu ermitteln, werden zunächst $\vec{u}$ und $\vec{v}$ nach der Definition 3.3 (S. 91) addiert; anschließend wird ein Repräsentant $\overrightarrow{AD}$ von $\lambda \cdot (\vec{u} + \vec{v})$ bestimmt, siehe Abb. 3.15. Für $\overrightarrow{AD}$ gilt:

- $|AD| = |\lambda| \cdot |AC|$;
- D liegt auf dem Strahl AC , falls $\lambda > 0$ und auf dem zu AC entgegengesetzten Strahl, falls $\lambda < 0$.

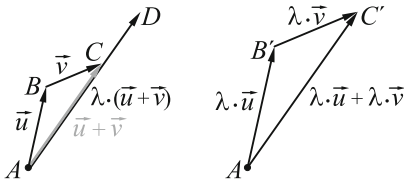


Abb. 3.15: 1. Distributivgesetz

Ist $\overrightarrow{AB}$ ein Repräsentant von $\lambda \cdot \vec{u}$ und $\overrightarrow{B'C'}$ ein Repräsentant von $\lambda \cdot \vec{v}$, so ist $\overrightarrow{AC'}$ ein Repräsentant von $\lambda \cdot \vec{u} + \lambda \cdot \vec{v}$. Das Dreieck $\triangle AB'C'$ geht durch eine zentrische Streckung mit dem Streckfaktor λ und dem Streckungszentrum A aus dem Dreieck $\triangle ABC$ hervor. Daher gilt:

- $|AC'| = |\lambda| \cdot |AC|$;
- für $\lambda > 0$ liegt C' auf dem Strahl AC , für $\lambda < 0$ auf dem zu AC entgegengesetzten Strahl.

Die Pfeile $\overrightarrow{AD}$ und $\overrightarrow{AC'}$ sind somit identisch, es gilt daher $\lambda(\vec{u} + \vec{v}) = \lambda\vec{u} + \lambda\vec{v}$.

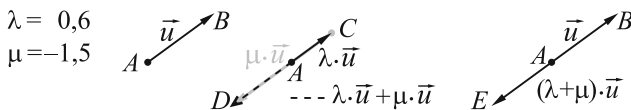
In dem bisher nicht betrachteten Fall, dass $\lambda = 0$ ist, gilt die Behauptung ebenfalls, da sowohl $\lambda \cdot (\vec{u} + \vec{v})$ als auch $\lambda \cdot \vec{u} + \lambda \cdot \vec{v}$ die Nullpfeilkategorie ist (siehe die Aufgabe 5 auf S. 99).

S4. Es müssen für λ und μ folgende Fälle betrachtet werden:

- $\lambda > 0, \mu > 0$;
- $\lambda < 0, \mu < 0$;
- eine der Zahlen λ und μ ist positiv, die andere negativ und $\lambda + \mu > 0$;
- eine der Zahlen λ und μ ist positiv, die andere negativ und $\lambda + \mu < 0$.

(Zusätzlich können noch die Fälle auftreten, dass λ oder μ oder $\lambda + \mu$ gleich Null ist; in diesen Fällen ist leicht zu überlegen, dass die Behauptung gilt.)

Im Folgenden wird der Beweis für den Fall $\lambda > 0, \mu < 0, \lambda + \mu < 0$ geführt, in den anderen Fällen ist analog vorzugehen (siehe die Aufgabe 7 auf S. 99).

Abb. 3.16:
2. Distributivgesetz

Ist $\overrightarrow{AB}$ ein Repräsentant von $\vec{u}$, so gilt für einen Repräsentanten $\overrightarrow{AC}$ von $\lambda \cdot \vec{u}$: $|AC| = |\lambda| \cdot |AB|$ und $\overrightarrow{AC}$ ist wegen $\lambda > 0$ gleich orientiert zu $\overrightarrow{AB}$. Für einen Repräsentanten $\overrightarrow{CD}$ von $\mu \cdot \vec{u}$ gilt $|CD| = |\mu| \cdot |AB|$ und $\overrightarrow{AC}$ und $\overrightarrow{CD}$ sind wegen $\mu < 0$ entgegengesetzt orientiert (siehe Abb. 3.16). Wegen $\lambda > 0, \mu < 0, \lambda + \mu < 0$ ist $|\mu| > |\lambda|$, daraus folgt $|CD| > |AC|$. Somit liegt D auf dem zu AC entgegengesetzten Strahl, der (da C auf dem Strahl AB liegt) auch zu AB entgegengesetzt ist. Es gilt $|AD| = |\mu| \cdot |AB| - |\lambda| \cdot |AB| = (|\mu| - |\lambda|) \cdot |AB|$.

Ein Repräsentant $\overrightarrow{AE}$ von $(\lambda + \mu) \cdot \vec{u}$ muss wegen $\lambda + \mu < 0$ entgegengesetzt orientiert zu $\overrightarrow{AB}$ sein, und es muss $|AE| = |\lambda + \mu| \cdot |AB|$ gelten.

Wegen $\mu < 0$ ist $|\mu| = -\mu$, also $|\lambda + \mu| = |\lambda - |\mu|| = ||\mu| - \lambda|$. Wegen $|\mu| > |\lambda|$ ist $|\mu| - \lambda$ positiv, also $||\mu| - \lambda| = |\mu| - \lambda$. Wegen $\lambda > 0$ lässt sich λ durch $|\lambda|$ ersetzen, es gilt also $|\lambda + \mu| = |\mu| - |\lambda|$. Die Repräsentanten $\overrightarrow{AD}$ von $\lambda \cdot \vec{u} + \mu \cdot \vec{u}$ und $\overrightarrow{AE}$ von $(\lambda + \mu) \cdot \vec{u}$ sind somit identisch, es gilt $\lambda \cdot \vec{u} + \mu \cdot \vec{u} = (\lambda + \mu) \cdot \vec{u}$. $\square$

Mit den Rechenregeln A1-A4 (siehe S. 93) der Pfeilklassenaddition und S1-S4 der Multiplikation von Pfeilklassen mit reellen Zahlen wurden grundlegende strukturelle Aussagen hergeleitet, deren Bedeutung weit über die hier betrachteten Pfeilklassen hinaus reicht. Die Eigenschaften wurden hier zunächst konkret für das Beispiel der Pfeilklassen auf elementargeometrischem Wege hergeleitet. Dieselben Rechenregeln werden im Folgenden für völlig andere Objekte (die zunächst keinerlei geometrische Bedeutung besitzen) erneut auftreten. Später wird sich sogar zeigen, dass die Regeln A1-A4 und S1-S4 ausreichend sind, um den Begriff des Vektorraumes zu bestimmen.

3.1.4 Aufgaben zu Abschnitt 3.1

1. Ein Flugzeug fliegt exakt nach Norden mit einer Eigengeschwindigkeit von $150 \frac{\text{km}}{\text{h}}$. Welchen „Kurs“ nimmt das Flugzeug, wenn Südost-Wind mit einer Geschwindigkeit von $30 \frac{\text{km}}{\text{h}}$ herrscht? Ermitteln Sie graphisch und rechnerisch den Winkel der Flugrichtung zur Nordrichtung sowie den Betrag der resultierenden Flugeschwindigkeit.
2. Ein 6000 kg schwerer Wagen (mit dem Betrag $|\vec{F}_G| \approx 59\,000 \text{ N}$ der Gewichtskraft) rollt eine geneigte Ebene mit einem Neigungswinkel von 30° hinunter. Bestimmen Sie graphisch und rechnerisch die Beträge der Hangabtriebskraft $\vec{F}_H$ und der Normalkraft $\vec{F}_N$ (siehe Abb. 3.3 auf S. 87).
3. Gegeben ist ein Viereck $ABCD$. Drücken Sie die durch $\overrightarrow{AC}$, $\overrightarrow{AD}$ und $\overrightarrow{BD}$ repräsentierten Pfeilklassen mithilfe der Pfeilklassen $\vec{a}$, $\vec{b}$ und $\vec{c}$ aus.

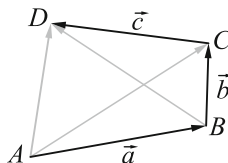


Abb. 3.17: Skizze zu Aufgabe 3

4. Beweisen Sie den Teil S2 des Satzes 3.5 (siehe S. 96) für die Fälle $\lambda > 0, \mu < 0$ und $\lambda < 0, \mu < 0$ sowie für den Fall, dass $\lambda = 0$ oder $\mu = 0$ ist.
5. Begründen Sie anhand der Definition 3.5 auf S. 95, dass für beliebige Pfeilklassen $\vec{u}$ gilt: $0 \cdot \vec{u}$ ist die Nullpfeilklass.
6. Das erste Distributivgesetz (Satz 3.5 S3, siehe S. 96) wurde in Abb. 3.15 für eine positive Zahl λ skizziert. Fertigen Sie eine Skizze für $\lambda < 0$ an. Überprüfen Sie, ob der auf S. 97 geführte Beweis auch für diesen Fall zutrifft.
7. Führen Sie Beweise der Eigenschaft S4 des Satzes 3.5 (siehe S. 96f.) für die Fälle $\lambda > 0, \mu > 0$ sowie $\lambda < 0, \mu < 0$.

3.2 Vektoren in arithmetischen Kontexten

3.2.1 Stücklisten, Farben

Produktsortimente und aus mehreren Komponenten bestehende Produkte wie Regalsysteme enthalten oft unterschiedliche Stückzahlen verschiedener Bauteile.

Beispiel 3.1

Ein Modellbahnhersteller bietet verschiedene Gleisbauteile sowohl einzeln als auch in Sortimentskästen an.

	Basis- sortiment	Ergänzungs- sortiment 1	Ergänzungs- sortiment 2
Gleisstück gerade (168,9 mm)	3	8	15
Gleisstück gebogen (45°)	8	4	8
Anschluss-Gleisstück	1	0	1
Weiche links	0	1	2
Weiche rechts	0	1	2
Weichenantrieb	0	2	4

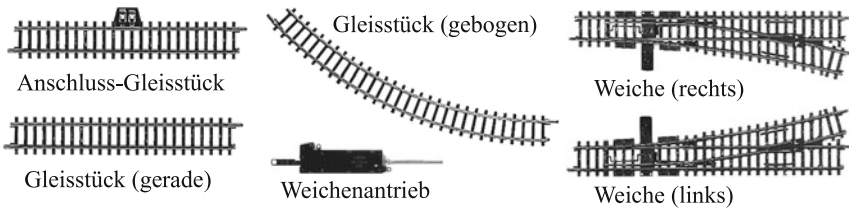


Abb. 3.18: Gleisteile für Modelleisenbahnen

Die Bestandteile der Sortimente lassen sich durch geordnete 6-Tupel natürlicher Zahlen angeben, die meist in Spaltenschreibweise geschrieben werden:

$$B = \begin{pmatrix} 3 \\ 8 \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \quad E_1 = \begin{pmatrix} 8 \\ 4 \\ 0 \\ 1 \\ 1 \\ 2 \end{pmatrix}, \quad E_2 = \begin{pmatrix} 15 \\ 8 \\ 1 \\ 2 \\ 2 \\ 4 \end{pmatrix}.$$

Unter Verwendung dieser 6-Tupel (bzw. allgemein mit n -Tupeln) lassen sich Rechenoperationen ausführen: Man addiert zwei n -Tupel, indem man jeweils zusammengehörige (an derselben Stelle stehende) Komponenten der beiden addiert und multipliziert ein n -Tupel mit einer Zahl k , indem man jede Komponente des n -Tupels mit k multipliziert. So lässt sich z. B. auf übersichtliche Weise darstellen, wie viele der jeweiligen Bauteile insgesamt in 18 Basissortimenten, 10 Ergänzungssortimenten 1 und 5 Ergänzungssortimenten 2 enthalten sind:

$$N = 18 \cdot \begin{pmatrix} 3 \\ 8 \\ 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} + 10 \cdot \begin{pmatrix} 8 \\ 4 \\ 0 \\ 1 \\ 1 \\ 2 \end{pmatrix} + 5 \cdot \begin{pmatrix} 15 \\ 8 \\ 1 \\ 2 \\ 2 \\ 4 \end{pmatrix} = \begin{pmatrix} 54 \\ 144 \\ 18 \\ 0 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 80 \\ 40 \\ 0 \\ 10 \\ 10 \\ 20 \end{pmatrix} + \begin{pmatrix} 75 \\ 40 \\ 5 \\ 10 \\ 10 \\ 20 \end{pmatrix} = \begin{pmatrix} 209 \\ 224 \\ 23 \\ 20 \\ 20 \\ 40 \end{pmatrix}.$$

Farbmischung in der elektronischen Bildwiedergabe

Fernsehgeräte und Monitore enthalten dicht benachbarte rote, grüne und blaue „Leuchtpunktchen“ (Subpixel), mithilfe derer sie für jeden Bildpunkt (Pixel) jede darstellbare Farbe mischen. Das zu Grunde liegende Farbmodell heißt wegen der verwendeten Grundfarben RGB-Modell. In diesem Modell wird jede Farbe durch ein Tripel $\begin{pmatrix} r \\ g \\ b \end{pmatrix}$ beschrieben.³ Im Grundzustand (kein Subpixel leuchtet) wird ein Bildschirm als schwarz angenommen. Durch die Lichtintensitäten der drei Subpixel, die durch die Komponenten r, g, b des RGB-Tripels beschrieben werden, werden Helligkeiten addiert und Farben gemischt.

Da das Auge nicht beliebig kleine Farbabweichungen wahrnimmt, genügt es im Allgemeinen, für jede Komponente 256 Werte zu unterscheiden.⁴ Die Komponenten r, g, b von Farbtripeln werden meist als natürliche Zahlen von 0 bis 255 oder (wie im Folgenden) als rationale Zahlen des Intervalls $[0, 1]$ dargestellt.

Die Addition von RGB-Tripeln ist durchaus eine sinnvolle Operation mit Farben. Die Addition zweier RGB-Tripel $\begin{pmatrix} r_1 \\ g_1 \\ b_1 \end{pmatrix}$ und $\begin{pmatrix} r_2 \\ g_2 \\ b_2 \end{pmatrix}$ lässt sich durch zwei Scheinwerfer dieser Lichtfarben veranschaulichen, die gemeinsam eine schwarze Fläche beleuchten. Dabei addieren sich die Intensitäten der drei Farbkomponenten beider Scheinwerfer. Jedoch kann keine der resultierenden Komponenten den Maximalwert 1 überschreiten und als Gesamtfarbe kann es keine hellere Farbe als reines Weiß geben. Eine sinnvolle Addition zweier Farbtripel ist daher durch

$$\begin{pmatrix} r_1 \\ g_1 \\ b_1 \end{pmatrix} + \begin{pmatrix} r_2 \\ g_2 \\ b_2 \end{pmatrix} := \begin{pmatrix} \min(r_1 + r_2, 1) \\ \min(g_1 + g_2, 1) \\ \min(b_1 + b_2, 1) \end{pmatrix}$$

definiert.

Die Addition von Farben ist in einem Bildbearbeitungsprogramm wie Adobe Photoshop oder der freien Software The Gimp nachvollziehbar. Werden drei Kreise mit den Grundfarben Rot, Grün und Blau auf jeweils eine Ebene über einer schwarzen Hintergrundebene gelegt und wird für diese Ebenen der Modus „Addition“ bzw. „Aufhellen“ (der ebenfalls die Addition der r -, g - und b -Komponenten übereinander liegender Pixel bewirkt) gewählt, so ergibt sich ein Bild wie in Abb. 3.19.⁵ Pixel, die im Durchschnitt aller drei Kreise liegen, nehmen die Farbe Weiß an:

$$R + G + B = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = W.$$

³Die rot-, grün- und blau-Komponenten von Farben werden mit Klein-, Farben selbst mit Großbuchstaben bezeichnet.

⁴Die Festlegung auf 256 Werte erfolgte, weil dadurch genau 8 Bit an Informationen für jede Komponente notwendig sind ($2^8 = 256$).

⁵Die Internetseite zu diesem Buch enthält eine Farbversion dieser Abbildung sowie Dateien, mithilfe derer die Farbadition in Bildbearbeitungssoftware nachvollzogen werden kann.

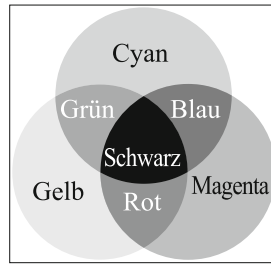
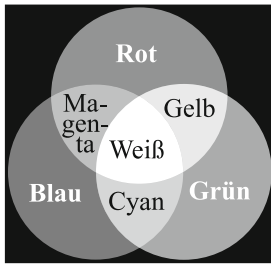


Abb. 3.19:
RGB-Farbmischung

Abb. 3.20:
CMY-Farbmischung

Als Summe von Rot und Grün ergibt sich die Farbe Gelb (Y, von Yellow), als Summe von Rot und Blau Magenta (M) sowie von Grün und Blau Cyan (C):

$$R+G = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} = Y, \quad R+B = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} = M, \quad G+B = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix} = C.$$

Die Summe zweier Grundfarben ist dabei jeweils die Komplementärfarbe der nicht an der Summenbildung beteiligten Grundfarbe; somit sind Gelb und Blau, Rot und Cyan sowie Grün und Magenta jeweils Paare von Komplementärfarben.

Die Komplementärfarben der RGB-Grundfarben sind vor allem für den Druck bedeutsam. Hier erfolgt die Farbmischung umgekehrt zur Darstellung auf Bildschirmen, da Papier im unbedruckten Zustand die größte Helligkeit aufweist. In seinem Ausgangszustand reflektiert ideales weißes Papier die gesamte auftreffende Helligkeit. Durch das Aufbringen von Farbpigmenten erfolgt eine Subtraktion von Helligkeitsanteilen, siehe Abb. 3.20. Werden Pigmente aufgetragen, die bestimmte Farben absorbieren, nimmt Papier die Farbe des „Restlichtes“ an, bei der es sich um die Komplementärfarbe der absorbierten Farbe handelt. Somit wird die Farbwiedergabe im Druck durch das subtraktive CMY-Modell mit den Grundfarben Cyan, Magenta und Gelb beschrieben. Zwischen dem RGB- und dem CMY-Tripel einer Farbe besteht der Zusammenhang

$$\begin{pmatrix} c \\ m \\ y \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} - \begin{pmatrix} r \\ g \\ b \end{pmatrix}.$$

Die Subtraktion erfolgt analog zur Addition von n -Tupeln komponentenweise. Jede der auftretenden Komponenten hat einen Wert von 0 bis 1.

In der Praxis ist die Transformation vom RGB- in den CMY-Farbraum (die durch den Computer bzw. Drucker vorgenommen werden muss, wenn z. B. ein mit einer Digitalkamera aufgenommenes Bild gedruckt wird) allerdings komplizierter, da weder Papier noch Farbstoffe ein ideales Verhalten aufweisen. Insbesondere kann eine reale Druckmaschine aus den drei Grundfarben kein wirkliches Schwarz erzeugen. Als vierte Farbe kommt daher Schwarz zur Anwendung – das CMY-Modell erfährt eine Erweiterung zum CMYK-Modell (K für Black).

Farben werden in diesem Modell durch CMYK-Quadrupel $\begin{pmatrix} c \\ m \\ y \\ k \end{pmatrix}$ beschrieben.

Die Transformation von Farben aus dem RGB- in das CMYK-Modell wird als *Farbseparation* bezeichnet, siehe z. B. Nyman (2004).

3.2.2 n -Tupel

In dem vorangegangenen Abschnitt wurden bereits einige Beispiele für Zahlen- n -Tupel behandelt, nämlich Stücklisten, RGB-Tripel und CMYK-Quadrupel. Als weitere Beispiele traten in den Abschnitten 1.1-1.3 Lösungspaare, -tripel und -quadrupel linearer Gleichungssysteme auf. Im Folgenden werden verallgemeinernd n -Tupel reeller Zahlen und Rechenoperationen mit ihnen betrachtet.

Unter einem n -Tupel reeller Zahlen versteht man eine geordnete Liste

$$\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \text{ mit } x_1, x_2, \dots, x_n \in \mathbb{R}.$$

Die Zahlen $x_1, x_2, \dots, x_n$ heißen Komponenten dieses n -Tupels.

Die Menge aller n -Tupel reeller Zahlen wird mit $\mathbb{R}^n$ bezeichnet. Dabei handelt es sich um das n -fache Mengenprodukt $\underbrace{\mathbb{R} \times \mathbb{R} \times \dots \times \mathbb{R}}_n$.

Im Folgenden werden n -Tupel mit den Symbolen bezeichnet, die bereits in dem

Abschnitt 3.1.2 für Pfeilklassen verwendet wurden, z. B. $\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$, $\vec{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$.

Dass diese Bezeichnungsweise sinnvoll ist, wird sich anhand der Tatsache zeigen, dass enge Beziehungen zwischen n -Tupeln und Pfeilklassen bestehen und dieselben Rechenregeln angewendet werden können. Später wird diese Symbolik dann allgemein für Vektoren verwendet, da sich heraus stellt, dass sowohl n -Tupel als auch Pfeilklassen Beispiele für Vektoren sind.

Addition von n -Tupeln, Multiplikation von n -Tupeln mit reellen Zahlen

Definition 3.6

Es seien $\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}$, $\vec{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$ n -Tupel und $\lambda \in \mathbb{R}$.

- Als *Summe der n -Tupel $\vec{x}$ und $\vec{y}$* bezeichnet man das n -Tupel, welches durch Addition der einander entsprechenden Komponenten von $\vec{x}$ und $\vec{y}$ entsteht:

$$\vec{x} + \vec{y} = \begin{pmatrix} x_1 + y_1 \\ x_2 + y_2 \\ \vdots \\ x_n + y_n \end{pmatrix}$$

- Als *Produkt des n -Tupels $\vec{x}$ mit der reellen Zahl λ* wird das n -Tupel bezeichnet, das durch Multiplikation jeder Komponente von $\vec{x}$ mit λ entsteht:

$$\lambda \cdot \vec{x} = \begin{pmatrix} \lambda x_1 \\ \lambda x_2 \\ \vdots \\ \lambda x_n \end{pmatrix}$$

◆

Für n -Tupel gelten dieselben Rechenregeln, die bereits in den Sätzen 3.3 (siehe S. 93) und 3.5 (S. 96) für Pfeilklassen aufgestellt wurden.

Satz 3.6

- A1. Für beliebige n -Tupel $\vec{x}$ und $\vec{y}$ gilt $\vec{x} + \vec{y} = \vec{y} + \vec{x}$
(Kommutativität der Addition).
- A2. Für beliebige n -Tupel $\vec{x}, \vec{y}, \vec{z}$ gilt $(\vec{x} + \vec{y}) + \vec{z} = \vec{x} + (\vec{y} + \vec{z})$
(Assoziativität der Addition).
- A3. Es existiert ein n -Tupel $\vec{o}$, so dass für jedes n -Tupel $\vec{x}$ gilt: $\vec{x} + \vec{o} = \vec{x}$
(Existenz eines „Null- n -Tupels“).
- A4. Zu jedem n -Tupel $\vec{x}$ existiert ein n -Tupel $-\vec{x}$ mit $\vec{x} + (-\vec{x}) = \vec{o}$
(Existenz eines „Gegen- n -Tupels“ zu jedem n -Tupel).
- S1. Für beliebige n -Tupel $\vec{x}$ gilt $1 \cdot \vec{x} = \vec{x}$.
- S2. Für beliebige n -Tupel $\vec{x}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda \cdot \mu) \cdot \vec{x} = \lambda \cdot (\mu \cdot \vec{x})$
(Assoziativität der Multiplikation von n -Tupeln mit reellen Zahlen).
- S3. Für beliebige n -Tupel $\vec{x}, \vec{y}$ und beliebige $\lambda \in \mathbb{R}$ gilt $\lambda \cdot (\vec{x} + \vec{y}) = \lambda \cdot \vec{x} + \lambda \cdot \vec{y}$
(1. Distributivgesetz).
- S4. Für beliebige n -Tupel $\vec{x}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda + \mu) \cdot \vec{x} = \lambda \cdot \vec{x} + \mu \cdot \vec{x}$
(2. Distributivgesetz).

Bemerkungen:

- In den Eigenschaften S2-S4 treten die Zeichen „+“ und „·“ in verschiedenen Bedeutungen auf, nämlich für die Addition reeller Zahlen und die Addition von n -Tupeln sowie für die Multiplikation reeller Zahlen und die Multiplikation von n -Tupeln mit reellen Zahlen (siehe hierzu auch die Bemerkung auf S. 97 zu dem Satz 3.5).
- Aufgrund von A1-A4 und der Tatsache, dass die Summe zweier n -Tupel stets wiederum ein n -Tupel ist (es sich also bei der in Def. 3.6 definierten Addition um eine Operation $+: \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ handelt) ist $\mathbb{R}^n$ mit der Addition von n -Tupeln eine *Abelsche* (kommutative) *Gruppe*, siehe S. 94.
- Die Eigenschaften A1-S4 treffen meist nur dann uneingeschränkt zu, wenn die n -Tupel-Addition und die Multiplikation mit reellen Zahlen auf der Menge $\mathbb{R}^n$ aller n -Tupel reeller Zahlen betrachtet werden. Bei den in dem Abschnitt 3.2.1 betrachteten Beispielen ist dies nicht der Fall. So sind Stücklisten mit negativen Komponenten nicht sinnvoll, womit die Eigenschaft A4 auf dieses Beispiel nicht zutrifft. Ebenso wenig können die rot-, grün oder blau-Komponenten von Farben negativ sein, zudem musste auf S. 101 eine spezielle Addition für RGB-Tripel definiert werden, da Komponenten von RGB-Tripeln dem Intervall $[0; 1]$ angehören. Trotz dieser Einschränkungen treffen viele der Aussagen des Satzes 3.6 auf das Rechnen mit Stücklisten und Farbtripeln zu.
- Es gibt auch spezielle Mengen von n -Tupeln, die zwar nur Teilmengen von $\mathbb{R}^n$ sind, auf denen aber die n -Tupel-Addition und die Multiplikation mit reellen Zahlen dennoch uneingeschränkt ausführbar sind und für die alle Aussagen des Satzes 3.6 gelten. Dazu zählen die Lösungsmengen homogener linearer Gleichungssysteme, vgl. Abschnitt 1.3.4 (siehe dazu Aufgabe 3 auf S. 105).

Beweis: Um den Satz 3.6 zu beweisen sind vor allem Rechengesetze für reelle Zahlen auf n -Tupel zu übertragen. Da dies für die meisten Aussagen des Satzes in recht analoger Weise erfolgt, werden hier nur einige der Aussagen bewiesen.

A2. Für beliebige n -Tupel $\vec{x}, \vec{y}, \vec{z}$ gilt nach Definition 3.6:

$$\begin{aligned} (\vec{x} + \vec{y}) + \vec{z} &= \left[\begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} + \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} \right] + \begin{pmatrix} z_1 \\ z_2 \\ \vdots \\ z_n \end{pmatrix} = \begin{pmatrix} x_1 + y_1 \\ x_2 + y_2 \\ \vdots \\ x_n + y_n \end{pmatrix} + \begin{pmatrix} z_1 \\ z_2 \\ \vdots \\ z_n \end{pmatrix} \\ &= \begin{pmatrix} x_1 + y_1 + z_1 \\ x_2 + y_2 + z_2 \\ \vdots \\ x_n + y_n + z_n \end{pmatrix} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} + \begin{pmatrix} y_1 + z_1 \\ y_2 + z_2 \\ \vdots \\ y_n + z_n \end{pmatrix} \\ &= \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} + \left[\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} + \begin{pmatrix} z_1 \\ z_2 \\ \vdots \\ z_n \end{pmatrix} \right] = \vec{x} + (\vec{y} + \vec{z}). \end{aligned}$$

Eine kürzere Begründung kann folgendermaßen gegeben werden:

Da für beliebige einander entsprechende Komponenten x_i, y_i, z_i der n -Tupel $\vec{x}, \vec{y}, \vec{z}$ gilt $(x_i + y_i) + z_i = x_i + (y_i + z_i)$, ist die Behauptung $(\vec{x} + \vec{y}) + \vec{z} = \vec{x} + (\vec{y} + \vec{z})$ erfüllt.

A3. Setzt man $\vec{o}$ als das n -Tupel, dessen sämtliche Komponenten Null sind, so gilt für jedes n -Tupel $\vec{x}$:

$$\vec{x} + \vec{o} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \vec{x}.$$

A4. Es sei $\vec{x}$ ein beliebiges n -Tupel mit den Komponenten $x_1, x_2, \dots, x_n$. Mit dem n -Tupel $-\vec{x}$ mit den Komponenten $-x_1, -x_2, \dots, -x_n$ gilt $\vec{x} + (-\vec{x}) = \vec{o}$.

S3. Für beliebige einander entsprechende Komponenten x_i, y_i ($i = 1 \dots n$) zweier n -Tupel $\vec{x}, \vec{y}$ und beliebige $\lambda \in \mathbb{R}$ gilt nach dem Distributivgesetz für reelle Zahlen $\lambda \cdot (x_i + y_i) = \lambda \cdot x_i + \lambda \cdot y_i$. Die n -Tupel $\lambda \cdot (\vec{x} + \vec{y})$ und $\lambda \cdot \vec{x} + \lambda \cdot \vec{y}$ stimmen somit in allen Komponenten überein, es gilt daher $\lambda \cdot (\vec{x} + \vec{y}) = \lambda \cdot \vec{x} + \lambda \cdot \vec{y}$. $\square$

3.2.3 Aufgaben zu Abschnitt 3.2

1. Beweisen Sie die Aussagen A1, S1, S2, und S4 des Satzes 3.6 auf S. 104.
2. Der Satz 3.6 sagt u. a. aus, dass in $\mathbb{R}^n$ ein „Null- n -Tupel“ $\vec{o}$ und für jedes n -Tupel $\vec{x}$ ein „Gegen- n -Tupel“ $-\vec{x}$ existiert. Er beinhaltet aber keine Eindeutigkeitsaussagen. Begründen Sie die Eindeutigkeit des „Null- n -Tupels“ sowie des „Gegen- n -Tupels“ zu jedem n -Tupel.
3. Es sei $L \subset \mathbb{R}^n$ die Lösungsmenge eines homogenen linearen Gleichungssystems in n Variablen (siehe Abschnitt 1.3.4). Begründen Sie, dass die Addition von n -Tupeln und die Multiplikation von n -Tupeln mit reellen Zahlen innerhalb von L uneingeschränkt ausführbar sind, für $\vec{x}, \vec{y} \in L, \lambda \in \mathbb{R}$ also stets $\vec{x} + \vec{y} \in L$ sowie $\lambda \cdot \vec{x} \in L$ gilt. Begründen Sie weiterhin, dass innerhalb von L alle Aussagen des Satzes 3.6 gelten.

3.3 Zusammenhänge zwischen Pfeilklassen und n -Tupeln – Vektoren

Die Sätze 3.3 und 3.5 sowie 3.6 haben gezeigt, dass für Pfeilklassen und n -Tupel dieselben grundlegenden Rechenregeln gelten, obwohl diese Objekte aus völlig unterschiedlichen Bereichen der Arithmetik und der Geometrie stammen und deshalb auch die Beweise der Rechenregeln auf grundverschiedenen Wegen geführt wurden. Im Folgenden werden nun weitere Zusammenhänge zwischen Pfeilklassen und n -Tupeln hergestellt, welche die strukturelle Gleichheit dieser Kategorien von Objekten noch untermauern.

3.3.1 Pfeilklassen im Koordinatensystem

Innerhalb eines kartesischen Koordinatensystems⁶ lassen sich den Anfangs- und Endpunkten von Pfeilen Koordinaten zuordnen. Dabei zeigt sich, dass Pfeile, die dieselbe Pfeilklassse repräsentieren (also parallelgleich sind), gleiche Koordinatendifferenzen der End- und Anfangspunkte haben, siehe Abb. 3.21. Für zwei beliebige parallelgleiche Pfeile $\overrightarrow{AB}$ und $\overrightarrow{CD}$ der Ebene oder des Raumes gilt

$$x_B - x_A = x_D - x_C, \quad y_B - y_A = y_D - y_C \quad (\text{und im Raum } z_B - z_A = z_D - z_C).$$

Ist insbesondere $\overrightarrow{OP}$ (wobei O der Koordinatenursprung ist) ein Pfeil, der derselben Pfeilklassse wie $\overrightarrow{AB}$ und $\overrightarrow{CD}$ angehört, so ist bei Pfeilen in der Ebene

$$x_P = x_P - 0 = x_B - x_A = x_D - x_C \quad \text{und} \quad y_P = y_P - 0 = y_B - y_A = y_D - y_C$$

sowie im Raum

$$x_P = x_B - x_A = x_D - x_C, \quad y_P = y_B - y_A = y_D - y_C, \quad z_P = z_B - z_A = z_D - z_C.$$

Die *Differenzen der End- und Anfangskoordinaten* sind somit ein gemeinsames Merkmal von allen Pfeilen einer Pfeilklassse, also *repräsentantenunabhängig*. Daraus ergibt sich eine umkehrbar eindeutige Zuordnung zwischen Pfeilklassen und n -Tupeln (Paaren bzw. Tripeln) reeller Zahlen.

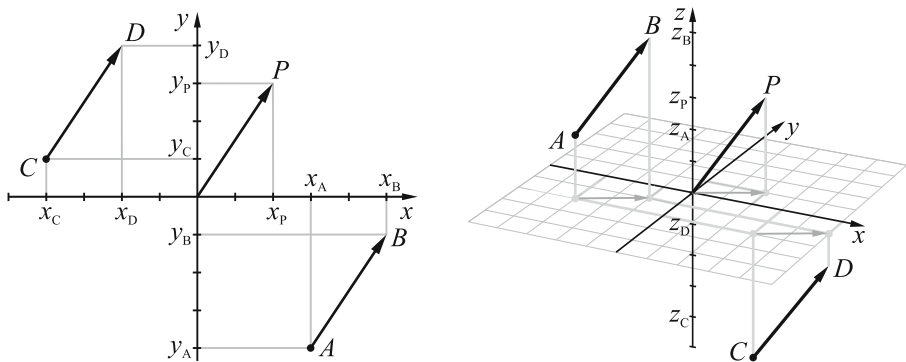


Abb. 3.21: Differenzen der Koordinaten der End- und Anfangspunkte von Pfeilen

⁶Unter einem kartesischen Koordinatensystem versteht man ein Koordinatensystem mit zueinander senkrechten Achsen mit gleicher Skalierung, wobei der Abstand zwischen zwei aufeinander folgenden ganzzahligen Achsenwerten der Länge einer Einheitsstrecke entspricht.

Gegeben sei ein kartesisches Koordinatensystem in der Ebene oder im Raum. Jeder Pfeilklass $\vec{p}$ lässt sich eindeutig ein Paar bzw. Tripel reeller Zahlen auf folgende Weise zuordnen:

Ist $\overrightarrow{AB}$ ein beliebiger Repräsentant von $\vec{p}$, so wird $\vec{p}$ das Paar bzw. Tripel der Koordinatendifferenzen der Punkte A und B zugeordnet:

$$\vec{p} \rightarrow \begin{pmatrix} x_B - x_A \\ y_B - y_A \end{pmatrix} \quad \text{bzw.} \quad \vec{p} \rightarrow \begin{pmatrix} x_B - x_A \\ y_B - y_A \\ z_B - z_A \end{pmatrix}.$$

Für spezielle Repräsentanten $\overrightarrow{OP}$ von $\vec{p}$, deren Anfangspunkt der Koordinatenursprung ist, vereinfacht sich diese Zuordnung:

$$\vec{p} \rightarrow \begin{pmatrix} x_P \\ y_P \end{pmatrix} \quad \text{bzw.} \quad \vec{p} \rightarrow \begin{pmatrix} x_P \\ y_P \\ z_P \end{pmatrix}.$$

Umgekehrt lässt sich jedem Zahlenpaar $\begin{pmatrix} x \\ y \end{pmatrix}$ bzw. -tripel $\begin{pmatrix} x \\ y \\ z \end{pmatrix}$ (mit $x, y, z \in \mathbb{R}$) bezüglich eines gegebenen Koordinatensystems eindeutig eine Pfeilklass der Ebene bzw. des Raumes zuordnen, die durch den Pfeil $\overrightarrow{OP}$ mit dem Anfangspunkt im Koordinatenursprung und dem Endpunkt $P(x; y)$ bzw. $P(x; y; z)$ repräsentiert wird.

Die Existenz eindeutiger Zuordnungen: Pfeilklassen $\rightarrow$ Zahlenpaare/-tripel sowie Zahlenpaare/-tripel $\rightarrow$ Pfeilklassen lässt sich wie folgt zusammenfassen:

Satz 3.7

Bezüglich eines gegebenen kartesischen Koordinatensystems existiert eine umkehrbar eindeutige Abbildung von der Menge aller Pfeilklassen der Ebene bzw. des Raumes auf $\mathbb{R}^2$ bzw. $\mathbb{R}^3$.

Eine Abbildung mit den in dem Satz 3.7 genannten Eigenschaften wird als *bijektive Abbildung* bzw. *Bijektion* bezeichnet. Genaue Definitionen dieser Begriffe werden in dem Kapitel 7 gegeben.

In den vorherigen Ausführungen und in dem Satz 3.7 wurde stets der Bezug auf ein (kartesisches) Koordinatensystem erwähnt. Es können verschiedene Koordinatensysteme der Ebene bzw. des Raumes betrachtet werden, wobei nicht einmal in allen Koordinatensystemen die Achsen senkrecht zueinander sind. Aber auch Koordinatensysteme mit senkrechten Achsen können sich durch die Skalierung der Achsen und ihre Lage in der Ebene bzw. im Raum unterschei-

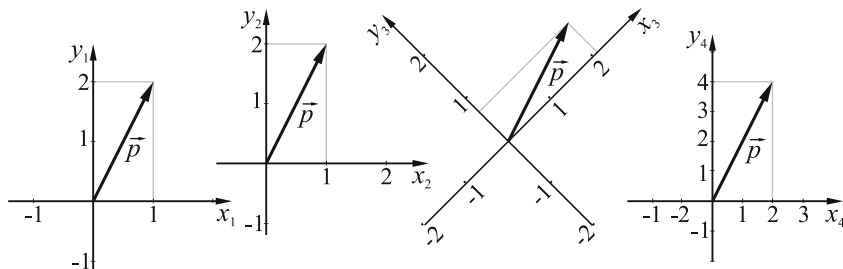


Abb. 3.22: Repräsentanten einer Pfeilklass in verschiedenen Koordinatensystemen

den. Welches n -Tupel einer Pfeilklassse zugeordnet wird, hängt von der Wahl des Koordinatensystems ab. Lediglich bezüglich zweier Koordinatensysteme mit jeweils zueinander parallelen Achsen und gleicher Skalierung wird einer Pfeilklassse nach dem beschriebenen Verfahren dasselbe n -Tupel zugeordnet. In Abb. 3.22 sind verschiedene Repräsentanten ein und derselben Pfeilklassse $\vec{p}$ dargestellt. Bezüglich der Koordinatensysteme mit den Achsen x_1, y_1 und x_2, y_2 wird $\vec{p}$ jeweils das Zahlenpaar $\begin{pmatrix} 1 \\ 2 \end{pmatrix}$ zugeordnet, bezüglich des Koordinatensystems mit den Achsen x_3, y_3 das Paar $\begin{pmatrix} \frac{3}{2}\sqrt{2} \\ \frac{1}{2}\sqrt{2} \end{pmatrix} \approx \begin{pmatrix} 2,12 \\ 0,71 \end{pmatrix}$ und bezüglich des Koordinatensystems mit den Achsen x_4, y_4 das Paar $\begin{pmatrix} 2 \\ 4 \end{pmatrix}$.

Der enge Zusammenhang zwischen Pfeilklassen und Zahlentupeln ist in beiden Richtungen von Nutzen. Die Zuordnung: Pfeilklassen $\rightarrow$ Zahlenpaare/-tripel ermöglicht es, geometrische Aufgaben auf arithmetischem bzw. algebraischem Wege zu lösen. Umgekehrt schafft die Interpretation von Zahlenpaaren oder -tripeln als Pfeilklassen mitunter interessante Veranschaulichungsmöglichkeiten für Sachverhalte, die zunächst keinerlei geometrischen Bezug zu haben scheinen.

Beispiel 3.2

Farbwürfel

In dem Abschnitt 3.2.1 wurde auf die Beschreibung von Farben durch Zahlentripel $\begin{pmatrix} r \\ g \\ b \end{pmatrix}$ eingegangen (siehe S. 101). Dabei ist die Menge aller RGB-Tripel eine Teilmenge von $\mathbb{R}^3$ mit $0 \leq r, g, b \leq 1$. Den RGB-Tripeln entsprechen Pfeilklassen. Betrachtet man Repräsentanten dieser Pfeilklassen mit Anfangspunkten im Koordinatenursprung, so liegen deren Endpunkte innerhalb eines Würfels mit den Eckpunktkoordinaten $(0; 0; 0)$, $(1; 0; 0)$, $(1; 1; 0)$, $(0; 1; 0)$, $(0; 0; 1)$, $(1; 0; 1)$, $(1; 1; 1)$ und $(0; 1; 1)$, siehe Abb. 3.23. Die Eckpunkte des als Farb- bzw. RGB-Würfel bezeichneten Würfels kennzeichnen somit die Grundfarben des RGB-Modells, deren Komplementärfarben sowie Schwarz und Weiß. Komplementären Farben entsprechen jeweils gegenüberliegende Punkte des Würfels. Auf den Seitenflächen sind die Verläufe zwischen verschiedenen Farben erkennbar. ■

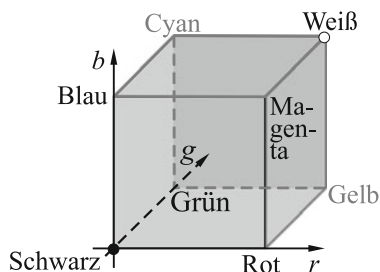


Abb. 3.23: RGB-Würfel

Auf der Internetseite zu diesem Buch befinden sich eine farbige Abbildung des RGB-Würfels sowie ein Video, das den RGB-Würfel aus verschiedenen Richtungen zeigt.

3.3.2 Die Isomorphie zwischen der Menge der Pfeilklassen der Ebene bzw. des Raumes und $\mathbb{R}^2$ bzw. $\mathbb{R}^3$

Pfeilklassen können nicht nur n -Tupel zugeordnet werden und umgekehrt, sondern auch die Pfeilklassenaddition und die Multiplikation mit reellen Zahlen lassen sich auf die Addition von n -Tupeln bzw. deren Multiplikation mit Zahlen zurückführen.

Satz 3.8

Es seien $\vec{u}$ und $\vec{v}$ Pfeilklassen, denen bezüglich eines Koordinatensystems Paare bzw. Tripel reeller Zahlen zugeordnet sind (wie auf S. 107 beschrieben):

$$\vec{u} \rightarrow \begin{pmatrix} x_u \\ y_u \end{pmatrix}, \quad \vec{v} \rightarrow \begin{pmatrix} x_v \\ y_v \end{pmatrix} \quad \text{bzw.} \quad \vec{u} \rightarrow \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}, \quad \vec{v} \rightarrow \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}.$$

Dann ergibt sich bei der „geometrischen“ Addition der Pfeilklassen $\vec{u}$ und $\vec{v}$ (nach Definition 3.3) eine Pfeilklass $\vec{u} + \vec{v}$, der das Paar/Tripel entspricht, das sich durch „arithmetische“ Addition (nach Definition 3.6) der $\vec{u}$ und $\vec{v}$ zugeordneten Paare/Tripel ergibt:

$$\vec{u} + \vec{v} \rightarrow \begin{pmatrix} x_u \\ y_u \end{pmatrix} + \begin{pmatrix} x_v \\ y_v \end{pmatrix} \quad \text{bzw.} \quad \vec{u} + \vec{v} \rightarrow \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix} + \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}.$$

Ist weiterhin λ eine reelle Zahl, so führt die „geometrische“ (durch die Definition 3.5 beschriebene) Multiplikation von $\vec{u}$ mit λ zu einer Pfeilklass $\lambda \cdot \vec{u}$, der das Paar/Tripel entspricht, welches sich durch „arithmetische“ Multiplikation (nach Definition 3.6) des $\vec{u}$ zugeordneten Paares/Tripels mit λ ergibt:

$$\lambda \cdot \vec{u} \rightarrow \lambda \cdot \begin{pmatrix} x_u \\ y_u \end{pmatrix} \quad \text{bzw.} \quad \lambda \cdot \vec{u} \rightarrow \lambda \cdot \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}.$$

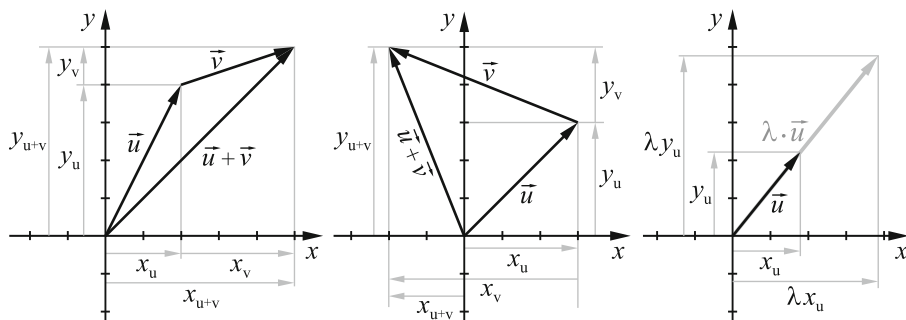


Abb. 3.24: Addition und skalare Multiplikation von Pfeilklassen sowie von n -Tupeln

Auf einen Beweis des Satzes 3.8 wird verzichtet; zur Veranschaulichung seiner Aussage sei auf Abb. 3.24 verwiesen, aus der sich auch Beweisideen für die beiden Teile des Satzes entnehmen lassen.

Die Sätze 3.7 und 3.8 beinhalten zusammen, dass (bezüglich eines Koordinatensystems) die Mengen aller Pfeilklassen der Ebene bzw. des Raumes und $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ bijektiv aufeinander abbildbar und dass die Addition sowie die Multiplikation mit Skalaren übertragbar sind. Dieser enge Zusammenhang zwischen

zunächst völlig unterschiedlichen Strukturen aus der Geometrie bzw. der Arithmetik wird als *Isomorphie* (Strukturgleichheit), die in dem Kasten auf S. 107 festgelegte Zuordnung als *Isomorphismus* bezeichnet. (Näher wird auf diese Begriffe in dem Kapitel 7 eingegangen.)

3.3.3 Vektoren

Die Strukturgleichheit zwischen Pfeilklassen und n -Tupeln (für $n = 2, 3$) ermöglicht es, diese Objekte „gleichzusetzen“. Wir schreiben daher ab sofort für Pfeilklassen $\vec{u}$ und zugehörige n -Tupel:

$$\vec{u} = \begin{pmatrix} x \\ y \end{pmatrix} \quad \text{bzw.} \quad \vec{u} = \begin{pmatrix} x \\ y \\ z \end{pmatrix}.$$

Wir unterscheiden künftig auch sprachlich nicht mehr zwischen Pfeilklassen und n -Tupeln, sondern bezeichnen diese nun als *Vektoren*.

Mit der Zusammenführung von geometrischen und arithmetischen Objekten zeichnet sich der Vektorbegriff bereits durch einen hohen Verallgemeinerungsgrad aus und eröffnet die Möglichkeit, rechnerische Methoden auf die Geometrie zu übertragen und umgekehrt. Eine weitere Verallgemeinerung erfolgt dann in dem Kapitel 5 mit der Behandlung von Vektorräumen. Erst auf dieser Grundlage ist dann eine allgemeine Definition des Begriffs Vektor möglich. Mit den in den Sätzen 3.3 und 3.5 sowie 3.6 enthaltenen Rechenregeln für Pfeilklassen und n -Tupel sind jedoch bereits die wichtigsten strukturellen Eigenschaften von Vektoren gegeben, die auch Vektorräume maßgeblich bestimmen werden. Im Folgenden werden diese Sätze zusammengefasst für Vektoren formuliert:

Satz 3.9

A1. Für beliebige Vektoren $\vec{u}$ und $\vec{v}$ gilt $\vec{u} + \vec{v} = \vec{v} + \vec{u}$

(Kommutativität der Vektoraddition).

A2. Für beliebige Vektoren $\vec{u}, \vec{v}, \vec{w}$ gilt $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$

(Assoziativität der Vektoraddition).

A3. Es existiert ein Vektor $\vec{o}$, so dass für jeden Vektor $\vec{u}$ gilt: $\vec{u} + \vec{o} = \vec{u}$

(Existenz eines Nullvektors).

A4. Zu jedem Vektor $\vec{u}$ existiert ein Vektor $-\vec{u}$ mit $\vec{u} + (-\vec{u}) = \vec{o}$

(Existenz eines Gegenvektors zu jedem Vektor).

S1. Für beliebige Vektoren $\vec{u}$ gilt $1 \cdot \vec{u} = \vec{u}$.

S2. Für beliebige Vektoren $\vec{u}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda \cdot \mu) \cdot \vec{u} = \lambda \cdot (\mu \cdot \vec{u})$

(Assoziativität der Multiplikation von Vektoren mit reellen Zahlen).

S3. Für beliebige Vektoren $\vec{u}, \vec{v}$ und beliebige $\lambda \in \mathbb{R}$ gilt $\lambda \cdot (\vec{u} + \vec{v}) = \lambda \cdot \vec{u} + \lambda \cdot \vec{v}$

(1. Distributivgesetz).

S4. Für beliebige Vektoren $\vec{u}$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda + \mu) \cdot \vec{u} = \lambda \cdot \vec{u} + \mu \cdot \vec{u}$

(2. Distributivgesetz).

Bewiesen wurde der Satz 3.9 bereits in den Abschnitten 3.1.3 (für Pfeilklassen) und 3.2.2 (für n -Tupel).

Mithilfe des Satzes 3.9 lassen sich weitere Eigenschaften von Vektoren herleiten. Dazu ist es nicht mehr notwendig, diese jeweils für Pfeilklassen und n -Tupel zu beweisen, sondern es genügt, auf die Aussagen des Satzes 3.9 zurückzugreifen.

Satz 3.10

1. Für alle Vektoren $\vec{u}$ ist $0 \cdot \vec{u} = \vec{o}$.
2. Der Nullvektor $\vec{o}$ ist eindeutig bestimmt.
3. Für alle $\lambda \in \mathbb{R}$ ist $\lambda \cdot \vec{o} = \vec{o}$.
4. Für jeden Vektor $\vec{u}$ ist der Gegenvektor $-\vec{u}$ eindeutig bestimmt.
5. Für alle Vektoren $\vec{u}$ und alle $\lambda \in \mathbb{R}$ ist $(-\lambda) \cdot \vec{u} = -(\lambda \cdot \vec{u})$.
6. Für alle Vektoren $\vec{u}$ und alle $\lambda \in \mathbb{R}$ ist $\lambda \cdot (-\vec{u}) = -(\lambda \cdot \vec{u})$.
7. Für alle Vektoren $\vec{u}$ und alle $\lambda \in \mathbb{R}$ gilt: Aus $\lambda \cdot \vec{u} = \vec{o}$ folgt $\lambda = 0$ oder $\vec{u} = \vec{o}$.

Bemerkung: Die Aussage 5 des Satzes beinhaltet den Spezialfall $(-1) \cdot \vec{u} = -\vec{u}$. Dies bedeutet, dass sich der Gegenvektor $-\vec{u}$ eines beliebigen Vektors $\vec{u}$ durch Multiplikation von $\vec{u}$ mit -1 ergibt.

Beweis: Es werden hier nicht alle Aussagen des Satzes 3.10 bewiesen. Da teilweise recht ähnliche Vorgehensweisen anzuwenden sind, kann die Leserin/der Leser die restlichen Aussagen selbst beweisen, siehe die Aufgabe 5 auf S. 112.

1. Wir setzen $\vec{v} := 0 \cdot \vec{u}$ und weisen nach, dass $\vec{v} = \vec{o}$ sein muss.

Wegen $0+0=0$ gilt $\vec{v} = 0 \cdot \vec{u} = (0+0) \cdot \vec{u}$. Nach dem 2. Distributivgesetz ist

$$(0+0) \cdot \vec{u} = 0 \cdot \vec{u} + 0 \cdot \vec{u} = \vec{v} + \vec{v},$$

also $\vec{v} = \vec{v} + \vec{v}$. Durch Addition des Gegenvektors $-\vec{v}$ auf beiden Seiten ergibt sich daraus $\vec{v} + (-\vec{v}) = \vec{v} + \vec{v} + (-\vec{v})$ und somit nach Satz 3.9 A4:

$$\vec{o} = \vec{v} + (\vec{v} + (-\vec{v})) = \vec{v} + \vec{o} = \vec{v} = 0 \cdot \vec{u}.$$

2. Wir nehmen an, es seien $\vec{o}$ und $\vec{o}^*$ zwei Nullvektoren. Dann gilt nach Satz 3.9 A3 sowohl $\vec{o} + \vec{o}^* = \vec{o}$ als auch $\vec{o}^* + \vec{o} = \vec{o}^*$. Somit ist also $\vec{o}^* = \vec{o}^* + \vec{o} = \vec{o} + \vec{o}^* = \vec{o}$.
3. Nach Satz 3.9 A3 ist $\vec{u} + \vec{o} = \vec{u}$ für alle Vektoren $\vec{u}$, woraus $\lambda \cdot (\vec{u} + \vec{o}) = \lambda \cdot \vec{u}$ folgt. Nach dem 1. Distributivgesetz ergibt sich daraus $\lambda \cdot \vec{u} + \lambda \cdot \vec{o} = \lambda \cdot \vec{u}$. Setzt man $\vec{v} = \lambda \cdot \vec{u}$, so ist also einerseits $\vec{v} + \lambda \cdot \vec{o} = \vec{v}$ und andererseits (nach Satz 3.9 A3) $\vec{v} + \vec{o} = \vec{v}$. Sowohl $\vec{o}$ als auch $\lambda \cdot \vec{o}$ sind somit Nullvektoren. Da der Nullvektor aber eindeutig bestimmt ist, gilt $\lambda \cdot \vec{o} = \vec{o}$.
5. Nach Satz 3.10, 1 ist $\vec{o} = 0 \cdot \vec{u} = (\lambda + (-\lambda)) \cdot \vec{u} = \lambda \cdot \vec{u} + (-\lambda) \cdot \vec{u}$. Somit ist $(-\lambda) \cdot \vec{u}$ der Gegenvektor zu $\lambda \cdot \vec{u}$. Wegen des Satzes 3.10, 4 ist dieser eindeutig bestimmt. Daher muss $-(\lambda \cdot \vec{u}) = (-\lambda) \cdot \vec{u}$ gelten.
7. Falls $\lambda = 0$ ist, so ist die Behauptung bereits erfüllt. Wir können also $\lambda \neq 0$ voraussetzen und zeigen, dass dann $\vec{u} = \vec{o}$ sein muss. Für $\lambda \neq 0$ existiert $\frac{1}{\lambda} \in \mathbb{R}$. Nach den Aussagen S1 und S2 des Satzes 3.9 sowie der Voraussetzung $\lambda \cdot \vec{u} = \vec{o}$ gilt: $\vec{u} = 1 \cdot \vec{u} = (\frac{1}{\lambda} \lambda) \cdot \vec{u} = \frac{1}{\lambda} \cdot (\lambda \cdot \vec{u}) = \frac{1}{\lambda} \cdot \vec{o} = \vec{o}$. □

3.3.4 Aufgaben zu Abschnitt 3.3

- Ein Vektor $\vec{v}$ wird durch einen Pfeil $\overrightarrow{AB}$ repräsentiert. Geben Sie $\vec{v}$ als Zahlentripel an.
 - $A(-8; 13; 5)$, $B(-5; 6; -11)$
 - $A(5; 7; 1)$, $B(3; 9; -2)$
- Gegeben ist eine Verschiebung $\vec{v}$ des Raumes durch einen Verschiebungspfeil $\overrightarrow{PP'}$ mit $P(2; 1; 3)$ und $P'(5; 3; -1)$.
 - Geben Sie den Verschiebungsvektor $\vec{v}$ als Zahlentripel an.
 - Geben Sie die Koordinaten der Bildpunkte der Punkte $A(3; -2; 4)$ und $B(3; 5; 2; 5; -5)$ bei der Verschiebung $\vec{v}$ an.
- Gegeben sind die Vektoren $\vec{u} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} -3 \\ 2 \end{pmatrix}$. Bestimmen Sie graphisch und rechnerisch den Vektor $\vec{w}$. Vergleichen Sie jeweils Ihre Ergebnisse.
 - $\vec{w} = 3 \cdot \vec{u} + 2 \cdot \vec{v}$
 - $\vec{w} = -4 \cdot \vec{u} + \frac{1}{2} \cdot \vec{v}$
- Durch $\vec{v}_1 = \begin{pmatrix} 5 \\ 3 \\ 2 \end{pmatrix}$ und $\vec{v}_2 = \begin{pmatrix} -4 \\ 2 \\ 4 \end{pmatrix}$ werden zwei Verschiebungen des Raumes beschrieben (siehe Abb. 3.25).
 - Der Punkt $P(-3; -3; 3)$ wird zunächst um $\vec{v}_1$ und dann um $\vec{v}_2$ verschoben. Geben Sie die Koordinaten der entstehenden Bildpunkte P' und P'' an.
 - Geben Sie den Verschiebungsvektor $\vec{v}$ an, der die Nacheinanderausführung der Verschiebungen $\vec{v}_1$ und $\vec{v}_2$ beschreibt.

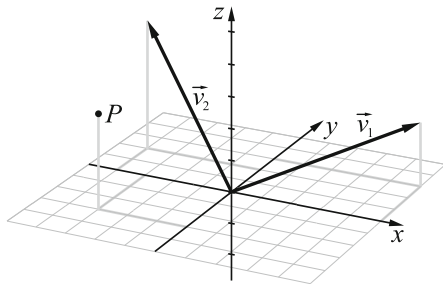


Abb. 3.25: Verschiebungen im Raum (Skizze zu Aufgabe 4)

- Beweisen Sie die Aussagen 4 und 6 des Satzes 3.10 auf S. 111.
- Weisen Sie nach: Die Vektorgleichung $\vec{u} + \vec{x} = \vec{v}$ ($\vec{u}, \vec{v}$ gegeben, $\vec{x}$ gesucht) besitzt eine eindeutige Lösung.
- Welcher Vektor $\vec{x}$ erfüllt die folgenden Gleichungen, wenn die Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ gegeben sind?
 - $2\vec{u} - \vec{v} + \vec{x} = \vec{w} + \vec{u} - \vec{v}$
 - $4(3\vec{u} - 2\vec{v} + 3\vec{x}) - 3(-2\vec{v} + 3\vec{w} - \vec{x}) = \vec{o}$

3.4 Linearkombinationen von Vektoren

3.4.1 Linearkombinationen

Sind zwei Vektoren $\vec{u}$ und $\vec{v}$ gegeben, so lassen sich weitere Vektoren $\vec{x}$ aus diesen durch Multiplikation mit reellen Zahlen und anschließende Addition erzeugen:

$$\vec{x} = \lambda \cdot \vec{u} + \mu \cdot \vec{v}.$$

Dieses als Linearkombination bezeichnete Prinzip der Bildung von Vektoren aus gegebenen Vektoren lässt sich für beliebig viele Vektoren verallgemeinern.

Definition 3.7

Als *Linearkombination* der Vektoren $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ bezeichnet man einen Vektor $\vec{x}$ mit

$$\vec{x} = \lambda_1 \cdot \vec{u}_1 + \lambda_2 \cdot \vec{u}_2 + \dots + \lambda_k \cdot \vec{u}_k = \sum_{i=1}^k \lambda_i \cdot \vec{u}_i \quad (\text{mit } \lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{R}). \quad \blacklozenge$$

Bemerkung: Da $\lambda_1, \lambda_2, \dots, \lambda_k$ beliebige reelle Zahlen sein können, gibt es zu jeder Menge von Vektoren, die nicht leer ist und nicht nur aus dem Nullvektor besteht, unendlich viele Linearkombinationen.

Beispiel 3.3

Der Vektor $\vec{x} = \begin{pmatrix} 5 \\ 2 \end{pmatrix}$ soll als Linearkombination der Vektoren $\vec{u} = \begin{pmatrix} -1 \\ 3 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 4 \\ -2 \end{pmatrix}$ dargestellt werden. Dazu sind Koeffizienten λ und μ so zu bestimmen, dass $\vec{x} = \lambda \vec{u} + \mu \vec{v}$ gilt, also

$$\begin{pmatrix} 5 \\ 2 \end{pmatrix} = \lambda \cdot \begin{pmatrix} -1 \\ 3 \end{pmatrix} + \mu \cdot \begin{pmatrix} 4 \\ -2 \end{pmatrix}.$$

Um die Koeffizienten λ und μ zu bestimmen, ist das lineare Gleichungssystem

$$\begin{aligned} -\lambda + 4 \cdot \mu &= 5 \\ 3 \cdot \lambda - 2 \cdot \mu &= 2 \end{aligned}$$

zu lösen, es ergibt sich dabei $\lambda = \frac{9}{5}$ und $\mu = \frac{17}{10}$. Der Vektor $\vec{x}$ lässt sich also in folgender Weise als Linearkombination der Vektoren $\vec{u}$ und $\vec{v}$ darstellen:

$$\vec{x} = \frac{9}{5} \cdot \vec{u} + \frac{17}{10} \cdot \vec{v}. \quad \blacksquare$$

Geometrisch lässt sich die Darstellung eines Vektors $\vec{x}$ als Linearkombination zweier Vektoren $\vec{u}$ und $\vec{v}$ folgendermaßen deuten: Die Vektoren $\vec{u}$ und $\vec{v}$ geben

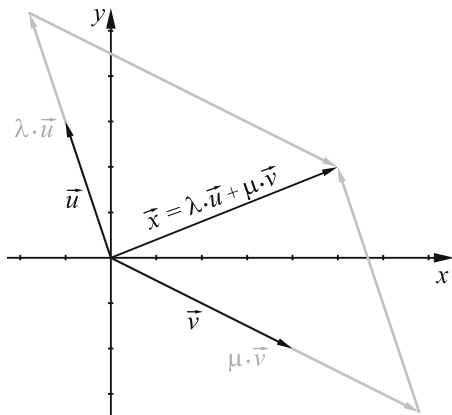


Abb. 3.26: Darstellung eines Vektors $\vec{x}$ als Linearkombination zweier Vektoren $\vec{u}$ und $\vec{v}$

die Richtungen der Seiten eines Parallelogramms an. Gesucht sind Koeffizienten λ, μ , mit denen $\vec{u}$ und $\vec{v}$ multipliziert werden müssen, damit der Vektor $\vec{x}$ eine Diagonale des von $\lambda\vec{u}$ und $\mu\vec{v}$ aufgespannten Parallelogramms beschreibt, siehe Abb. 3.26. In dem Beispiel 3.3 wurde diese Aufgabe rechnerisch gelöst.

Nicht immer lässt sich ein Vektor als eine Linearkombination zweier gegebener Vektoren darstellen und nicht in allen Fällen, in denen dies möglich ist, sind die Koeffizienten eindeutig bestimmt.

Beispiel 3.4

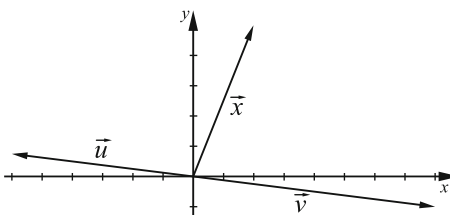
Der Vektor $\vec{x} = \begin{pmatrix} 5 \\ 2 \end{pmatrix}$ soll als Linearkombination der Vektoren $\vec{u} = \begin{pmatrix} -6 \\ 3 \\ 4 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 8 \\ -1 \end{pmatrix}$ dargestellt werden.

Um Koeffizienten λ und μ so zu finden, dass $\vec{x} = \lambda\vec{u} + \mu\vec{v}$ gilt, wird wie in dem Beispiel 3.3 ein lineares Gleichungssystem aufgestellt:

$$\begin{aligned} -6 \cdot \lambda + 8 \cdot \mu &= 5 \\ \frac{3}{4} \cdot \lambda - \mu &= 2 \end{aligned}$$

Nach einem Umformungsschritt ergibt sich daraus $0 = 21$, das LGS ist nicht lösbar.

Der Vektor $\vec{x}$ kann daher nicht als Linearkombination von $\vec{u}$ und $\vec{v}$ dargestellt werden. Geometrisch lässt sich dies dadurch veranschaulichen, dass der Vektor $\vec{v}$ durch Multiplikation des Vektors $\vec{u}$ mit $-\frac{4}{3}$ entsteht. Beide Vektoren werden somit durch parallele Pfeile dargestellt. Da der Vektor $\vec{x}$ nicht durch Pfeile repräsentiert wird, die dazu parallel sind, lässt er sich nicht in der Form $\lambda\vec{u} + \mu\vec{v}$ darstellen (siehe Abb. 3.27).



Der Vektor $\vec{x} = \begin{pmatrix} 4 \\ 3 \end{pmatrix}$ soll als Linearkombination der Vektoren $\vec{u} = \begin{pmatrix} -2 \\ 3 \\ 4 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 2 \\ 6 \end{pmatrix}$ und $\vec{w} = \begin{pmatrix} -1 \\ 6 \end{pmatrix}$ dargestellt werden.

Um Koeffizienten für eine Linearkombination $\vec{x} = \lambda\vec{u} + \mu\vec{v} + \nu\vec{w}$ zu bestimmen, ist folgendes LGS zu lösen:

$$\begin{aligned} -2 \cdot \lambda + 2 \cdot \mu - \nu &= 4 \\ \frac{3}{4} \cdot \lambda + 6 \cdot \mu + 6 \cdot \nu &= 3. \end{aligned}$$

Umformung dieses LGS führt zu

$$\begin{aligned} -2 \cdot \lambda + 2 \cdot \mu - \nu &= 4 \\ 6 \cdot \mu + 5 \cdot \nu &= 4. \end{aligned}$$

Es existiert somit eine einparametrische Lösungsmenge. Setzt man $t = \nu$, so ist $\mu = \frac{2}{3} - \frac{5}{6}t$ und $\lambda = -\frac{4}{3} - \frac{4}{3}t$. Es gilt also $\vec{x} = \left(-\frac{4}{3} - \frac{4}{3}t\right) \cdot \vec{u} + \left(\frac{2}{3} - \frac{5}{6}t\right) \cdot \vec{v} + t \cdot \vec{w}$ für beliebige $t \in \mathbb{R}$.

Prinzipiell lässt sich ein Vektor $\vec{x} \in \mathbb{R}^2$ niemals eindeutig als Linearkombination von mehr als zwei Vektoren darstellen, da hierbei stets LGS mit zwei Gleichungen und mehr als zwei Variablen entstehen. Diese besitzen stets unendlich viele Lösungen oder sind nicht lösbar (vgl. Abschnitt 1.2). ■

Abb. 3.27: Beispiel für eine nicht mögliche Linearkombination

Beispiel 3.5Linearkombination von Vektoren in $\mathbb{R}^3$

Der Vektor $\vec{x} = \begin{pmatrix} 4 \\ 8 \\ 3 \end{pmatrix}$ wird als Linearkombination dreier Vektoren $\vec{u} = \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 1 \\ 4 \\ -1 \end{pmatrix}$ und $\vec{w} = \begin{pmatrix} 5 \\ -1 \\ -1 \end{pmatrix}$ in der Form $\vec{x} = \lambda \vec{u} + \mu \vec{v} + \nu \vec{w}$ dargestellt. Dazu ist das folgende lineare Gleichungssystem zu lösen:

$$\begin{aligned} 1 \cdot \mu + 5 \cdot \nu &= 4 \\ \lambda + 4 \cdot \mu - 1 \cdot \nu &= 8 \\ 2 \cdot \lambda - 1 \cdot \mu - 1 \cdot \nu &= 3. \end{aligned}$$

Mithilfe des Gauß-Algorithmus (siehe Abschnitt 1.2.2) erhält man die (eindeutige) Lösung $\lambda = \frac{5}{2}$, $\mu = \frac{3}{2}$, $\nu = \frac{1}{2}$. Es gilt also

$$\vec{x} = \frac{5}{2} \cdot \vec{u} + \frac{3}{2} \cdot \vec{v} + \frac{1}{2} \cdot \vec{w}.$$

Geometrisch lässt sich die Darstellung eines Vektors $\vec{x} \in \mathbb{R}^3$ als Linearkombination dreier Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ folgendermaßen veranschaulichen: Die Vektoren $\vec{u}$, $\vec{v}$, $\vec{w}$ geben die Richtungen der Seiten eines Parallelepipeds (Spats) vor. Darunter versteht man einen Körper, der von sechs paarweise kongruenten und in parallelen Ebenen liegenden Parallelogrammen begrenzt wird, siehe Abb. 3.28. Werden $\vec{u}$, $\vec{v}$, $\vec{w}$ mit den in dem Beispiel 3.5 ermittelten Koeffizienten λ , μ , ν multipliziert, so spannen die Vektoren $\lambda \vec{u}$, $\mu \vec{v}$ und $\nu \vec{w}$ ein Parallelepiped auf, in dem der Vektor $\vec{x}$ eine Diagonale beschreibt.⁷

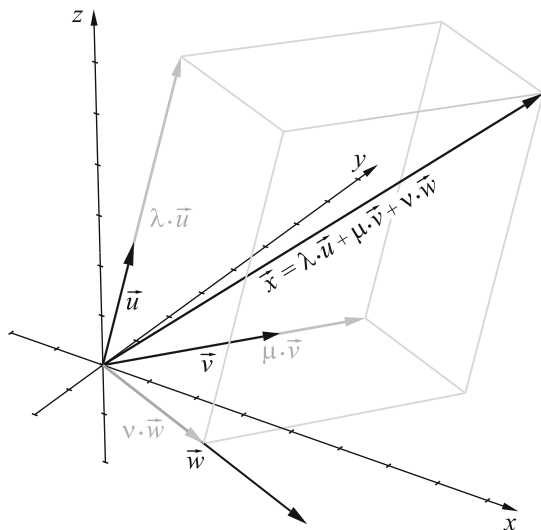


Abb. 3.28: Darstellung eines Vektors $\vec{x} \in \mathbb{R}^3$ als Linearkombination dreier Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$

⁷Genauer müsste davon gesprochen werden, dass ein den Vektor $\vec{x}$ repräsentierender Pfeil eine Diagonale eines Parallelepipeds bildet, welches von Pfeilen aufgespannt wird, die Repräsentanten der Vektoren $\lambda \vec{u}$, $\mu \vec{v}$ und $\nu \vec{w}$ sind. Vektoren (in der geometrischen Interpretation als Pfeilklassen) verfügen nicht über eine Lage im Raum und können deshalb auch keinen Körper begrenzen. Um eine zu komplizierte Sprechweise zu vermeiden, wird dennoch gesagt, dass Vektoren ein Parallelogramm oder Parallelepiped aufspannen, wobei dann jeweils repräsentierende Pfeile gemeint sind.

3.4.2 Kollineare und komplanare Vektoren

In gewissen Fällen lässt sich ein Vektor $\vec{x} \in \mathbb{R}^3$ auch als Linearkombination von nur zwei Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^3$ darstellen.

Beispiel 3.6

Der Vektor $\vec{x} = \begin{pmatrix} 3,5 \\ 5 \\ 0,5 \end{pmatrix}$ wird als Linearkombination von $\vec{u} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix}$ dargestellt (siehe Abb. 3.29). Das lineare Gleichungssystem

$$\begin{aligned} \lambda + \mu &= 3,5 \\ 2 \cdot \lambda + \mu &= 5 \\ 3 \cdot \lambda - 2 \cdot \mu &= 0,5. \end{aligned}$$

ist lösbar, als Lösung erhält man $\lambda = \frac{3}{2}$, $\mu = 2$. Es gilt also

$$\vec{x} = \frac{3}{2} \cdot \vec{u} + 2 \cdot \vec{v}.$$

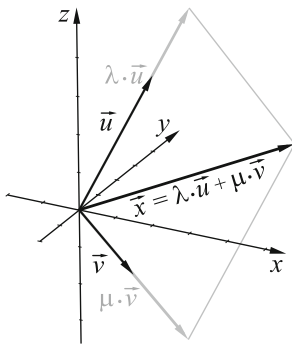


Abb. 3.29: Darstellung eines Vektors $\vec{x} \in \mathbb{R}^3$ als Linearkombination zweier Vektoren $\vec{u}, \vec{v}$

Die drei Vektoren $\vec{u}, \vec{v}$ und $\vec{x}$ in dem Beispiel 3.6 werden durch Pfeile repräsentiert, die in einer Ebene liegen, man sagt $\vec{u}, \vec{v}$ und $\vec{x}$ sind *komplanar*. Hingegen trifft für keine Auswahl von drei der vier Vektoren $\vec{u}, \vec{v}, \vec{w}$ und $\vec{x}$ aus dem Beispiel 3.5 zu, dass sie komplanar sind.

So wie drei Vektoren als komplanar bezeichnet werden, wenn sie (bzw. repräsentierende Pfeile) in einer Ebene liegen, werden zwei Vektoren als *kollinear* bezeichnet, wenn zugehörige Pfeile auf einer Geraden liegen oder parallel sind.

Definition 3.8

Zwei Vektoren $\vec{u}$ und $\vec{v}$ heißen *kollinear*, wenn einer der Vektoren ein Vielfaches des anderen Vektors ist, d. h. wenn z. B. $\lambda \in \mathbb{R}$ mit $\vec{u} = \lambda \cdot \vec{v}$ existiert.

Drei Vektoren $\vec{u}, \vec{v}$ und $\vec{w}$ heißen *komplanar*, wenn sich einer der Vektoren als Linearkombination der beiden anderen Vektoren darstellen lässt, d. h. wenn z. B. $\lambda, \mu \in \mathbb{R}$ existieren mit $\vec{w} = \lambda \cdot \vec{u} + \mu \cdot \vec{v}$. ♦

Bemerkungen:

- Die Definition 3.8 schließt den Nullvektor ein. Zwei Vektoren $\vec{u}, \vec{v}$, von denen einer der Nullvektor ist (z. B. $\vec{u} = \vec{0}$), sind stets kollinear, denn es gilt $\vec{u} = 0 \cdot \vec{v}$. Jedoch ist dann der andere Vektor $\vec{v}$ kein Vielfaches von $\vec{u}$. Ist hingegen keiner der beiden Vektoren $\vec{u}, \vec{v}$ der Nullvektor, so folgt aus der Existenz einer Zahl λ mit $\vec{u} = \lambda \cdot \vec{v}$, dass auch $\vec{v}$ ein Vielfaches von $\vec{u}$ ist, nämlich $\vec{v} = \frac{1}{\lambda} \cdot \vec{u}$.

- Drei Vektoren $\vec{u}$, $\vec{v}$, $\vec{w}$, von denen einer der Nullvektor ist (z. B. $\vec{w} = \vec{0}$), sind komplanar, denn es gilt $\vec{w} = 0 \cdot \vec{u} + 0 \cdot \vec{v}$. Weiterhin sind drei Vektoren stets komplanar, wenn zwei dieser Vektoren kollinear sind (Aufgabe 5, S. 121).

Satz 3.11

Sind $\vec{u}$, $\vec{v}$ und $\vec{w}$ drei komplanare Vektoren und gibt es unter diesen drei Vektoren kein Paar kollinear Vektoren, so lässt sich jeder der Vektoren $\vec{u}$, $\vec{v}$, $\vec{w}$ als Linearkombination der jeweils beiden anderen Vektoren darstellen.

Beweis: Nach Voraussetzung und Definition 3.8 lässt sich *einer* der Vektoren als Linearkombination der beiden anderen darstellen, z. B.

$$(*) \quad \vec{w} = \lambda \cdot \vec{u} + \mu \cdot \vec{v}.$$

Es muss $\lambda \neq 0$ sein, denn ansonsten wäre $\vec{w} = \mu \vec{v}$, was der Voraussetzung widerspricht, dass $\vec{v}$ und $\vec{w}$ nicht kollinear sind. Ebenso muss $\mu \neq 0$ gelten, ansonsten wären wegen $\vec{w} = \lambda \vec{u}$ die Vektoren $\vec{u}$ und $\vec{w}$ kollinear. Somit existieren $\frac{1}{\lambda}, \frac{1}{\mu} \in \mathbb{R}$, und die Gleichung $(*)$ lässt sich in die folgenden Gleichungen umformen:

$$\vec{u} = -\frac{\mu}{\lambda} \cdot \vec{v} + \frac{1}{\lambda} \cdot \vec{w}, \quad \vec{v} = -\frac{\lambda}{\mu} \cdot \vec{u} + \frac{1}{\mu} \cdot \vec{w}.$$

Somit ist also $\vec{u}$ als Linearkombination der Vektoren $\vec{v}$ und $\vec{w}$ sowie $\vec{v}$ als Linearkombination der Vektoren $\vec{u}$ und $\vec{w}$ darstellbar. $\square$

Satz 3.12

Drei Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ sind genau dann komplanar, wenn sich der Nullvektor auf nicht triviale Weise als Linearkombination dieser drei Vektoren darstellen lässt, d. h. falls λ, μ, ν existieren, die nicht alle gleich Null sind, so dass gilt:

$$\vec{0} = \lambda \cdot \vec{u} + \mu \cdot \vec{v} + \nu \cdot \vec{w}.$$

Bemerkung: Der Nullvektor lässt sich stets „auf triviale Weise“ als Linearkombination beliebiger Vektoren schreiben, indem für alle Koeffizienten Null gesetzt wird, z. B. $\vec{0} = 0 \cdot \vec{u} + 0 \cdot \vec{v} + 0 \cdot \vec{w}$. Hingegen verlangt Satz 3.12 die Existenz einer Darstellung mit Koeffizienten, von denen mindestens einer von Null verschieden sein muss, dies wird als *nicht triviale Linearkombination* bezeichnet.

Beweis: Da es sich bei dem Satz 3.12 um eine „genau-dann-wenn-Aussage“ handelt, sind zwei Richtungen zu beweisen.

- Falls $\vec{u}$, $\vec{v}$ und $\vec{w}$ komplanar sind, lässt sich einer der drei Vektoren als Linearkombination der beiden anderen Vektoren darstellen, z. B. $\vec{w} = \lambda \cdot \vec{u} + \mu \cdot \vec{v}$. Durch Umstellen ergibt sich daraus

$$\vec{0} = \lambda \cdot \vec{u} + \mu \cdot \vec{v} + (-1) \cdot \vec{w}.$$

Da der zu $\vec{w}$ gehörige Koeffizient -1 von Null verschieden ist, ist der Nullvektor somit als nicht triviale Linearkombination von $\vec{u}$, $\vec{v}$ und $\vec{w}$ dargestellt.

- Es sei $\vec{0}$ in der Form $\vec{0} = \lambda \cdot \vec{u} + \mu \cdot \vec{v} + \nu \cdot \vec{w}$ darstellbar, wobei mindestens einer der Koeffizienten λ, μ, ν von Null verschieden ist. Falls $\lambda \neq 0$ ist, so folgt

$$\vec{u} = -\frac{\mu}{\lambda} \cdot \vec{v} - \frac{\nu}{\lambda} \cdot \vec{w},$$

$\vec{u}$ ist also eine Linearkombination von $\vec{v}$ und $\vec{w}$. Ist $\mu \neq 0$ oder $\nu \neq 0$, so lässt sich in analoger Weise $\vec{v}$ oder $\vec{w}$ als Linearkombination der anderen Vektoren darstellen. Somit sind $\vec{u}$, $\vec{v}$ und $\vec{w}$ in jedem Falle komplanar. $\square$

3.4.3 Anwendungen von Linearkombinationen in der Geometrie

Mithilfe von Vektoren, insbesondere durch Linearkombinationen, lassen sich recht effektiv viele Sätze der Geometrie beweisen. Dadurch dass Vektoren Arithmetik und Geometrie miteinander verknüpfen, ergeben sich zudem interessante Bezüge, z. B. zwischen Schwerpunkten von Dreiecken und Tetraedern sowie den arithmetischen Mitteln der Koordinaten ihrer Eckpunkte.

Beispiel 3.7

Satz von Varignon: Die Mittelpunkte M_{AB} , M_{BC} , M_{CD} und M_{DA} der Seiten eines beliebigen Vierecks $ABCD$ bilden ein Parallelogramm (siehe Abb. 3.30).

Um diesen Satz mit vektoriellen Mitteln zu beweisen, genügt es, zu zeigen, dass $\overrightarrow{M_{AB}M_{BC}} = \overrightarrow{M_{DA}M_{CD}}$ gilt, denn daraus folgt, dass die gegenüberliegenden Seiten $\overline{M_{AB}M_{BC}}$ und $\overline{M_{DA}M_{CD}}$ des Vierecks $M_{AB}M_{BC}M_{CD}M_{DA}$ parallel und gleich lang sind.

Da M_{AB} der Mittelpunkt von $\overline{AB}$ und M_{BC} der Mittelpunkt von $\overline{BC}$ ist, gilt

$$\overrightarrow{M_{AB}M_{BC}} = \overrightarrow{M_{AB}B} + \overrightarrow{BM_{BC}} = \frac{1}{2} \overrightarrow{AB} + \frac{1}{2} \overrightarrow{BC} = \frac{1}{2} \overrightarrow{AC}.$$

Analog gilt

$$\overrightarrow{M_{DA}M_{CD}} = \overrightarrow{M_{DA}D} + \overrightarrow{DM_{CD}} = \frac{1}{2} \overrightarrow{AD} + \frac{1}{2} \overrightarrow{DC} = \frac{1}{2} \overrightarrow{AC}.$$

Damit ist die Behauptung $\overrightarrow{M_{AB}M_{BC}} = \overrightarrow{M_{DA}M_{CD}}$ erfüllt. ■

Beispiel 3.8

In jedem Parallelogramm halbieren die Diagonalen einander.

Um diesen Satz zu beweisen, betrachten wir ein Parallelogramm $ABCD$ mit dem Schnittpunkt E der Diagonalen, siehe Abb. 3.31. Da gegenüberliegende Seiten in einem Parallelogramm parallel und gleich lang sind, repräsentieren $\overrightarrow{AB}$ und $\overrightarrow{DC}$ sowie $\overrightarrow{BC}$ und $\overrightarrow{AD}$ (bei Bezeichnung der Punkte wie in Abb. 3.31) jeweils denselben Vektor, es gilt somit $\overrightarrow{BC} = \overrightarrow{AD}$ und $\overrightarrow{AB} = -\overrightarrow{CD}$.

Da der Diagonalschnittpunkt E den beiden Strecken $\overline{AC}$ und $\overline{BD}$ angehört, sind die Vektoren $\overrightarrow{AE}$ und $\overrightarrow{AC}$ sowie $\overrightarrow{BE}$ und $\overrightarrow{BD}$ jeweils kollinear. Somit existieren also $\lambda, \mu \in \mathbb{R}$ mit

$$\overrightarrow{AE} = \lambda \cdot \overrightarrow{AC} \quad \text{und} \quad \overrightarrow{BE} = \mu \cdot \overrightarrow{BD}.$$

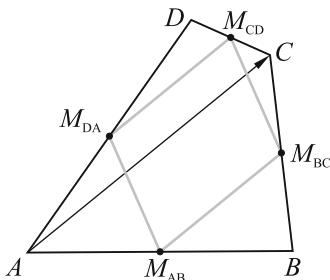


Abb. 3.30: Seitenmittenviereck

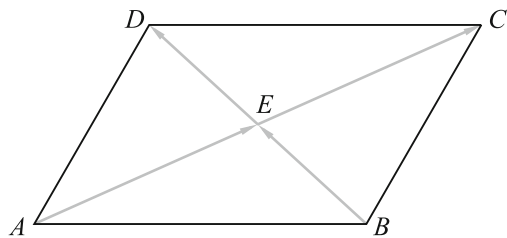


Abb. 3.31: Diagonalen eines Parallelogramms

Wir weisen nach, dass $\lambda = \mu = \frac{1}{2}$ sein muss. Dazu nutzen wir Beziehungen zwischen Vektoren, die aus der Addition von repräsentierenden Pfeilen innerhalb des Parallelogramms $ABCD$ resultieren (siehe Definition 3.3 auf S. 91). Es gilt:

$$\begin{aligned}\overrightarrow{AB} &= \overrightarrow{AE} + \overrightarrow{EB} = \lambda \cdot \overrightarrow{AC} + (-\mu) \cdot \overrightarrow{BD} = \lambda \cdot (\overrightarrow{AB} + \overrightarrow{BC}) - \mu \cdot (\overrightarrow{BC} + \overrightarrow{CD}) \\ &= \lambda \cdot (\overrightarrow{AB} + \overrightarrow{BC}) + \mu \cdot (-\overrightarrow{BC} + \overrightarrow{AB}) = \lambda \cdot \overrightarrow{AB} + \mu \cdot \overrightarrow{AB} + \lambda \cdot \overrightarrow{BC} - \mu \cdot \overrightarrow{BC} \\ &= (\lambda + \mu) \cdot \overrightarrow{AB} + (\lambda - \mu) \cdot \overrightarrow{BC}\end{aligned}$$

bzw.

$$(\lambda + \mu - 1) \cdot \overrightarrow{AB} + (\lambda - \mu) \cdot \overrightarrow{BC} = \vec{0}.$$

Die Vektoren $\overrightarrow{AB}$ und $\overrightarrow{BC}$ sind, da sie ein Parallelogramm aufspannen, nicht kollinear. Deshalb kann der Nullvektor nur auf triviale Weise als Linearkombination dieser Vektoren entstehen (siehe die Aufgabe 3 auf S. 121). Die Koeffizienten vor $\overrightarrow{AB}$ und $\overrightarrow{BC}$ in der obigen Gleichung müssen also jeweils Null sein:

$$\begin{aligned}\lambda + \mu - 1 &= 0 \\ \lambda - \mu &= 0\end{aligned}$$

Durch Lösen dieses LGS ergibt sich $\lambda = \frac{1}{2}$, $\mu = \frac{1}{2}$ und somit die Behauptung $\overrightarrow{AE} = \frac{1}{2} \cdot \overrightarrow{AC}$ sowie $\overrightarrow{BE} = \frac{1}{2} \cdot \overrightarrow{BD}$. ■

Beispiel 3.9

In einem beliebigen Dreieck $\triangle ABC$ schneiden sich die Seitenhalbierenden in einem Punkt. Dieser teilt die Seitenhalbierenden jeweils im Verhältnis 2:1.

Wir zeigen, dass alle Seitenhalbierenden durch den Punkt S mit

$$\overrightarrow{AS} = \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC}$$

verlaufen, und dass S alle drei Seitenhalbierenden im Verhältnis 2:1 teilt. Dies ist genau dann der Fall, wenn (mit den Bezeichnungen in Abb. 3.32) gilt:

$$\overrightarrow{AS} = \frac{2}{3} \overrightarrow{AM_{BC}}, \quad \overrightarrow{BS} = \frac{2}{3} \overrightarrow{BM_{AC}} \quad \text{sowie} \quad \overrightarrow{CS} = \frac{2}{3} \overrightarrow{CM_{AB}}.$$

Wegen $\overrightarrow{AM_{BC}} = \overrightarrow{AB} + \overrightarrow{BM_{BC}} = \overrightarrow{AB} + \frac{1}{2} \overrightarrow{BC}$ und $\overrightarrow{BC} = \overrightarrow{BA} + \overrightarrow{AC} = \overrightarrow{AC} - \overrightarrow{AB}$ ist $\frac{2}{3} \overrightarrow{AM_{BC}} = \frac{2}{3} \overrightarrow{AB} + \frac{2}{3} \cdot \frac{1}{2} \overrightarrow{BC} = \frac{2}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC} - \frac{1}{3} \overrightarrow{AB} = \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC} = \overrightarrow{AS}$. Auf analoge Weise lässt sich zeigen, dass $\overrightarrow{BS} = \frac{2}{3} \overrightarrow{BM_{AC}}$ und $\overrightarrow{CS} = \frac{2}{3} \overrightarrow{CM_{AB}}$ gilt, der Punkt S also alle drei Seitenhalbierenden im Verhältnis 2:1 teilt.

Der Schnittpunkt S der Seitenhalbierenden eines Dreiecks wird auch als *Schwerpunkt* bezeichnet. Aus der Darstellung $\overrightarrow{AS} = \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC}$ folgt für die Koordinaten des Schwerpunktes (siehe Abschnitt 3.3.1):

$$\begin{pmatrix} x_S - x_A \\ y_S - y_A \end{pmatrix} = \frac{1}{3} \begin{pmatrix} x_B - x_A \\ y_B - y_A \end{pmatrix} + \frac{1}{3} \begin{pmatrix} x_C - x_A \\ y_C - y_A \end{pmatrix}$$

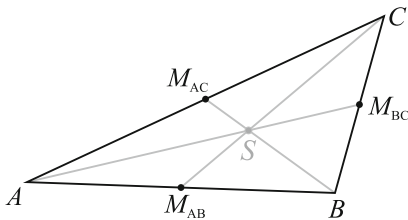


Abb. 3.32: Schwerpunkt eines Dreiecks

und somit $x_s = \frac{1}{3} (x_A + x_B + x_C)$, $y_s = \frac{1}{3} (y_A + y_B + y_C)$.

Die Koordinaten des Schwerpunktes sind also die arithmetischen Mittel der Koordinaten der Eckpunkte eines Dreiecks. ■

Eine besondere Stärke vektorieller Methoden in der Geometrie besteht darin, dass in vielen Fällen Vorgehensweisen, die bei ebenen Figuren angewendet werden, leicht auf dazu analoge Körper im Raum übertragen werden können.

Beispiel 3.10

In einem beliebigen Tetraeder $ABCD$ schneiden sich die vier Schwerlinien, d. h. die Verbindungsstrecken $\overline{AS_{BCD}}$, $\overline{BS_{ACD}}$, $\overline{CS_{ABD}}$ und $\overline{DS_{ABC}}$ der Eckpunkte mit den Schwerpunkten der jeweils gegenüberliegenden Seitenflächen, in einem Punkt. Dieser teilt die Schwerlinien jeweils im Verhältnis 3:1.

Wir zeigen, dass alle Schwerlinien durch den Punkt S (den Schwerpunkt des Tetraeders) mit

$$\overrightarrow{AS} = \frac{1}{4} \overrightarrow{AB} + \frac{1}{4} \overrightarrow{AC} + \frac{1}{4} \overrightarrow{AD}.$$

verlaufen und dass dieser Punkt alle drei Schwerlinien im Verhältnis 3:1 teilt. Für die Schwerpunkte der Seitenflächen gilt (mit den Bezeichnungen in Abb. 3.33):

$$\begin{aligned} \overrightarrow{AS_{ABC}} &= \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC}, & \overrightarrow{AS_{ABD}} &= \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AD}, \\ \overrightarrow{AS_{ACD}} &= \frac{1}{3} \overrightarrow{AC} + \frac{1}{3} \overrightarrow{AD}, & \overrightarrow{BS_{BCD}} &= \frac{1}{3} \overrightarrow{BC} + \frac{1}{3} \overrightarrow{BD} \end{aligned}$$

(vgl. Beispiel 3.9). Um die Behauptung nachzuweisen, ist zu zeigen:

$$\overrightarrow{AS_{BCD}} = \frac{4}{3} \overrightarrow{AS}, \quad \overrightarrow{BS_{ACD}} = \frac{4}{3} \overrightarrow{BS}, \quad \overrightarrow{CS_{ABD}} = \frac{4}{3} \overrightarrow{CS}, \quad \overrightarrow{DS_{ABC}} = \frac{4}{3} \overrightarrow{DS}.$$

Es gilt:

$$\begin{aligned} \overrightarrow{AS_{BCD}} &= \overrightarrow{AB} + \overrightarrow{BS_{BCD}} = \overrightarrow{AB} + \frac{1}{3} \overrightarrow{BC} + \frac{1}{3} \overrightarrow{BD} \\ &= \overrightarrow{AB} + \frac{1}{3} (-\overrightarrow{AB} + \overrightarrow{AC}) + \frac{1}{3} (-\overrightarrow{AB} + \overrightarrow{AD}) \\ &= \frac{1}{3} \overrightarrow{AB} + \frac{1}{3} \overrightarrow{AC} + \frac{1}{3} \overrightarrow{AD}, \end{aligned}$$

$$\frac{4}{3} \overrightarrow{AS} = \frac{4}{3} \left(\frac{1}{4} \overrightarrow{AB} + \frac{1}{4} \overrightarrow{AC} + \frac{1}{4} \overrightarrow{AD} \right) = \frac{1}{3} (\overrightarrow{AB} + \overrightarrow{AC} + \overrightarrow{AD}).$$

Somit ist also $\overrightarrow{AS_{BCD}} = \frac{4}{3} \overrightarrow{AS}$. Auf dieselbe Weise lässt sich nachweisen, dass auch $\overrightarrow{BS_{ACD}} = \frac{4}{3} \overrightarrow{BS}$, $\overrightarrow{CS_{ABD}} = \frac{4}{3} \overrightarrow{CS}$ und $\overrightarrow{DS_{ABC}} = \frac{4}{3} \overrightarrow{DS}$ gilt, der Punkt S also alle vier Schwerlinien im Verhältnis 3:1 teilt.

Wie in dem Beispiel 3.9 für den Schwerpunkt eines Dreiecks gezeigt wurde, lassen sich auch die Koordinaten des Schwerpunktes eines Tetraeders als die arithmetischen Mittel der Koordinaten der Eckpunkte berechnen. ■

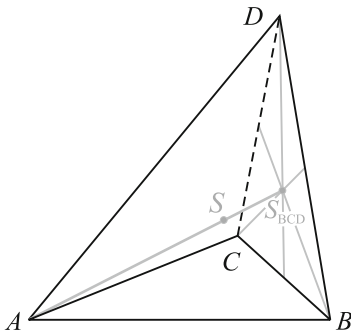


Abb. 3.33: Schwerpunkt eines Tetraeders

3.4.4 Aufgaben zu Abschnitt 3.4

- Überprüfen Sie, ob sich der Vektor $\vec{x}$ als Linearkombination der Vektoren $\vec{u}$ und $\vec{v}$ bzw. $\vec{u}$, $\vec{v}$ und $\vec{w}$ darstellen lässt. Geben Sie, falls dies möglich ist, die Koeffizienten der Linearkombination an.
 - $\vec{x} = \begin{pmatrix} 4 \\ -5 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 12 \\ -3 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 3 \\ -2 \end{pmatrix}$ b) $\vec{x} = \begin{pmatrix} 13 \\ 1 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 2 \\ -\frac{1}{3} \end{pmatrix}$, $\vec{v} = \begin{pmatrix} -5 \\ \frac{5}{6} \end{pmatrix}$
 - $\vec{x} = \begin{pmatrix} -8 \\ 3 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 5 \\ -9 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 10 \\ 6 \end{pmatrix}$, $\vec{w} = \begin{pmatrix} -9 \\ 2 \end{pmatrix}$
- Stellen Sie den Vektor $\vec{x} \in \mathbb{R}^3$ als Linearkombination der Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ dar, falls dies möglich ist.
 - $\vec{x} = \begin{pmatrix} 0 \\ 4 \\ -2 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix}$, $\vec{w} = \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix}$
 - $\vec{x} = \begin{pmatrix} -2 \\ 1 \\ 1 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 4 \\ -9 \\ 2 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 3 \\ 6 \\ -12 \end{pmatrix}$, $\vec{w} = \begin{pmatrix} 2 \\ -8 \\ 2 \end{pmatrix}$
 - $\vec{x} = \begin{pmatrix} 1 \\ 3 \\ 2 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} -1 \\ 3 \\ -4 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 3 \\ -5 \\ -2 \end{pmatrix}$, $\vec{w} = \begin{pmatrix} 7 \\ -9 \\ -14 \end{pmatrix}$
- Weisen Sie nach: Sind zwei vom Nullvektor verschiedene Vektoren $\vec{u}$ und $\vec{v}$ nicht kollinear und gilt $\vec{o} = \lambda \vec{u} + \mu \vec{v}$, so ist $\lambda = \mu = 0$.
- Welche der folgenden Mengen von Vektoren sind komplanar?
 - $M_2 = \left\{ \begin{pmatrix} 3 \\ 5 \\ 2 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}; \begin{pmatrix} 3 \\ 6 \\ \frac{3}{2} \end{pmatrix} \right\}$ b) $M_3 = \left\{ \begin{pmatrix} 3 \\ 5 \\ 2 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ -1 \end{pmatrix}; \begin{pmatrix} 2 \\ 4 \\ 1 \end{pmatrix} \right\}$
- Weisen Sie nach: Wenn es unter drei Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ zwei kollineare Vektoren gibt, so sind $\vec{u}$, $\vec{v}$ und $\vec{w}$ komplanar.
- Als Seitenmittendreieck eines Dreiecks $\triangle ABC$ bezeichnet man das Dreieck, dessen Eckpunkte die Mittelpunkte M_{AB} , M_{BC} und M_{CA} der Seiten des Dreiecks $\triangle ABC$ sind, siehe Abb. 3.34. Beweisen Sie, dass für ein beliebiges Dreieck $\triangle ABC$ der Schwerpunkt des Dreiecks $\triangle ABC$ und der Schwerpunkt seines Seitenmittendreiecks $\triangle M_{AB}M_{BC}M_{CA}$ identisch sind.
- Weisen Sie nach: In jedem Parallelogramm schneiden sich die Verbindungsstrecke eines beliebigen Eckpunktes mit dem Mittelpunkt der gegenüberliegenden Seite und die Diagonale durch die benachbarten Eckpunkte im Verhältnis 1:2 (siehe Abb. 3.35).

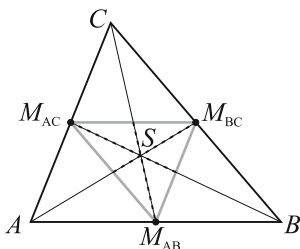


Abb. 3.34: Seitenmittendreieck

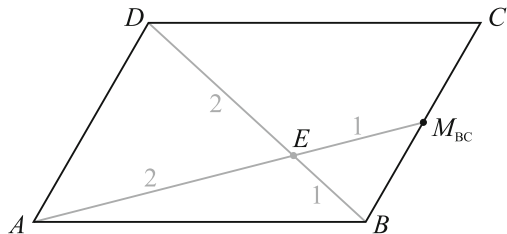


Abb. 3.35: Skizze zu Aufgabe 7

3.5 Das Skalarprodukt zweier Vektoren

3.5.1 Arithmetische Einführung des Skalarprodukts

Beispiel 3.11

In dem Beispiel 3.1 auf S.100 wurden Stücklisten als Beispiele für n -Tupel betrachtet. Wir beziehen nun die Stückpreise der einzelnen Gleisbauteile ein.

Gleisstück gerade	2,40 €	Weiche links	17,98 €
Gleisstück gebogen	2,70 €	Weiche rechts	17,98 €
Anschluss-Gleisstück	6,29 €	Weichenantrieb	12,98 €

Daraus ergibt sich (bei Einhaltung der Reihenfolge der Teile, die in dem Bei-

spiel 3.1 für die Stücklisten gewählt wurde) der *Preisvektor* $\vec{p} = \begin{pmatrix} 2,40 \\ 2,70 \\ 6,29 \\ 17,98 \\ 17,98 \\ 12,98 \end{pmatrix}$.

Unter Heranziehung der Stückliste eines Sortiments, die wir im Folgenden als

Teilevektor bezeichnen, z. B. $\vec{e}_2 = \begin{pmatrix} 15 \\ 8 \\ 1 \\ 2 \\ 2 \\ 4 \end{pmatrix}$ für das Ergänzungssortiment 2, lässt

sich der Gesamtpreis eines Sortiments durch die Summe der Produkte einander entsprechender Komponenten des Teile- und des Preisvektors ausdrücken:

$$15 \cdot 2,40 + 8 \cdot 2,70 + 1 \cdot 6,29 + 2 \cdot 17,98 + 2 \cdot 17,98 + 4 \cdot 12,98 = 187,73 \quad (\text{in } \text{€}).$$

Dieser Gesamtpreis bezieht sich darauf, dass die Teile einzeln gekauft werden. Beim Kauf von Sortimentskästen, werden i. Allg. andere, meist niedrigere, Preise berechnet.

Diese Produktsomme, welche den Gesamtpreis ergibt, werden wir im Folgenden als Skalarprodukt der Vektoren $\vec{e}_2$ und $\vec{p}$ bezeichnen. ■

Definition 3.9

Als *Skalarprodukt* zweier Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^n$ mit $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix}$

wird die Summe der Produkte der einander entsprechenden Komponenten von $\vec{u}$ und $\vec{v}$ bezeichnet:

$$\vec{u} \cdot \vec{v} = u_1 \cdot v_1 + u_2 \cdot v_2 + \dots + u_n \cdot v_n = \sum_{i=1}^n u_i \cdot v_i. \quad \blacklozenge$$

Bemerkung: Der Name *Skalarprodukt* rührt daher, dass zwei Vektoren ein Skalar, d. h. eine reelle Zahl zugeordnet wird. Beim Rechnen mit Vektoren treten damit nunmehr drei Arten von Produkten auf, für die jeweils das Zeichen „ $\cdot$ “ verwendet wird:

- das Produkt zweier reeller Zahlen,
- das Produkt eines Vektors mit einer reellen Zahl,
- das Skalarprodukt zweier Vektoren.

Welche Bedeutung das Zeichen „ $\cdot$ “ jeweils hat, geht aus der Art von Objekten (Vektoren oder reelle Zahlen) hervor, zwischen denen es steht, siehe hierzu die Aufgabe 1 auf S. 131.

Rechenregeln für das Skalarprodukt

Satz 3.13

Für beliebige Vektoren $\vec{u}$, $\vec{v}$, $\vec{w}$ und alle $\lambda \in \mathbb{R}$ gilt:

1. $\vec{u} \cdot \vec{v} = \vec{v} \cdot \vec{u}$ (Kommutativgesetz),
2. $(\lambda \vec{u}) \cdot \vec{v} = \lambda (\vec{u} \cdot \vec{v})$ bzw. $\vec{u} \cdot (\lambda \vec{v}) = \lambda (\vec{u} \cdot \vec{v})$,
3. $(\vec{u} + \vec{v}) \cdot \vec{w} = \vec{u} \cdot \vec{w} + \vec{v} \cdot \vec{w}$ (Distributivgesetz),
4. $\vec{u} \cdot \vec{u} \geq 0$ und $\vec{u} \cdot \vec{u} = 0 \Leftrightarrow \vec{u} = \vec{0}$.

Beweis:

1. Es werden die Definition 3.9 und die Kommutativität der Multiplikation reeller Zahlen angewendet:

$$\vec{u} \cdot \vec{v} = u_1 \cdot v_1 + u_2 \cdot v_2 + \dots + u_n \cdot v_n = v_1 \cdot u_1 + v_2 \cdot u_2 + \dots + v_n \cdot u_n = \vec{v} \cdot \vec{u},$$

$$\text{bzw. in Kurzschreibweise: } \vec{u} \cdot \vec{v} = \sum_{i=1}^n u_i \cdot v_i = \sum_{i=1}^n v_i \cdot u_i = \vec{v} \cdot \vec{u}.$$

2. Die Aussage lässt sich aus der Kommutativität und Assoziativität der Multiplikation in $\mathbb{R}$ ableiten. Nach Definition 3.6 auf S. 103 ist $\lambda \cdot \vec{u} = \begin{pmatrix} \lambda u_1 \\ \lambda u_2 \\ \vdots \\ \lambda u_n \end{pmatrix}$.

Damit gilt $(\lambda \vec{u}) \cdot \vec{v} = \sum_{i=1}^n (\lambda \cdot u_i) \cdot v_i = \sum_{i=1}^n \lambda \cdot (u_i \cdot v_i) = \lambda (\vec{u} \cdot \vec{v})$. Auf völlig analoge Weise wird gezeigt, dass $\vec{u} \cdot (\lambda \vec{v}) = \lambda (\vec{u} \cdot \vec{v})$ gilt.

3. Unter Verwendung der Definitionen 3.6 (S. 103) und 3.9 sowie des Distributivgesetzes der reellen Zahlen lassen sich folgende Umformungen durchführen:

$$\begin{aligned} (\vec{u} + \vec{v}) \cdot \vec{w} &= \sum_{i=1}^n (u_i + v_i) \cdot w_i = \sum_{i=1}^n u_i \cdot w_i + v_i \cdot w_i \\ &= \sum_{i=1}^n u_i \cdot w_i + \sum_{i=1}^n v_i \cdot w_i = \vec{u} \cdot \vec{w} + \vec{v} \cdot \vec{w}. \end{aligned}$$

4. Nach Definition 3.9 ist $\vec{u} \cdot \vec{u} = \sum_{i=1}^n u_i \cdot u_i = \sum_{i=1}^n u_i^2$. Da eine Summe von Qua-

draten niemals negativ sein kann, gilt $\vec{u} \cdot \vec{u} \geq 0$. Die Summe $\sum_{i=1}^n u_i^2$ ist genau dann gleich Null, wenn jeder ihrer Summanden Null ist, also $u_i^2 = 0$ und damit $u_i = 0$ für alle i ($i = 1 \dots n$) gilt. Dies bedeutet, dass $\vec{u}$ der Nullvektor ist, womit auch der zweite Teil der Behauptung bewiesen wurde. $\square$

Das Skalarprodukt $\vec{u} \cdot \vec{u}$ eines Vektors mit sich selbst wird auch mit $\vec{u}^2$ bezeichnet, siehe hierzu die Aufgabe 4 b auf S. 132.

3.5.2 Geometrische Deutung des Skalarprodukts

Der Betrag eines Vektors

In dem Abschnitt 2.2 wurde herausgearbeitet, dass der Abstand zweier Punkte $P_1(x_1, y_1)$, $P_2(x_2, y_2)$ der Ebene $\sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2}$ ist und der Abstand zweier Punkte des Raumes $\sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2}$ beträgt. Dieser Abstand entspricht der Länge des Pfeils $\overrightarrow{P_1 P_2}$. Wird ein Vektor $\vec{p}$ durch den Pfeil $\overrightarrow{P_1 P_2}$ repräsentiert, so lässt sich dieser Vektor auch durch das Paar bzw. Tripel $\begin{pmatrix} x_p \\ y_p \end{pmatrix} = \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \end{pmatrix}$ bzw. $\begin{pmatrix} x_p \\ y_p \\ z_p \end{pmatrix} = \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \\ z_2 - z_1 \end{pmatrix}$ darstellen, siehe Abschnitt 3.3.1.

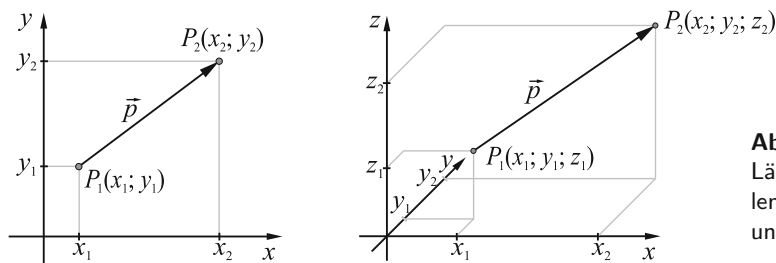


Abb. 3.36:
Länge von Pfeilen in der Ebene und im Raum

Bildet man nach der Definition 3.9 das Skalarprodukt eines Vektors $\vec{p} = \begin{pmatrix} x_p \\ y_p \\ z_p \end{pmatrix}$ mit sich selbst, so erhält man $\vec{p} \cdot \vec{p} = x_p \cdot x_p + y_p \cdot y_p + z_p \cdot z_p = x_p^2 + y_p^2 + z_p^2$ bzw. mit den Koordinaten des Anfangs- und des Endpunktes des Pfeils $\overrightarrow{P_1 P_2}$ (siehe oben): $\vec{p} \cdot \vec{p} = (x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2$.

Das Skalarprodukt eines Vektors mit sich selbst ist also gleich dem Quadrat der Länge eines dem Vektor zugeordneten Pfeils; umgekehrt errechnet sich die Länge eines Pfeils als Wurzel des Skalarproduktes des zugehörigen Vektors mit sich selbst. Wir verallgemeinern dieses Konzept nun für beliebige Vektoren in $\mathbb{R}^n$, wobei wir von Beträgen von Vektoren statt von Längen entsprechender Pfeile sprechen.

Definition 3.10

Als *Betrag eines Vektors* $\vec{u}$ bezeichnet man die Wurzel aus dem Skalarprodukt dieses Vektors mit sich selbst:

$$|\vec{u}| = \sqrt{\vec{u} \cdot \vec{u}} = \sqrt{u^2}.$$



Von fundamentaler Bedeutung ist die folgende Beziehung zwischen dem Skalarprodukt zweier Vektoren und ihren Beträgen.

Satz 3.14

Cauchy-Schwarzsche Ungleichung:

$$\text{Für beliebige Vektoren } \vec{u} \text{ und } \vec{v} \text{ gilt } |\vec{u} \cdot \vec{v}| \leq |\vec{u}| \cdot |\vec{v}|.$$

Bemerkung: Das Multiplikationszeichen und die Betragsstriche treten in dem Satz 3.14 jeweils in zwei Bedeutungen auf. Auf der linken Seite der Ungleichung wird zunächst das Skalarprodukt der Vektoren $\vec{u}$ und $\vec{v}$ berechnet, dabei handelt

es sich um eine reelle Zahl; von dieser wird der Absolutbetrag gebildet. Auf der rechten Seite werden demgegenüber zunächst Beträge von Vektoren berechnet, wobei es sich um Zahlen handelt; anschließend werden diese multipliziert.

Beweis: Falls einer der beiden Vektoren $\vec{u}$ und $\vec{v}$ der Nullvektor ist, ergibt sich nach Definition 3.9 und Satz 3.13 sofort, dass auf beiden Seiten der Cauchy-Schwarzschen Ungleichung Null steht, die Ungleichung somit erfüllt ist. Wir betrachten im Folgenden nur den Fall, dass $\vec{u} \neq \vec{o}$, $\vec{v} \neq \vec{o}$ und somit $|\vec{u}| \neq 0$, $|\vec{v}| \neq 0$ ist. Damit existieren die Reziproken $\frac{1}{|\vec{u}|}$, $\frac{1}{|\vec{v}|}$ der Beträge beider Vektoren. Wir verwenden im Folgenden die abkürzende Schreibweise $\frac{\vec{u}}{|\vec{u}|}$ für $\frac{1}{|\vec{u}|} \vec{u}$.

Nach dem Satz 3.13, 4 ist das Skalarprodukt eines beliebigen Vektors mit sich selbst niemals negativ. Da dies auch für den Vektor $\frac{\vec{u}}{|\vec{u}|} - \frac{\vec{v}}{|\vec{v}|}$ zutrifft, gilt

$$\left(\frac{\vec{u}}{|\vec{u}|} - \frac{\vec{v}}{|\vec{v}|} \right) \cdot \left(\frac{\vec{u}}{|\vec{u}|} - \frac{\vec{v}}{|\vec{v}|} \right) \geq 0.$$

Die linke Seite dieser Ungleichung wird nun mithilfe von Satz 3.13 umgeformt:

$$\begin{aligned} 0 &\leq \left(\frac{\vec{u}}{|\vec{u}|} - \frac{\vec{v}}{|\vec{v}|} \right) \cdot \left(\frac{\vec{u}}{|\vec{u}|} - \frac{\vec{v}}{|\vec{v}|} \right) = \frac{\vec{u}}{|\vec{u}|} \cdot \frac{\vec{u}}{|\vec{u}|} + \frac{\vec{v}}{|\vec{v}|} \cdot \frac{\vec{v}}{|\vec{v}|} - 2 \cdot \frac{\vec{u}}{|\vec{u}|} \cdot \frac{\vec{v}}{|\vec{v}|} \\ &= \frac{\vec{u} \cdot \vec{u}}{|\vec{u}|^2} + \frac{\vec{v} \cdot \vec{v}}{|\vec{v}|^2} - 2 \cdot \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|}. \end{aligned}$$

Wegen Definition 3.10 ist $\vec{u} \cdot \vec{u} = |\vec{u}|^2$ und $\vec{v} \cdot \vec{v} = |\vec{v}|^2$, die obere Ungleichung lässt sich also in der Form

$$0 \leq 2 - 2 \cdot \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|} \quad \text{bzw.} \quad \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|} \leq 1$$

schreiben. Somit gilt also $\vec{u} \cdot \vec{v} \leq |\vec{u}| \cdot |\vec{v}|$. Auf analoge Weise lässt sich (z. B. durch Einsetzen des Gegenvektors $-\vec{u}$ an Stelle von $\vec{u}$) zeigen, dass auch $-\vec{u} \cdot \vec{v} \leq |\vec{u}| \cdot |\vec{v}|$ ist. Aus beiden Ungleichungen zusammen ergibt sich die Behauptung

$$|\vec{u} \cdot \vec{v}| \leq |\vec{u}| \cdot |\vec{v}|. \quad \square$$

Das Skalarprodukt kollinearer Vektoren

Satz 3.15

Zwei Vektoren $\vec{u}$ und $\vec{v}$ seien kollinear. Sind $\vec{u}$ und $\vec{v}$ gleich orientiert, so gilt:

$$\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}|.$$

Falls $\vec{u}$ und $\vec{v}$ entgegengesetzt orientiert sind, so gilt:

$$\vec{u} \cdot \vec{v} = -|\vec{u}| \cdot |\vec{v}|.$$

Beweis: Nach Definition 3.8 sind zwei Vektoren $\vec{u}$ und $\vec{v}$ kollinear, wenn einer der Vektoren ein Vielfaches des anderen Vektors ist, also z. B. $\lambda \in \mathbb{R}$ mit $\vec{v} = \lambda \cdot \vec{u}$ existiert. Für das Skalarprodukt beider Vektoren gilt dann nach Satz 3.13, 2:

$$\vec{u} \cdot \vec{v} = \vec{u} \cdot (\lambda \vec{u}) = \lambda \cdot (\vec{u} \cdot \vec{u}).$$

Für den Betrag des Vektors $\vec{v}$ gilt nach Definition 3.10 und Satz 3.13, 2:

$$|\vec{v}| = |\lambda \vec{u}| = \sqrt{(\lambda \vec{u}) \cdot (\lambda \vec{u})} = \sqrt{\lambda \lambda \cdot (\vec{u} \cdot \vec{u})} = |\lambda| \cdot \sqrt{\vec{u} \cdot \vec{u}} = |\lambda| \cdot |\vec{u}|.$$

Somit gilt $|\vec{u}| \cdot |\vec{v}| = |\lambda| \cdot |\vec{u}| \cdot |\vec{u}| = |\lambda| \cdot (\vec{u} \cdot \vec{u})$. Wegen $\vec{u} \cdot \vec{v} = \lambda \cdot (\vec{u} \cdot \vec{u})$ ist für $\lambda > 0$ (was nach Definition 3.5 bedeutet, dass $\vec{u}$ und $\vec{v}$ gleich orientiert sind) $\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}|$. Sind $\vec{u}$ und $\vec{v}$ entgegengesetzt orientiert, also $\lambda < 0$, so ist $\vec{u} \cdot \vec{v} = -|\vec{u}| \cdot |\vec{v}|$. $\square$

Ohne Beweis sei angemerkt, dass auch die Umkehrung des Satzes 3.15 gilt, für zwei Vektoren $\vec{u}$ und $\vec{v}$ also *nur* dann $\vec{u} \cdot \vec{v} = \pm |\vec{u}| \cdot |\vec{v}|$ ist, wenn die beiden Vektoren kollinear sind.

Das Skalarprodukt orthogonaler Vektoren

Wir nennen zwei Vektoren orthogonal, wenn sie repräsentierende Pfeile besitzen, die auf zueinander senkrechten Geraden liegen oder wenn mindestens einer der beiden Vektoren der Nullvektor ist.

Satz 3.16

Zwei Vektoren $\vec{u}$ und $\vec{v}$ sind genau dann orthogonal, wenn $\vec{u} \cdot \vec{v} = 0$ gilt.

Beweis: Wir führen zwei Beweise für diesen Satz und wenden dabei sehr unterschiedlichen Vorgehensweisen an. Der erste Beweis beschränkt sich auf Vektoren in $\mathbb{R}^2$, während der zweite Beweis auch Pfeilklassen des Raumes erfasst.

1. Wir betrachten zwei Vektoren $\vec{u} = \begin{pmatrix} x_u \\ y_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \end{pmatrix}$.

- Ist $\vec{u}$ oder $\vec{v}$ der Nullvektor, so folgt $\vec{u} \cdot \vec{v} = 0$ direkt aus Definition 3.9.
- Es sei keiner der beiden Vektoren der Nullvektor, aber eine Komponente eines der Vektoren sei gleich Null (was bedeutet, dass dieser Vektor die Richtung einer der Koordinatenachsen angibt), z. B. $\vec{u} = \begin{pmatrix} x_u \\ 0 \end{pmatrix}$. Die beiden Vektoren $\vec{u}$ und $\vec{v}$ sind unter dieser Voraussetzung genau dann orthogonal, wenn $\vec{v}$ die Richtung der y -Achse angibt, also $x_v = 0$ ist. In diesem Falle ist $\vec{u} \cdot \vec{v} = x_u x_v + y_u y_v = x_u \cdot 0 + 0 \cdot y_v = 0$. Umgekehrt folgt aus $\vec{u} \cdot \vec{v} = x_u x_v + y_u y_v = 0$ mit $y_u = 0$, dass $x_u x_v = 0$ sein muss. Da wegen $\vec{u} \neq \vec{0}$ nicht $x_u = 0$ sein kann, ist $x_v = 0$. $\vec{u}$ und $\vec{v}$ sind daher orthogonal.
- Es ist noch der Fall zu untersuchen, dass keine Komponente eines der Vektoren $\vec{u}$ und $\vec{v}$ Null ist. Wir betrachten Pfeile, welche diese beiden Vektoren repräsentieren und ihren Anfangspunkt im Koordinatenursprung haben sowie die Geraden g und h , welchen diese Pfeile angehören, siehe Abb. 3.37. Die Anstiege dieser Geraden sind $m_g = \frac{y_u}{x_u}$ und $m_h = \frac{y_v}{x_v}$. Nach dem Satz 2.2 auf S. 54 sind die Geraden g und h genau dann senkrecht zueinander, wenn $m_h = -\frac{1}{m_g}$ gilt. Dies ist gleichbedeutend mit $\frac{y_v}{x_v} = -\frac{x_u}{y_u} \Leftrightarrow y_v \cdot y_u = -x_u \cdot x_v \Leftrightarrow y_u \cdot y_v + x_u \cdot x_v = 0 \Leftrightarrow \vec{u} \cdot \vec{v} = 0$.

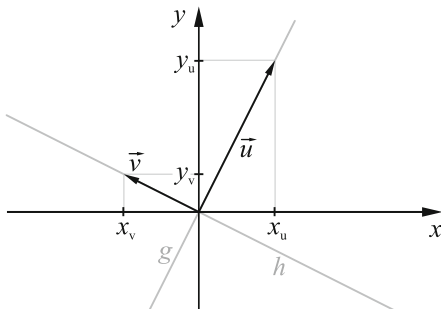


Abb. 3.37:
Orthogonale Vektoren in $\mathbb{R}^2$

2. Wir beweisen den Satz nun für Vektoren von $\mathbb{R}^3$ bzw. Pfeilklassen im Raum.

Es seien zwei Vektoren $\vec{u} = \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}$ gegeben, unter denen nicht der Nullvektor ist. (Ist $\vec{u} = \vec{o}$ oder $\vec{v} = \vec{o}$, so folgt $\vec{u} \cdot \vec{v} = 0$ wiederum direkt aus der Definition 3.9.) Wir betrachten zwei Pfeile $\overrightarrow{OU}$ und $\overrightarrow{OV}$, welche Repräsentanten von $\vec{u}$ bzw. $\vec{v}$ sind und deren Anfangspunkt der Koordinatenursprung ist, siehe Abb. 3.38. Die Vektoren sind genau dann orthogonal, wenn das Dreieck ΔUVO bei O rechtwinklig ist. Nach dem Satz des Pythagoras und seiner Umkehrung ist dies genau dann der Fall, wenn $|\overrightarrow{OU}|^2 + |\overrightarrow{OV}|^2 = |\overrightarrow{UV}|^2$ gilt. Diese Abstandsquadrate lassen sich folgendermaßen mithilfe der Vektoren $\vec{u}$ und $\vec{v}$ bzw. ihrer Komponenten ausdrücken:

$$|\overrightarrow{OU}|^2 = |\vec{u}|^2$$

$$|\overrightarrow{OV}|^2 = |\vec{v}|^2$$

$$\begin{aligned} |\overrightarrow{UV}|^2 &= |\vec{v} - \vec{u}|^2 = (x_v - x_u)^2 + (y_v - y_u)^2 + (z_v - z_u)^2 \\ &= x_v^2 + x_u^2 - 2x_u x_v + y_v^2 + y_u^2 - 2y_u y_v + z_v^2 + z_u^2 - 2z_u z_v \\ &= x_u^2 + y_u^2 + z_u^2 + x_v^2 + y_v^2 + z_v^2 - 2 \cdot (x_u x_v + y_u y_v + z_u z_v) \\ &= |\vec{u}|^2 + |\vec{v}|^2 - 2 \cdot \vec{u} \cdot \vec{v}. \end{aligned}$$

Es gilt somit $|\overrightarrow{OU}|^2 + |\overrightarrow{OV}|^2 = |\overrightarrow{UV}|^2$ genau dann, wenn $\vec{u} \cdot \vec{v} = 0$ ist, und genau in diesem Falle sind die Vektoren $\vec{u}$ und $\vec{v}$ orthogonal. $\square$

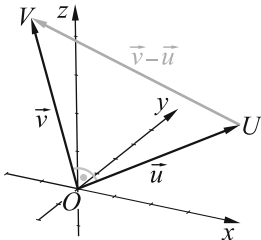


Abb. 3.38:
Orthogonale Vektoren im Raum

Bestimmung eines Vektors, der zu zwei gegebenen Vektoren orthogonal ist

Beispiel 3.12

Gesucht ist ein Vektor $\vec{n} = \begin{pmatrix} x_n \\ y_n \\ z_n \end{pmatrix}$, der zu $\vec{u} = \begin{pmatrix} 3 \\ -1 \\ 7 \end{pmatrix}$ und zu $\vec{v} = \begin{pmatrix} 2 \\ -5 \\ 9 \end{pmatrix}$ orthogonal ist. Dazu wird der Satz 3.16 herangezogen. Die Bedingungen $\vec{u} \cdot \vec{n} = 0$ und $\vec{v} \cdot \vec{n} = 0$ führen zu dem linearen Gleichungssystem

$$\begin{aligned} 3x_n - y_n + 7z_n &= 0 \\ 2x_n - 5y_n + 9z_n &= 0. \end{aligned}$$

Die Lösungen dieses LGS lassen sich in der Form $\begin{pmatrix} x_n \\ y_n \\ z_n \end{pmatrix} = t \cdot \begin{pmatrix} -2 \\ 1 \\ 1 \end{pmatrix}$ mit $t \in \mathbb{R}$ darstellen. Alle Vektoren dieser Form sind sowohl zu $\vec{u}$ als auch zu $\vec{v}$ orthogonal, so zum Beispiel der Vektor $\vec{n} = \begin{pmatrix} -2 \\ 1 \\ 1 \end{pmatrix}$. $\blacksquare$

Skalarprodukt und Winkel zwischen Vektoren

Mit den Sätzen 3.14, 3.15 und 3.16 ist nun bekannt, dass der Betrag des Skalarproduktes zweier Vektoren maximal ist, wenn diese Vektoren kollinear sind und minimal (nämlich Null), wenn die Vektoren orthogonal sind. Damit liegt die Vermutung nahe, dass das Skalarprodukt zweier Vektoren von ihren Beträgen und von dem Winkel zwischen den Vektoren abhängt. (Als Winkel zweier Vektoren betrachten wir den Winkel von zwei Geraden, auf denen repräsentierende Pfeile der betrachteten Vektoren liegen.)

Um den Zusammenhang zwischen dem Skalarprodukt zweier Vektoren und ihrem Winkel zu untersuchen, betrachten wir zwei Vektoren $\vec{u}$ und $\vec{v}$, die weder kollinear noch orthogonal zueinander sind, sowie Pfeile $\overrightarrow{PU}$ und $\overrightarrow{PV}$ mit einem gemeinsamen Anfangspunkt P , die Repräsentanten von $\vec{u}$ bzw. $\vec{v}$ sind. Es ist unerheblich, ob $\vec{u}$ und $\vec{v}$ Vektoren der Ebene oder des Raumes sind, denn auch in letzterem Falle liegen die Pfeile $\overrightarrow{PU}$ und $\overrightarrow{PV}$ in einer Ebene (Abb. 3.39 a).

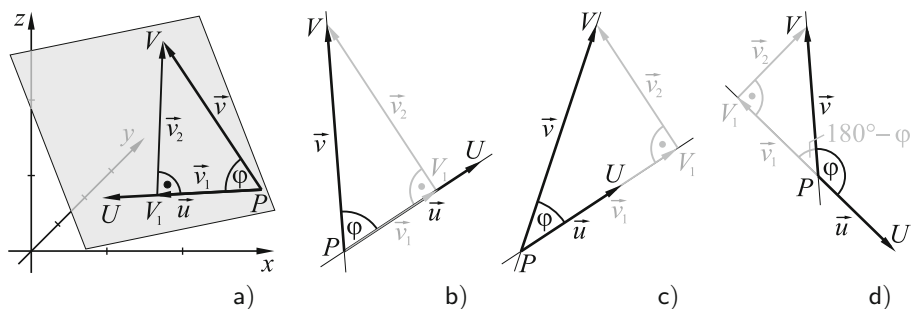


Abb. 3.39: Winkel zwischen zwei Vektoren im Raum und in der Ebene

Von dem Endpunkt V des zu $\vec{v}$ gehörigen Pfeils $\overrightarrow{PV}$ fallen wir das Lot auf die Gerade, die durch $\overrightarrow{PU}$ verläuft und bezeichnen den Lotfußpunkt mit V_1 . Der Vektor $\vec{v}$ wird dadurch in einen zu $\vec{u}$ kollinearen Vektor $\vec{v}_1 = \overrightarrow{PV_1}$ und einen zu $\vec{u}$ orthogonalen Vektor $\vec{v}_2 = \overrightarrow{V_1V}$ zerlegt. Es gilt $\vec{v} = \vec{v}_1 + \vec{v}_2$ und somit

$$\vec{u} \cdot \vec{v} = \vec{u} \cdot (\vec{v}_1 + \vec{v}_2) = \vec{u} \cdot \vec{v}_1 + \vec{u} \cdot \vec{v}_2.$$

Da $\vec{u}$ und $\vec{v}_2$ orthogonal sind, ist $\vec{u} \cdot \vec{v}_2 = 0$ und deshalb

$$\vec{u} \cdot \vec{v} = \vec{u} \cdot \vec{v}_1.$$

Die Vektoren $\vec{u}$ und $\vec{v}_1$ können, je nachdem, ob der Winkel φ zwischen den Vektoren $\vec{u}$ und $\vec{v}$ spitz oder stumpf ist, gleich oder entgegengesetzt orientiert sein (siehe Abb. 3.39 b, c bzw. d). Diese beiden Fälle werden nun untersucht.

1. Es seien $\vec{u}$ und $\vec{v}_1$ *gleich orientiert*. Dann ist wegen Satz 3.15 $\vec{u} \cdot \vec{v}_1 = |\vec{u}| \cdot |\vec{v}_1|$. In dem rechtwinkligen Dreieck ΔPV_1V (Abb. 3.39 b, c) gilt $\cos \varphi = \frac{|PV_1|}{|PV|}$ bzw. $|\vec{v}_1| = |\vec{v}| \cdot \cos \varphi$. Wegen $\vec{u} \cdot \vec{v} = \vec{u} \cdot \vec{v}_1 = |\vec{u}| \cdot |\vec{v}_1|$ folgt

$$\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}| \cdot \cos \varphi.$$

2. Falls $\vec{u}$ und $\vec{v}_1$ *entgegengesetzt orientiert* sind, so ist $\vec{u} \cdot \vec{v}_1 = -|\vec{u}| \cdot |\vec{v}_1|$. In dem Dreieck ΔPV_1V (siehe Abb. 3.39 d) gilt dann $\cos(180^\circ - \varphi) = \frac{|PV_1|}{|PV|}$

bzw. $|\vec{v}_1| = |\vec{v}| \cdot \cos(180^\circ - \varphi)$. Wegen $\cos(180^\circ - \varphi) = -\cos \varphi$ folgt daraus $|\vec{v}_1| = -|\vec{v}| \cdot \cos \varphi$. Da $\vec{u} \cdot \vec{v} = \vec{u} \cdot \vec{v}_1 = -|\vec{u}| \cdot |\vec{v}_1|$ ist, gilt auch in diesem Falle

$$\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}| \cdot \cos \varphi.$$

In beiden Fällen besteht zwischen dem Skalarprodukt zweier Vektoren $\vec{u}$ und $\vec{v}$ sowie dem Winkel zwischen ihnen die Beziehung $\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}| \cdot \cos \varphi$. Diese Beziehung gilt auch, wenn $\vec{u}$ und $\vec{v}$ kollinear oder orthogonal sind. Sind die Vektoren nämlich kollinear und gleich bzw. entgegengesetzt orientiert, so beträgt ihr Winkel 0° bzw. 180° und sein Kosinus ± 1 ; dies stimmt mit dem Satz 3.15 überein, wonach $\vec{u} \cdot \vec{v} = \pm |\vec{u}| \cdot |\vec{v}|$ gilt. Sind $\vec{u}$ und $\vec{v}$ orthogonal zueinander, so ist der Kosinus ihres Winkels Null und ihr Skalarprodukt ist nach dem Satz 3.16 ebenfalls Null. Es gilt also in jedem Falle der folgende Satz.

Satz 3.17

Für das Skalarprodukt zweier Vektoren $\vec{u}$ und $\vec{v}$, ihre Beträge $|\vec{u}|$, $|\vec{v}|$ und den Winkel $\varphi = \angle(\vec{u}, \vec{v})$ gilt

$$\vec{u} \cdot \vec{v} = |\vec{u}| \cdot |\vec{v}| \cdot \cos \varphi.$$

Der Satz 3.17 ermöglicht es, Winkel zwischen Vektoren zu berechnen, die als Paare bzw. Tripel reeller Zahlen gegeben sind.

Winkel zwischen zwei Vektoren

Sind $\vec{u} = \begin{pmatrix} x_u \\ y_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \end{pmatrix}$ zwei Vektoren der Ebene, so ist

$$\cos \angle(\vec{u}, \vec{v}) = \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|} = \frac{x_u x_v + y_u y_v}{\sqrt{x_u^2 + y_u^2} \cdot \sqrt{x_v^2 + y_v^2}}.$$

Sind $\vec{u} = \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}$ zwei Vektoren des Raumes, so ist

$$\cos \angle(\vec{u}, \vec{v}) = \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|} = \frac{x_u x_v + y_u y_v + z_u z_v}{\sqrt{x_u^2 + y_u^2 + z_u^2} \cdot \sqrt{x_v^2 + y_v^2 + z_v^2}}.$$

Beispiel 3.13

In einem Quader $ABCDEFGH$ mit den Seitenlängen $a = 4$ cm, $b = 3$ cm und $c = 2$ cm, soll der Winkel zwischen der Seite $\overline{AB}$ und der Raumdiagonalen $\overline{AG}$ berechnet werden. Dazu wird ein Koordinatensystem gewählt, in dem sich die Seiten und Diagonalen des Quaders leicht durch Pfeile und entsprechende Zahlentripel beschreiben lassen (siehe Abb. 3.40): $\overrightarrow{AB} = \begin{pmatrix} 4 \\ 0 \\ 0 \end{pmatrix}$, $\overrightarrow{AG} = \begin{pmatrix} 4 \\ 3 \\ 2 \end{pmatrix}$.

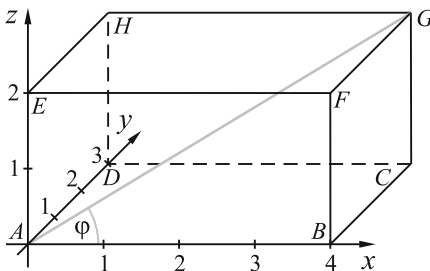


Abb. 3.40:

Berechnung eines Winkels in einem Quader

Für den Winkel zwischen diesen beiden Vektoren gilt:

$$\cos \angle (\overrightarrow{AB}, \overrightarrow{AG}) = \frac{\overrightarrow{AB} \cdot \overrightarrow{AG}}{|\overrightarrow{AB}| \cdot |\overrightarrow{AG}|} = \frac{4 \cdot 4 + 0 \cdot 3 + 0 \cdot 2}{\sqrt{4^2} \cdot \sqrt{4^2 + 3^2 + 2^2}} = \frac{16}{4 \cdot \sqrt{29}} \approx 0,743.$$

Es ergibt sich daraus für den gesuchten Winkel $\angle (\overrightarrow{AB}, \overrightarrow{AG}) \approx 42^\circ$. ■

3.5.3 Anwendungen des Skalarproduktes in der Geometrie und der Physik

Mithilfe des Skalarproduktes lassen sich recht leicht viele bekannte Sätze der Geometrie beweisen, in denen Aussagen über Winkel oder über Beziehungen zwischen Seitenlängen und Winkelgrößen getroffen werden.

Beispiel 3.14

Satz des Thales: Liegt ein Punkt C eines Dreiecks $\triangle ABC$ auf der Peripherie eines Kreises mit dem Durchmesser $\overline{AB}$, so ist das Dreieck bei C rechtwinklig.

Beweis: In vektorieller Schreibweise lautet die Behauptung des Satzes (mit den Bezeichnungen in Abb. 3.41) $\vec{a} \cdot \vec{b} = 0$. Die Vektoren $\vec{a}$ und $\vec{b}$ lassen sich durch $\vec{r}$ und $\vec{m}$ ausdrücken: $\vec{a} = -\vec{r} + \vec{m}$, $\vec{b} = \vec{r} + \vec{m}$. Somit ist

$$\vec{a} \cdot \vec{b} = (-\vec{r} + \vec{m}) \cdot (\vec{r} + \vec{m}) = \vec{m}^2 - \vec{r}^2.$$

Da $\vec{r}$ und $\vec{m}$ Radiusvektoren sind, somit also $|\vec{r}| = |\vec{m}| = r$ und deshalb $\vec{m}^2 - \vec{r}^2 = 0$ gilt, ist die Behauptung erfüllt. ■

Beispiel 3.15

Sinussatz: In jedem Dreieck ist der Quotient zweier Seitenlängen gleich dem Quotienten der Sinus der jeweils gegenüberliegenden Innenwinkel: $\frac{a}{b} = \frac{\sin \alpha}{\sin \beta}$.

Beweis: Durch Multiplikation der Vektorgleichung $\vec{a} + \vec{b} = \vec{c}$ mit dem Höhenvektor $\vec{h}$ (siehe Abb. 3.42) ergibt sich zunächst

$$\vec{a} \cdot \vec{h} + \vec{b} \cdot \vec{h} = \vec{c} \cdot \vec{h} = 0,$$

(da $\vec{h} \perp \vec{c}$). Wegen des Satzes 3.17 ist dies gleichbedeutend mit

$$|\vec{a}| \cdot |\vec{h}| \cdot \cos \gamma_2 + |\vec{b}| \cdot |\vec{h}| \cdot \cos (180^\circ - \gamma_1) = 0$$

bzw.

$$a \cdot \cos \gamma_2 - b \cdot \cos \gamma_1 = 0.$$

Wegen $\gamma_1 = 90^\circ - \alpha$ und $\gamma_2 = 90^\circ - \beta$ folgt daraus

$$a \cdot \cos (90^\circ - \beta) = b \cdot \cos (90^\circ - \alpha) \quad \text{bzw.} \quad a \cdot \sin \beta = b \cdot \sin \alpha$$

und somit die Behauptung $\frac{a}{b} = \frac{\sin \alpha}{\sin \beta}$. ■

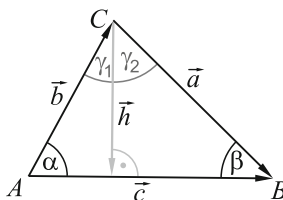
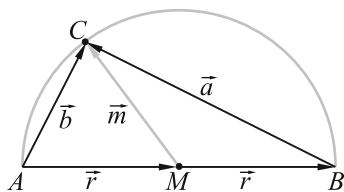


Abb. 3.41:
Satz des Thales
(links)

Abb. 3.42:
Sinussatz
(rechts)

Mechanische Arbeit

Im Physikunterricht der Sekundarstufe I wird die Hubarbeit als Produkt aus der Länge des zurückgelegten Weges und dem Betrag der in Richtung des Weges wirkenden Gewichtskraft berechnet: $W = |\vec{F}_G| \cdot |\vec{s}|$. Damit kann die Arbeit berechnet werden, wenn die zu überwindende Kraft $\vec{F}_G$ und der zurückgelegte Weg $\vec{s}$ gleich gerichtet sind. Haben die Gewichtskraft und der Weg wie bei einer geneigten Ebene verschiedene Richtungen, so muss der Winkel zwischen $\vec{F}_G$ und $\vec{s}$ berücksichtigt werden. Die Arbeit ist sehr gering wenn dieser Winkel nahezu ein rechter ist (z. B. bei einer sehr flach geneigten Ebene). Allgemein gilt:

$$W = \vec{F} \cdot \vec{s} = |\vec{F}_G| \cdot |\vec{s}| \cdot \cos \angle(\vec{F}_G, \vec{s}).$$

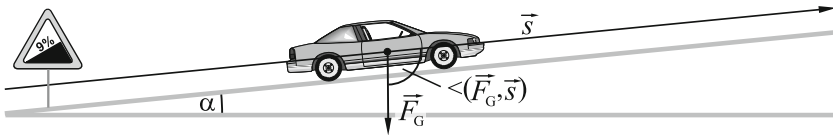


Abb. 3.43: Arbeit an der geneigten Ebene

Beispiel 3.16

Ein PKW der Masse 1300 kg (mit Insassen) fährt eine 750 Meter lange Straße mit dem konstanten Anstieg 9% hinauf. Es soll die dabei verrichtete Arbeit berechnet werden (ohne Berücksichtigung von Reibung und Luftwiderstand).

- Wir berechnen zunächst aus dem gegebenen Anstieg von 9% (d. h. 0,09) den Winkel α der geneigten Ebene: $\tan \alpha = 0,09$; es ergibt sich $\alpha \approx 5,14^\circ$.
- Der Winkel zwischen der Gewichtskraft $\vec{F}_G$ und $\vec{s}$ ergibt sich aus α durch $\angle(\vec{F}_G, \vec{s}) = 90^\circ + \alpha = 95,14^\circ$ (siehe Abb. 3.43).
- Mit $|\vec{F}_G| = m \cdot g \approx 1300 \text{ kg} \cdot 9,81 \frac{\text{m}}{\text{s}^2} \approx 12,75 \text{ kN}$ erhalten wir:
 $W = \vec{F}_G \cdot \vec{s} = |\vec{F}_G| \cdot |\vec{s}| \cdot \cos \angle(\vec{F}_G, \vec{s}) \approx 12,75 \text{ kN} \cdot 750 \text{ m} \cdot \cos 95,14^\circ \approx -857 \text{ kJ}.$
 Das negative Vorzeichen rührt daher, dass Arbeit verrichtet (Energie zugeführt) werden muss. Die zu verrichtende Arbeit beträgt somit 857 kJ. ■

3.5.4 Aufgaben zu Abschnitt 3.5

1. Überprüfen Sie, ob folgende Terme sinnvoll sind und – falls ja – ob sich als Ergebnisse Vektoren oder Zahlen ergeben.
 - a) $-\vec{a} + \vec{a} \cdot \vec{b}$ b) $\vec{a} + \vec{b}$ c) $\vec{a} + \vec{a} \cdot \vec{b}$ d) $\frac{\vec{a}}{\vec{b}}$ e) $\frac{a}{b}$ f) $\vec{a} \cdot (\vec{a} \cdot \vec{b})$
2. Berechnen Sie die Skalarprodukte folgender Paare von Vektoren.
 - a) $\vec{a} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}, \vec{b} = \begin{pmatrix} 1,5 \\ 3 \end{pmatrix}$ b) $\vec{x} = \begin{pmatrix} 1 \\ 2 \end{pmatrix}, \vec{y} = \begin{pmatrix} -2 \\ 1 \end{pmatrix}$ c) $\vec{u} = \frac{1}{4} \begin{pmatrix} 2 \\ -3 \\ 5 \end{pmatrix}, \vec{v} = \begin{pmatrix} -2 \\ 2 \\ 2 \end{pmatrix}$
3. Zeigen Sie mithilfe des Satzes 3.13 auf S. 123, dass für beliebige Vektoren $\vec{a}, \vec{b}, \vec{c}$ und $\vec{d}$ gilt: $(\vec{a} + \vec{b}) \cdot (\vec{c} + \vec{d}) = \vec{a} \cdot \vec{c} + \vec{b} \cdot \vec{c} + \vec{a} \cdot \vec{d} + \vec{b} \cdot \vec{d}.$

4. Finden Sie Beispiele dafür, dass die folgenden Regeln für die skalare Multiplikation von Vektoren *nicht* gelten.
- a) $(\vec{u} \cdot \vec{v}) \cdot \vec{w} = \vec{u} \cdot (\vec{v} \cdot \vec{w})$ b) $(\vec{u} \cdot \vec{v})^2 = \vec{u}^2 \cdot \vec{v}^2$
5. Für welche Parameter $t \in \mathbb{R}$ sind die folgenden Gleichungen erfüllt?
- a) $\begin{pmatrix} 2t \\ -7 \\ 2t^2 \end{pmatrix} \cdot \begin{pmatrix} -t \\ t \\ 1 \end{pmatrix} = 35$ b) $\begin{pmatrix} -3t \\ -7t \\ 4 \end{pmatrix} \cdot \begin{pmatrix} 2t \\ 2 \\ 5 \end{pmatrix} = -4t^2$ c) $\begin{pmatrix} -11 \\ 5t \\ 12t \end{pmatrix} \cdot \begin{pmatrix} t^2 \\ 3t \\ 4 \end{pmatrix} = -t^2 - 2t$
6. Vergleichen Sie für die in der Aufgabe 2 gegebenen Paare von Vektoren die Absolutbeträge $|\vec{a} \cdot \vec{b}|$, $|\vec{x} \cdot \vec{y}|$, $|\vec{u} \cdot \vec{v}|$ der Skalarprodukte jeweils mit den entsprechenden Produkten $|\vec{a}| \cdot |\vec{b}|$, $|\vec{x}| \cdot |\vec{y}|$, $|\vec{u}| \cdot |\vec{v}|$ der Beträge der Vektoren.
7. Begründen Sie, dass der Betrag eines Vektors genau dann Null ist, wenn es sich um den Nullvektor handelt.
8. Bestimmen Sie einen Vektor $\vec{n}$, der zu den Vektoren $\vec{u}$ und $\vec{v}$ orthogonal ist.
- a) $\vec{u} = \begin{pmatrix} 6 \\ 0 \\ -1 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 0 \\ 1 \\ 6 \end{pmatrix}$ b) $\vec{u} = \begin{pmatrix} 2 \\ -3 \\ -5 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 2 \\ 7 \\ 0 \end{pmatrix}$ c) $\vec{u} = \begin{pmatrix} 2 \\ 6 \\ 6 \end{pmatrix}$, $\vec{v} = \begin{pmatrix} 3 \\ -6 \\ -9 \end{pmatrix}$
9. Berechnen Sie für die in Abb. 3.44 dargestellte Pyramide
- a) die Kantenlängen,
b) die Höhen der Seitenflächen und
c) die Winkel $\angle(\vec{AC}, \vec{AS})$, $\angle(\vec{SA}, \vec{SB})$, $\angle(\vec{M_1M_2}, \vec{M_1S})$, $\angle(\vec{SM_1}, \vec{SM_2})$.
10. Beweisen Sie die folgenden Sätze (siehe Abb. 3.45):
- a) Wenn die Diagonalen eines Parallelogramms gleich lang sind, so ist es ein Rechteck.
b) In jedem Rechteck sind die Diagonalen gleich lang.
c) In jedem Rhombus (Raute) stehen die Diagonalen senkrecht aufeinander.
11. Welche Arbeit ist zu verrichten, wenn ein Bierfass der Masse 130 kg um 8 m eine 25° geneigte Ebene hinaufgerollt wird? Die Reibung wird vernachlässigt.
12. Ein Segelboot wird so gesteuert, dass der Wind mit einem Winkel von 25° zur Fahrtrichtung einfällt. Die Kraft, die der Wind auf das Boot ausübt, beträgt 1800 N. Wie groß ist die vom Wind nach einer Strecke von 2 km am Boot verrichtete Arbeit?

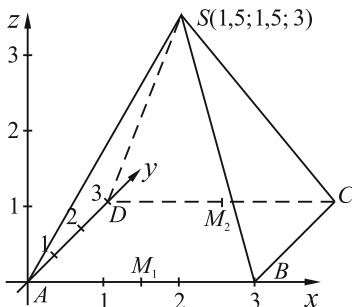


Abb. 3.44: Skizze zu Aufgabe 9 (links)

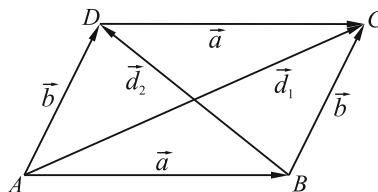


Abb. 3.45: Skizze zu Aufgabe 10

3.6 Vektor- und Spatprodukt

3.6.1 Das Vektorprodukt

Im Gegensatz zum Skalarprodukt kann das Vektorprodukt nur für Vektoren des Raumes bzw. in $\mathbb{R}^3$ gebildet werden. Durch dieses Produkt, mitunter auch als „Kreuzprodukt“ bezeichnet, wird zwei Vektoren ein Vektor zugeordnet.

Definition 3.11

Als Vektorprodukt zweier Vektoren $\vec{u}$ und $\vec{v}$ wird der Vektor $\vec{u} \times \vec{v}$ mit folgenden Eigenschaften bezeichnet:

1. $\vec{u} \times \vec{v}$ ist sowohl zu $\vec{u}$ als auch zu $\vec{v}$ orthogonal.
2. Der Betrag von $\vec{u} \times \vec{v}$ ist gleich dem Flächeninhalt des von $\vec{u}$ und $\vec{v}$ aufgespannten Parallelogramms (siehe Abb. 3.47), d. h.

$$|\vec{u} \times \vec{v}| = |\vec{u}| \cdot |\vec{v}| \cdot \sin \angle(\vec{u}, \vec{v}).$$
3. $\vec{u}$, $\vec{v}$ und $\vec{u} \times \vec{v}$ bilden ein *Rechtssystem*, d. h. ihre Orientierung entspricht der Orientierung der x -, y - und z -Achse eines kartesischen Koordinatensystems.

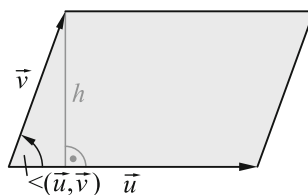
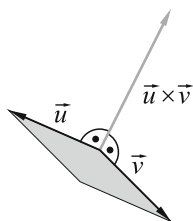
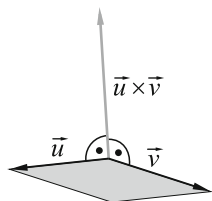


Abb. 3.46: Vektorprodukt

Abb. 3.47: Flächeninhalt eines Parallelogramms

Der Begriff der Orientierung kann an dieser Stelle noch nicht exakt definiert werden (siehe dazu den Abschnitt 6.3, speziell S. 255). Eine Veranschaulichung ist aber durch die *Rechte-Hand-Regel* möglich: Zeigt $\vec{u}$ in Richtung des Daumens und $\vec{v}$ in Richtung des Zeigefingers der rechten Hand, so zeigt $\vec{u} \times \vec{v}$ in die Richtung des Mittelfingers, siehe Abb. 3.48.

Vertauscht man bei der Bildung des Vektorprodukts die Reihenfolge, in der zwei Vektoren $\vec{u}$ und $\vec{v}$ multipliziert werden, so müssen die Vektoren $\vec{v}$, $\vec{u}$ und $\vec{v} \times \vec{u}$ in dieser Reihenfolge ein Rechtssystem bilden. Dies ist der Fall, wenn der Vektor $\vec{v} \times \vec{u}$ entgegengesetzt zu dem Vektor $\vec{u} \times \vec{v}$ orientiert ist, siehe Abb. 3.49. Es gilt also der folgende Satz.

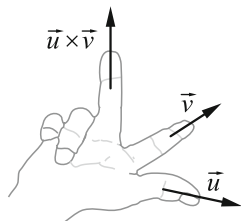


Abb. 3.48: Rechte-Hand-Regel

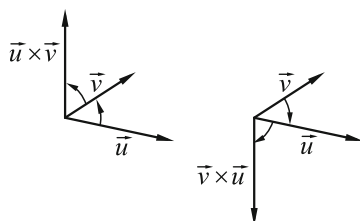
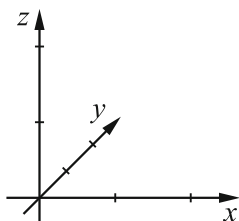


Abb. 3.49: „Anti-Kommutativität“

Satz 3.18

„Anti-Kommutativgesetz“: Für beliebige Vektoren $\vec{u}$ und $\vec{v}$ gilt $\vec{v} \times \vec{u} = -\vec{u} \times \vec{v}$.

In einigen Spezialfällen lassen sich recht leicht Aussagen über das Vektorprodukt zweier Vektoren machen:

Satz 3.19

1. Für beliebige Vektoren $\vec{u}$ gilt $\vec{u} \times \vec{o} = \vec{o} \times \vec{u} = \vec{o}$.
2. Für beliebige Vektoren $\vec{u}$ gilt $\vec{u} \times \vec{u} = \vec{o}$.
3. Sind zwei Vektoren $\vec{u}$ und $\vec{v}$ kollinear, so ist $\vec{u} \times \vec{v} = \vec{o}$.
4. Sind zwei Vektoren $\vec{u}$ und $\vec{v}$ orthogonal, so gilt $|\vec{u} \times \vec{v}| = |\vec{u}| \cdot |\vec{v}|$.

Alle vier Aussagen des Satzes 3.19 lassen sich recht leicht anhand der Definition 3.11 begründen (siehe die Aufgabe 1 auf S. 137).

Ohne Beweis sei vermerkt, dass für das Vektorprodukt die folgenden Rechenregeln gelten.

Satz 3.20

1. Für beliebige Vektoren $\vec{u}, \vec{v}, \vec{w}$ des Raumes gilt $\vec{u} \times (\vec{v} + \vec{w}) = \vec{u} \times \vec{v} + \vec{u} \times \vec{w}$ sowie $(\vec{v} + \vec{w}) \times \vec{u} = \vec{v} \times \vec{u} + \vec{w} \times \vec{u}$ (Distributivgesetze).
2. Für beliebige Vektoren $\vec{u}, \vec{v}$ des Raumes und beliebige $\lambda \in \mathbb{R}$ gilt $\lambda \cdot (\vec{u} \times \vec{v}) = (\lambda \cdot \vec{u}) \times \vec{v} = \vec{u} \times (\lambda \cdot \vec{v})$.

Bemerkungen:

- Aus dem Grunde, dass das Vektorprodukt nicht kommutativ ist, müssen zwei Distributivgesetze formuliert werden.
- Der Satz 3.20, 2 beinhaltet, dass es zu demselben Ergebnis führt, wenn zunächst $\vec{u}$ oder $\vec{v}$ mit einer reellen Zahl λ multipliziert oder wenn zuerst das Vektorprodukt gebildet und dieses mit λ multipliziert wird. Die Reihenfolge der Anordnung von λ und $\vec{u}$ kann also vertauscht werden, das bedeutet jedoch keine Kommutativität des Vektorprodukts (siehe hierzu den Satz 3.18).
- Da wegen Satz 3.20 die Klammern in $\lambda \cdot (\vec{u} \times \vec{v})$ unbedeutend sind, schreiben wir dafür im Folgenden einfach $\lambda \cdot \vec{u} \times \vec{v}$ oder $\lambda \vec{u} \times \vec{v}$.
- Ein Assoziativgesetz für das Vektorprodukt gilt nicht; im Allgemeinen ist $\vec{u} \times (\vec{v} \times \vec{w}) \neq (\vec{u} \times \vec{v}) \times \vec{w}$ (siehe Aufgabe 4 auf S. 137).

Komponentendarstellung des Vektorprodukts

Im Folgenden wird eine Möglichkeit hergeleitet, das Vektorprodukt zweier Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^3$ anhand ihrer Komponenten zu berechnen. Dazu betrachten wir zunächst die drei Einheitsvektoren in den Richtungen der Achsen eines kartesischen Koordinatensystems: $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ und $\vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$. Diese drei Vektoren haben jeweils den Betrag Eins und sind paarweise orthogonal zueinander. Ihre Vektorprodukte haben somit ebenfalls jeweils den Betrag Eins. Unter Berücksichtigung der Orientierung (siehe S. 133) ergibt sich:

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

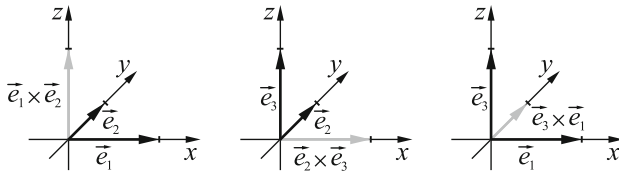


Abb. 3.50:
Vektorprodukte orthogonaler Einheitsvektoren

(siehe Abb. 3.50). Sind nun $\vec{u} = \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}$ beliebige Vektoren von $\mathbb{R}^3$, so lassen sie sich durch $\vec{e}_1$, $\vec{e}_2$ und $\vec{e}_3$ darstellen:

$$\vec{u} = x_u \cdot \vec{e}_1 + y_u \cdot \vec{e}_2 + z_u \cdot \vec{e}_3, \quad \vec{v} = x_v \cdot \vec{e}_1 + y_v \cdot \vec{e}_2 + z_v \cdot \vec{e}_3.$$

Unter Anwendung der Sätze 3.20, 3.19, 2 und 3.18 erhält man:

$$\begin{aligned} \vec{u} \times \vec{v} &= (x_u \cdot \vec{e}_1 + y_u \cdot \vec{e}_2 + z_u \cdot \vec{e}_3) \times (x_v \cdot \vec{e}_1 + y_v \cdot \vec{e}_2 + z_v \cdot \vec{e}_3) \\ &= x_u x_v \vec{e}_1 \times \vec{e}_1 + x_u y_v \vec{e}_1 \times \vec{e}_2 + x_u z_v \vec{e}_1 \times \vec{e}_3 \\ &\quad + y_u x_v \vec{e}_2 \times \vec{e}_1 + y_u y_v \vec{e}_2 \times \vec{e}_2 + y_u z_v \vec{e}_2 \times \vec{e}_3 \\ &\quad + z_u x_v \vec{e}_3 \times \vec{e}_1 + z_u y_v \vec{e}_3 \times \vec{e}_2 + z_u z_v \vec{e}_3 \times \vec{e}_3 \\ &= \vec{0} + x_u y_v \vec{e}_3 + x_u z_v (-\vec{e}_2) + y_u x_v (-\vec{e}_3) + \vec{0} + y_u z_v \vec{e}_1 \\ &\quad + z_u x_v \vec{e}_2 + z_u y_v (-\vec{e}_1) + \vec{0} \\ &= (y_u z_v - z_u y_v) \vec{e}_1 + (z_u x_v - x_u z_v) \vec{e}_2 + (x_u y_v - y_u x_v) \vec{e}_3. \end{aligned}$$

Schreibt man dieses Ergebnis in Komponentenschreibweise, so ergibt sich der folgende Satz.

Satz 3.21

Für beliebige Vektoren $\vec{u} = \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}$ gilt $\vec{u} \times \vec{v} = \begin{pmatrix} y_u z_v - z_u y_v \\ z_u x_v - x_u z_v \\ x_u y_v - y_u x_v \end{pmatrix}$.

Das Vektorprodukt kann verwendet werden, um auf recht einfache Weise einen Vektor zu bestimmen, der zu zwei gegebenen Vektoren orthogonal ist (siehe hierzu die Aufgaben 8 auf S. 132 und 3 auf S. 137). Außerdem lassen sich mithilfe der Definition 3.11 und des Satzes 3.21 Flächeninhalte von Parallelogrammen und Dreiecken im Raum ermitteln, deren Eckpunktkoordinaten gegeben sind.

Beispiel 3.17

Es soll der Flächeninhalt des Dreiecks $\triangle ABC$ mit $A(4; 5; 4)$, $B(5, 1, 2)$ und $C(0, 3, 3)$ bestimmt werden (siehe Abb. 3.51 auf S. 136).

Die Vektoren $\overrightarrow{AB} = \begin{pmatrix} 1 \\ -4 \\ -2 \end{pmatrix}$, $\overrightarrow{AC} = \begin{pmatrix} -4 \\ -2 \\ -1 \end{pmatrix}$ spannen ein Parallelogramm $ABCD$ auf. Wir berechnen zunächst das Vektorprodukt der Vektoren $\overrightarrow{AB}$ und $\overrightarrow{AC}$.

$$\overrightarrow{AB} \times \overrightarrow{AC} = \begin{pmatrix} 1 \\ -4 \\ -2 \end{pmatrix} \times \begin{pmatrix} -4 \\ -2 \\ -1 \end{pmatrix} = \begin{pmatrix} y_{AB} z_{AC} - z_{AB} y_{AC} \\ z_{AB} x_{AC} - x_{AB} z_{AC} \\ x_{AB} y_{AC} - y_{AB} x_{AC} \end{pmatrix} = \begin{pmatrix} 0 \\ 9 \\ -18 \end{pmatrix}.$$

Der Flächeninhalt des Dreiecks $\triangle ABC$ ist halb so groß wie der Flächeninhalt des Parallelogramms $ABCD$ und somit nach Definition 3.11 gleich der Hälfte des Betrags des Vektorprodukts $\overrightarrow{AB} \times \overrightarrow{AC}$:

$$A_{\triangle ABC} = \frac{1}{2} \cdot \left| \overrightarrow{AB} \times \overrightarrow{AC} \right| = \frac{1}{2} \sqrt{9^2 + 18^2} \approx 10,06. \quad \blacksquare$$

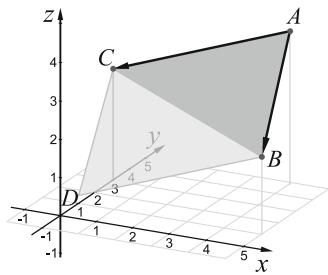


Abb. 3.51: Dreieck im Raum

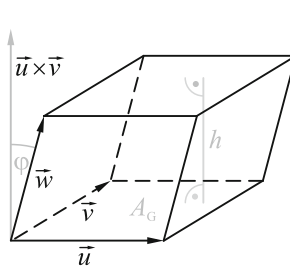


Abb. 3.52: Spatprodukt

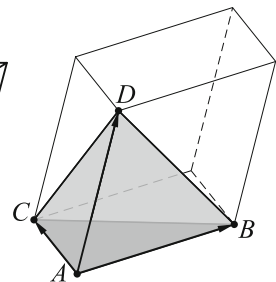


Abb. 3.53: Tetraeder

3.6.2 Das Spatprodukt – Berechnung von Volumina

Drei nicht komplanare Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ spannen ein Parallelepiped bzw. einen Spat auf, siehe Abb. 3.52. Alle Seitenflächen eines Spats sind Parallelogramme, die von je zwei der Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ aufgespannt werden. Das Volumen eines Spats ist das Produkt des Flächeninhalts A_G einer (als Grundfläche betrachteten) Seitenfläche mit der zugehörigen Höhe h . Für das von den Vektoren $\vec{u}$ und $\vec{v}$ aufgespannte Parallelogramm gilt nach Definition 3.11:

$$A_G = |\vec{u} \times \vec{v}|.$$

Die zugehörige Höhe lässt sich durch

$$h = |\vec{w}| \cdot \cos \varphi = |\vec{w}| \cdot |\cos \angle(\vec{u} \times \vec{v}, \vec{w})|$$

ausdrücken. Es muss der Betrag von $\cos \angle(\vec{u} \times \vec{v}, \vec{w})$ verwendet werden, da bei anderer Orientierung von $\vec{u} \times \vec{v}$ ein stumpfer Winkel entsteht. In diesem Falle ist mit $\varphi = 180^\circ - \angle(\vec{u} \times \vec{v}, \vec{w})$ zu rechnen, d. h. $\cos \varphi = -\cos \angle(\vec{u} \times \vec{v}, \vec{w})$.

Für das Volumen des Spats folgt aus den vorherigen Beziehungen:

$$V = A_G \cdot h = |\vec{u} \times \vec{v}| \cdot |\vec{w}| \cdot |\cos \angle(\vec{u} \times \vec{v}, \vec{w})| = |(\vec{u} \times \vec{v}) \cdot \vec{w}|.$$

Definition 3.12

Als Spatprodukt dreier Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ bezeichnet man $(\vec{u} \times \vec{v}) \cdot \vec{w}$. ◆

Das Volumen eines Spats ist daher der Absolutbetrag des Spatprodukts der drei aufspannenden Vektoren. Damit lassen sich auch Volumina von Pyramiden mit parallelogrammförmiger Grundfläche berechnen, diese betragen ein Drittel der Volumina von Prismen bzw. Spaten gleicher Grundfläche und gleicher Höhe.

Beispiel 3.18

Es soll das Volumen des Tetraeders mit den Eckpunkten $A(2; -4; 2)$, $B(4; 0; 2)$, $C(-1; 2; 1)$, $D(2; -1; 5)$ bestimmt werden (Abb. 3.53). Dazu berechnet man das Volumen des Spats, das von den Vektoren $\overrightarrow{AB}$, $\overrightarrow{AC}$ und $\overrightarrow{AD}$ aufgespannt wird:

$$V_P = \left| \left(\overrightarrow{AB} \times \overrightarrow{AC} \right) \cdot \overrightarrow{AD} \right| = \left| \left(\begin{pmatrix} 2 \\ 4 \\ 0 \end{pmatrix} \times \begin{pmatrix} -3 \\ 6 \\ -1 \end{pmatrix} \right) \cdot \begin{pmatrix} 0 \\ 3 \\ 3 \end{pmatrix} \right| = \left| \begin{pmatrix} -4 \\ 2 \\ 24 \end{pmatrix} \cdot \begin{pmatrix} 0 \\ 3 \\ 3 \end{pmatrix} \right| = 78.$$

Das Volumen der vierseitigen Pyramide mit dem von $\overrightarrow{AB}$ und $\overrightarrow{AC}$ aufgespannten Parallelogramm als Grundfläche und der Spitze D beträgt ein Drittel dieses Wertes. Da die Grundfläche des Tetraeders $ABCD$ jedoch demgegenüber nur halb so groß ist, hat dieses das Volumen $\frac{1}{6} \cdot 78 = 13$. ■

3.6.3 Aufgaben zu Abschnitt 3.6

1. Begründen Sie die Aussagen des Satzes 3.19 auf S. 134.
2. Beweisen Sie den folgenden Zusammenhang zwischen dem Vektor- und dem Skalarprodukt für beliebige Vektoren $\vec{u}$ und $\vec{v}$: $|\vec{u} \times \vec{v}|^2 = |\vec{u}|^2 \cdot |\vec{v}|^2 - (\vec{u} \cdot \vec{v})^2$.
3. Berechnen Sie die Vektorprodukte der in der Aufgabe 8 auf S. 132 gegebenen Paare von Vektoren.
4. Begründen Sie anhand eines Gegenbeispiels, dass für das Vektorprodukt kein Assoziativgesetz der Form $\vec{u} \times (\vec{v} \times \vec{w}) = (\vec{u} \times \vec{v}) \times \vec{w}$ gilt.
5. Begründen Sie, dass das Viereck $ABCD$ mit $A(3; 2; -4)$, $B(8; -2; -5)$, $C(2; 5; -6)$ und $D(7; 1; -7)$ ein Parallelogramm ist und berechnen Sie seinen Flächeninhalt.
6. Berechnen Sie das Volumen der Pyramide mit dem Parallelogramm $ABCD$ aus der Aufgabe 5 als Grundfläche und der Spitze $S(-8; 9; 10)$.
7. Begründen Sie mittels geometrischer Überlegungen:

$$|(\vec{u} \times \vec{v}) \cdot \vec{w}| = |(\vec{v} \times \vec{w}) \cdot \vec{u}| = |(\vec{w} \times \vec{u}) \cdot \vec{v}|.$$

3.7 Vektorrechnung und -darstellung mithilfe eines Computeralgebrasystems

Vektoren werden in dem CAS Maxima (siehe Abschnitt 1.5) als n -Tupel in eckigen Klammern eingegeben, z. B. $[x, y]$ und $[x, y, z]$. Diese lassen sich addieren und mit reellen Zahlen multiplizieren. So gibt Maxima z. B. auf die Eingaben

```
u:[3,2,12]$ v:[-8,12.5,13]$ lambda:4/3$
u+v; lambda*u; u+lambda*v;
```

die Ergebnisse $[-5, 14.5, 25]$, $[4, \frac{8}{3}, 16]$ und $[-\frac{23}{3}, 18.666, \frac{88}{3}]$ aus.

Graphische Darstellungen von Vektoren als Pfeile

Mithilfe von `vector([x0,y0],[x,y])` bzw. `vector([x0,y0,z0],[x,y,z])` lässt sich ein Vektor $\begin{pmatrix} x \\ y \end{pmatrix}$ der Ebene bzw. $\begin{pmatrix} x \\ y \\ z \end{pmatrix}$ des Raumes als Pfeil mit dem Anfangspunkt $P_0(x_0; y_0)$ bzw. $P_0(x_0; y_0; z_0)$ darstellen. Dazu werden die auf S. 47/48 beschriebenen Umgebungen `draw2d` bzw. `draw3d` genutzt. Beispiele für die Darstellung von Vektoren in der Ebene und im Raum enthält die Internetseite zu diesem Buch. Mithilfe der folgenden Eingaben wurde Abb. 3.54 generiert:

```
load(draw)$
o:[0,0,0]$ u:[1,3,-1]$ v:[-3,-1,2]$ w:[-1,3,1]$
draw3d( user_preamble = ["set size ratio 1"], grid=true,
        xyplane=0, xaxis=true, yaxis=true, zaxis=true,
        xrange = [-6,6], yrange = [-6,6], zrange = [-4,4],
        color = black, vector(o,u), vector(o,v), vector(o,w),
        color = red, vector(u,u+v), vector(v,u+v), vector(w,u+v)
)
```

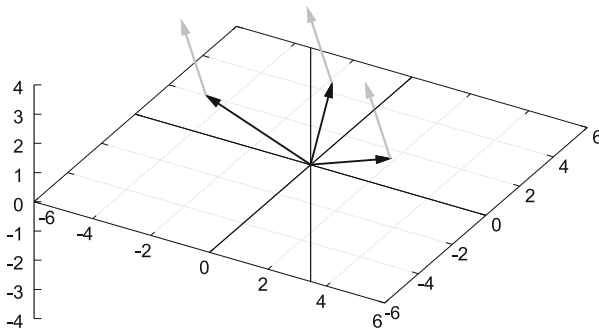


Abb. 3.54:
Darstellung von Vektoren
im Raum mithilfe des CAS
Maxima

Die Werte für **xrange**, **yrange**, **zrange** sind anzupassen, wenn (in Abhängigkeit von den Komponenten der zu visualisierenden Vektoren) andere Ausschnitte des Raumes dargestellt werden sollen. Die Anschaulichkeit räumlicher Darstellungen ist insbesondere dadurch gegeben, dass diese mit der Maus „gedreht“ und aus verschiedenen Richtungen betrachtet werden können.

Linearkombinationen

Um Vektoren als Linearkombinationen anderer Vektoren darzustellen, sind lineare Gleichungssysteme zu lösen, vgl. Abschnitt 3.4 (zum Lösen von LGS mit Maxima siehe Abschnitt 1.5). Um die Komponenten von Vektoren nicht einzeln eingeben zu müssen, kann in Maxima auf Komponenten in der Form $[x,y]$ bzw. $[x,y,z]$ eingegebener Vektoren zurückgegriffen werden, z. B. mittels $v[2]$ auf die zweite (y -) Komponente eines zuvor angegebenen Vektors $\vec{v}$. Dies wird in dem folgenden Beispiel genutzt, um die Koeffizienten bei der Darstellung des Vektors $\vec{x}$ als Linearkombination der Vektoren $\vec{u}$, $\vec{v}$ und $\vec{w}$ zu ermitteln.

```
x:[1,3,2]$   u:[-1,3,-4]$   v:[3,-5,-2]$   w:[7,-9,-4]$
solve([ lambda * u[1] + mu * v[1] + nu * w[1]= x[1],
        lambda * u[2] + mu * v[2] + nu * w[2]= x[2],
        lambda * u[3] + mu * v[3] + nu * w[3]= x[3] ],
      [lambda, mu, nu] );
```

Als Ausgabe erhält man: $[[\text{lambda}=-3/10, \text{mu}=-21/5, \text{nu}=19/10]]$.

Skalarprodukt

Für das Skalarprodukt zweier Vektoren wird ein Punkt $.$ (dot) verwendet.

```
Eingabe: u:[3,4,-1]$   v:[-12,5,7]$           Ausgabe: -23
          u.v;
```

Vektorprodukt

Das Vektorprodukt zweier Vektoren $\vec{u}$ und $\vec{v}$ wird mittels **express**($u \sim v$) berechnet. Dazu muss vorher mittels **load**("vect") das Paket „vect“ geladen werden.

```
Eingabe: load("vect")$           Ausgabe: [0,-10,10]
          u:[1,2,2]$   v:[-4,2,2]$
          express(u~v);
```

Die hier beschriebenen sowie weitere Beispiele enthält eine Maxima-Datei, die auf der Internetseite zu diesem Buch zur Verfügung steht.

4 Vektorielle Raumgeometrie

Übersicht

4.1	Parameterdarstellungen von Geraden und Kurven	140
4.2	Ebenen	151
4.3	Metrische Geometrie von Geraden und Ebenen.....	156

Die in dem Kapitel 2 zunächst nur mithilfe von Koordinaten behandelte analytische Geometrie wird im Folgenden mit den mächtigen Hilfsmitteln der Vektorrechnung weitergeführt. Diese ermöglichen es, Aufgaben der Raumgeometrie zu bearbeiten, die mit reinen Koordinatenmethoden nicht oder nur schwer lösbar sind. In vielen Fällen lassen sich unter Anwendung vektorieller Methoden geometrische Objekte und Beziehungen im dreidimensionalen Raum auf analoge Weise zu Sachverhalten der ebenen Geometrie untersuchen. Dabei wird auch auf Inhalte des Kapitels 1 zurückgegriffen, in dem bereits Geraden und Ebenen als Lösungsmengen linearer Gleichungen bzw. Gleichungssysteme auftraten.

In den ersten beiden Abschnitten dieses Kapitels wird auf der Grundlage der Vektorrechnung zunächst *affine Geometrie* von Geraden und Ebenen betrieben, d. h. es werden Lagebeziehungen untersucht, ohne dabei Maße zu verwenden. Diese sind Gegenstand des Abschnitts 4.3 zur *metrischen Geometrie*, in dessen Mittelpunkt die Berechnung von Abständen und Winkelmaßen im Raum unter Nutzung des Skalarproduktes steht. Da es sich hierbei um Standardinhalte des Mathematikunterrichts der gymnasialen Oberstufe handelt, wird eine recht knappe Darstellung gegeben; eine Vielzahl von Aufgaben zu diesem Abschnitt dient u. a. dazu, Schulwissen zu reaktivieren.

Bei der vektoriellen Beschreibung von Geraden und Ebenen sind *Parameterdarstellungen* von zentraler Bedeutung, die im Mittelpunkt der ersten beiden Abschnitte dieses Kapitels stehen. Parameterdarstellungen ermöglichen aber nicht nur die Beschreibung linearer Gebilde, sondern auch interessanter Kurven (sowie auch Flächen); hierauf wird ansatzweise in den Abschnitten 4.1.4 und 4.1.5 eingegangen.

4.1 Parameterdarstellungen von Geraden und Kurven

4.1.1 Beschreibung von Geraden durch Parametergleichungen

In dem Abschnitt 1.1.1 stellte sich bereits heraus, dass Lösungsmengen linearer Gleichungen mit zwei Variablen im Allgemeinen durch Geraden der Ebene dargestellt und in der Form

$$\begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} m_1 \\ m_2 \end{pmatrix} + \begin{pmatrix} n_1 \\ n_2 \end{pmatrix}$$

mit $t \in \mathbb{R}$ beschrieben werden können (siehe S.4). Mithilfe der Vektorrechnung lassen sich Geraden in dieser Darstellung nun näher untersuchen, und vor allem lässt sich diese Art, Geraden zu beschreiben, leicht auf den Raum übertragen.

Beispiel 4.1

Gegeben sind der Vektor $\vec{p}_0 = \begin{pmatrix} 1 \\ -1 \\ 3 \end{pmatrix}$ und der Vektor $\vec{a} = \begin{pmatrix} -2,5 \\ 1 \\ -1,5 \end{pmatrix}$. Es werden die Vektoren $\vec{p}_{0,5} = \vec{p}_0 + 0,5\vec{a}$, $\vec{p}_1 = \vec{p}_0 + 1\vec{a}$, $\vec{p}_{1,5} = \vec{p}_0 + 1,5\vec{a}$, $\vec{p}_2 = \vec{p}_0 + 2\vec{a}$, $\vec{p}_{-0,5} = \vec{p}_0 - 0,5\vec{a}$, $\vec{p}_{-1} = \vec{p}_0 - 1\vec{a}$, $\vec{p}_{-1,5} = \vec{p}_0 - 1,5\vec{a}$ und $\vec{p}_{-2} = \vec{p}_0 - 2\vec{a}$ berechnet. In ein Koordinatensystem werden eingezeichnet:

- der Pfeil $\overrightarrow{OP_0}$ als Repräsentant des Vektors $\vec{p}_0$ mit dem Anfangspunkt im Koordinatenursprung,
- ein Pfeil mit dem Anfangspunkt P_0 , der den Vektor $\vec{a}$ repräsentiert,
- Pfeile $\overrightarrow{OP_{-2}}, \overrightarrow{OP_{-1,5}}, \dots, \overrightarrow{OP_2}$ als Repräsentanten der Vektoren $\vec{p}_{-2}, \dots, \vec{p}_2$.

Die Endpunkte $P_{-2}, P_{-1,5}, \dots, P_2$ der entsprechenden Pfeile werden markiert, siehe Abb. 4.1 (links). Offensichtlich liegen alle diese Punkte auf einer Geraden. Durch die Darstellung einer größeren Anzahl von Punkten nach derselben Vorgehensweise wie in Abb. 4.1 (rechts) wird dies noch deutlicher sichtbar. ■

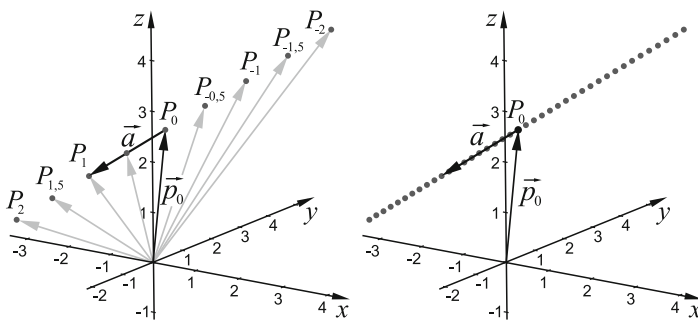


Abb. 4.1:
Punkte einer durch eine Parameterdarstellung gegebenen Geraden

Interaktive Versionen dieser Abbildungen, die aus verschiedenen Richtungen betrachtet werden können, enthält die Internetseite zu diesem Buch.

Alle Endpunkte von Pfeilen $\overrightarrow{OP}$ der Ebene oder des Raumes, welche Vektoren repräsentieren, die sich in der Form $\vec{p} = \vec{p}_0 + t\vec{a}$ (mit festen Vektoren $\vec{p}_0$ und $\vec{a}$ sowie $t \in \mathbb{R}$) darstellen lassen, liegen auf einer Geraden. Umgekehrt lassen sich alle Punkte P der Geraden, die durch einen Punkt P_0 verläuft und deren Richtung durch einen Vektor $\vec{a}$ gegeben ist, in der Form $\vec{p} = \vec{p}_0 + t\vec{a}$ (mit

$t \in \mathbb{R}$) beschreiben, wobei $\vec{p}$ der durch den Pfeil $\overrightarrow{OP}$ und $\vec{p}_0$ der durch $\overrightarrow{OP_0}$ repräsentierte Vektor ist. Dabei werden folgende Bezeichnungen verwendet:

- Ein Vektor $\vec{p}$ heißt *Ortsvektor* eines Punktes P , falls $\vec{p}$ einen repräsentierenden Pfeil $\overrightarrow{OP}$ besitzt, dessen Anfangspunkt der Koordinatenursprung und dessen Endpunkt P ist.
- Ein bekannter (fester) Punkt P_0 einer Geraden g heißt *Aufpunkt* von g . Es kann jeder Punkt einer Geraden als Aufpunkt verwendet werden.
- Der Ortsvektor eines Aufpunktes einer Geraden wird als *Stützvektor* dieser Geraden bezeichnet.
- Ein Vektor $\vec{a}$, der die Richtung einer Geraden g beschreibt, heißt *Richtungsvektor* von g .
- Die Variable t in der Darstellung $\vec{p} = \vec{p}_0 + t\vec{a}$ einer Geraden heißt *Parameter*.

Parameterdarstellungen von Geraden

Jede Gerade g in der Ebene oder im Raum lässt sich durch eine Gleichung der Form

$$\vec{p} = \vec{p}_0 + t\vec{a} \quad (4.1)$$

beschreiben, die als *Parameterdarstellung* oder *Parametergleichung* von g bezeichnet wird. Dabei ist $\vec{a}$ ein Richtungsvektor und $\vec{p}_0$ ein Stützvektor der Geraden g . Der Parameter t durchläuft den gesamten Bereich der reellen Zahlen, wobei jeder Zahl t ein Punkt P der Geraden mit dem Ortsvektor $\vec{p} = \vec{p}_0 + t\vec{a}$ zugeordnet ist.

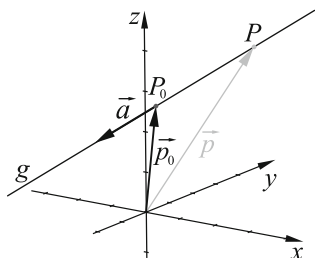


Abb. 4.2: Gerade mit Pfeilen, die einen Stützvektor und einen Richtungsvektor repräsentieren, sowie einem (willkürlich gewählten) Punkt P

Während eine Parameterdarstellung der Form (4.1) stets genau eine Gerade beschreibt, kann ein und dieselbe Gerade durch verschiedene Parameterdarstellungen beschrieben werden.

Beispiel 4.2

Durch die Parameterdarstellungen $\vec{p} = \vec{p}_0 + t\vec{a}$ und $\vec{p} = \vec{q}_0 + s\vec{b}$ mit

$$\vec{p}_0 = \begin{pmatrix} 4 \\ -2 \\ -2 \end{pmatrix}, \vec{q}_0 = \begin{pmatrix} -4 \\ 6 \\ 2 \end{pmatrix}, \vec{a} = \begin{pmatrix} -2 \\ 2 \\ 1 \end{pmatrix} \text{ und } \vec{b} = \begin{pmatrix} 3 \\ -3 \\ -1,5 \end{pmatrix}$$

wird dieselbe Gerade beschrieben, siehe Abb. 4.3.

Die Vektoren $\vec{a}$ und $\vec{b}$ sind kollinear, denn es gilt $\vec{b} = -1,5 \cdot \vec{a}$. Außerdem liegt der Punkt Q_0 auf der durch $\vec{p} = \vec{p}_0 + t\vec{a}$ beschriebenen Geraden. Dies lässt sich mittels einer so genannten *Punktprobe* überprüfen, indem $\vec{q}_0$ in diese Parameterdarstellung eingesetzt wird:

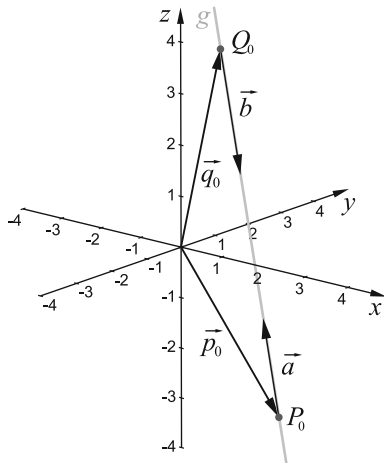


Abb. 4.3: Zwei Stütz- und zwei Richtungsvektoren einer Geraden

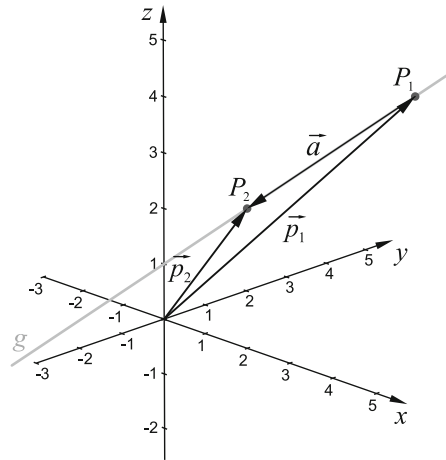


Abb. 4.4: Gerade durch zwei gegebene Punkte

$$\begin{pmatrix} -4 \\ 6 \\ 2 \end{pmatrix} = \begin{pmatrix} 4 \\ -2 \\ -2 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ 2 \\ 1 \end{pmatrix}.$$

Mit $t=4$ ist diese Vektorgleichung offensichtlich erfüllt. Somit liegt Q_0 auf der durch die Parameterdarstellung $\vec{p} = \vec{p}_0 + t\vec{a}$ beschriebenen Geraden (es ließe sich analog zeigen, dass P_0 auf der durch $\vec{p} = \vec{q}_0 + s\vec{b}$ erzeugten Geraden liegt, dies ist aber nicht mehr notwendig). Da außerdem die Richtungsvektoren $\vec{a}$ und $\vec{b}$ kollinear sind, ergibt sich durch beide Parameterdarstellungen dieselbe Gerade. Diese Gerade wird zudem auch durch die Parameterdarstellungen $\vec{p} = \vec{q}_0 + u\vec{a}$ und $\vec{p} = \vec{p}_0 + v\vec{b}$ mit $u, v \in \mathbb{R}$ beschrieben. ■

Aufstellen einer Parameterdarstellung einer Geraden anhand zweier Punkte

Beispiel 4.3

Gegeben sind zwei Punkte $P_1(3; 3; 4)$ und $P_2(1; 1; 2)$, gesucht ist eine Parameterdarstellung der durch P_1 und P_2 verlaufenden Geraden g , siehe Abb. 4.4. Ein Richtungsvektor $\vec{a}$ von g ist der Vektor, der sich als Differenz der Ortsvektoren der Punkte P_1 und P_2 ergibt: $\vec{a} = \vec{p}_2 - \vec{p}_1 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} - \begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix} = \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix}$. Als Stützvektor kann $\vec{p}_2$ oder $\vec{p}_1$ verwendet werden, Parameterdarstellungen von g sind somit z. B. $\vec{p} = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ -2 \\ -2 \end{pmatrix}$ und $\vec{p} = \begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$. ■

4.1.2 Parameter- und Koordinatengleichungen bzw. LGS

In dem Abschnitt 1.1.1 wurde bereits festgestellt, dass Lösungsmengen von linearen Gleichungen mit zwei Variablen in den meisten Fällen Geraden in der Ebene beschreiben, und es wurden diese Lösungsmengen als Parameterdarstellungen angegeben, siehe das Beispiel 1.2 auf S. 4 sowie die Zusammenfassung auf S. 5. Umgekehrt lassen sich auch Geraden der Ebene, die durch Parametergleichungen gegeben sind, durch Gleichungen der Form $ax + by = c$ beschreiben.

Beispiel 4.4

Eine Gerade ist durch die Parameterdarstellung

$$\vec{p} = \begin{pmatrix} x \\ y \end{pmatrix} = t \cdot \begin{pmatrix} 3 \\ -8 \end{pmatrix} + \begin{pmatrix} -5 \\ 3 \end{pmatrix}$$

gegeben. Diese lässt sich auch in Form von zwei Gleichungen schreiben:

$$\begin{aligned} x &= 3t - 5 \\ y &= -8t + 3. \end{aligned}$$

Aus diesem LGS wird nun der Parameter t eliminiert. Multipliziert man die erste Gleichung mit 8, die zweite Gleichung mit 3 und addiert anschließend beide Gleichungen, so ergibt sich

$$8x + 3y = -31$$

als Koordinatengleichung der gegebenen Geraden. ■

Geraden im Raum lassen sich nicht durch einzelne lineare Gleichungen darstellen, da deren Lösungsmengen Ebenen beschreiben oder leer sind, siehe Abschnitt 1.2.1. Hingegen können Lösungsmengen von linearen Gleichungssystemen mit drei Variablen und zwei Gleichungen, die den Rang 2 haben, geometrisch als Geraden interpretiert werden, siehe das Beispiel 1.16 a auf S. 19. Für dieses Beispiel wird im Folgenden eine Parametergleichung aufgestellt.

Beispiel 4.5

Eine Parametergleichung der Geraden, welche die Lösungsmenge des LGS

$$\begin{aligned} \frac{3}{2}x - 2y + 4z &= 1 \\ x + y - 3z &= 2 \end{aligned}$$

beschreibt, erhält man durch Lösung dieses LGS (z. B. mithilfe des Gauß-Algorithmus, siehe Beispiel 1.23 auf S. 27). Es ergibt sich eine Lösungsmenge mit einem Parameter:

$$\begin{aligned} x &= \frac{4}{7}t + \frac{10}{7} \\ y &= \frac{17}{7}t + \frac{4}{7} \quad (t \in \mathbb{R}). \\ z &= t \end{aligned}$$

Eine Parameterdarstellung der durch das LGS beschriebenen Geraden ist also

$$\vec{p} = t \cdot \begin{pmatrix} \frac{4}{7} \\ \frac{17}{7} \\ 1 \end{pmatrix} + \begin{pmatrix} \frac{10}{7} \\ \frac{4}{7} \\ 0 \end{pmatrix}. \quad \blacksquare$$

In Umkehrung zu dem Beispiel 4.5 lassen sich Geraden des Raumes, die durch Parameterdarstellungen gegeben sind, auch durch lineare Gleichungssysteme mit zwei Gleichungen beschreiben.

Beispiel 4.6

Die Gerade aus dem Beispiel 4.2 (Abb. 4.3) wird durch die Parametergleichung

$$\vec{p} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 4 \\ -2 \\ -2 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ 2 \\ 1 \end{pmatrix}$$

beschrieben, die sich auch in einzelnen Gleichungen schreiben lässt:

$$\begin{aligned} x &= 4 - 2t \\ y &= -2 + 2t \\ z &= -2 + t. \end{aligned}$$

Aus zwei Paaren dieser Gleichungen (z. B. aus der ersten und der zweiten sowie der zweiten und der dritten Gleichung) wird der Parameter t eliminiert (wie in Beispiel 4.4 beschrieben). Es ergibt sich das LGS

$$\begin{array}{rcl} x + y & = & 2 \\ -y + 2z & = & -2 \end{array}$$

als Beschreibung der Geraden. ■

4.1.3 Lagebeziehungen und Schnittpunkte von Geraden

Nachdem in dem Abschnitt 1.1.2 bereits Lagebeziehungen von Geraden in der Ebene untersucht und (falls vorhanden) ihre Schnittpunkte bestimmt wurden, ist dies nun anhand von Parameterdarstellungen auch für Geraden im Raum möglich. Zunächst ist leicht zu überlegen, dass zwei Geraden parallel oder sogar identisch sind, falls ihre Richtungsvektoren kollinear sind.

Beispiel 4.7

Die Richtungsvektoren der drei Geraden g , h und k mit

$g: \vec{p} = \vec{p}_0 + t\vec{a} = \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix} + t \cdot \begin{pmatrix} -1 \\ -1 \\ -0,5 \end{pmatrix}$, $h: \vec{p} = \vec{q}_0 + s\vec{b} = \begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix} + s \cdot \begin{pmatrix} -2 \\ -2 \\ -1 \end{pmatrix}$ und $k: \vec{p} = \vec{r}_0 + r\vec{c} = \begin{pmatrix} 4 \\ -3 \\ 3 \end{pmatrix} + r \cdot \begin{pmatrix} 2 \\ 2 \\ 1 \end{pmatrix}$ sind jeweils kollinear: $\vec{c} = -\vec{b} = -2\vec{a}$. Die Geraden g , h und k sind daher jeweils zueinander parallel oder sogar identisch.

Um festzustellen, ob g und h identisch sind, reicht es aus, zu prüfen, ob ein einziger Punkt von g auf h liegt oder umgekehrt, ob also z. B. $t \in \mathbb{R}$ existiert mit $\vec{q}_0 = \vec{p}_0 + t\vec{a}$. Dies ist der Fall, denn mit $t = -4$ ist $\begin{pmatrix} 3 \\ 3 \\ 4 \end{pmatrix} = \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix} - 4 \cdot \begin{pmatrix} -1 \\ -1 \\ -0,5 \end{pmatrix}$.

Die Geraden g und h sind somit identisch. Zu einem anderen Ergebnis führt die Überprüfung, ob g und k identisch sind. Dazu müsste ein $t \in \mathbb{R}$ mit $\vec{r}_0 = \vec{p}_0 + t\vec{a}$ existieren, also mit $\begin{pmatrix} 4 \\ -3 \\ 3 \end{pmatrix} = \begin{pmatrix} -1 \\ -1 \\ 2 \end{pmatrix} + t \cdot \begin{pmatrix} -1 \\ -1 \\ -0,5 \end{pmatrix}$. Damit diese Vektorgleichung für die erste Komponente erfüllt ist, müsste $t = -5$ sein, für die zweite Komponente hingegen $t = 2$. Somit erfüllt keine reelle Zahl t diese Vektorgleichung; die Gerade k ist nicht identisch mit g und h , sondern lediglich parallel dazu. ■

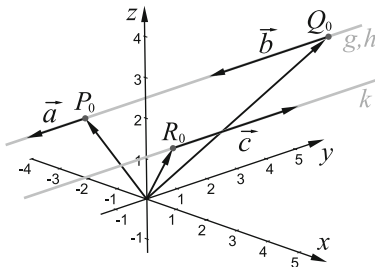


Abb. 4.5: Parallele und identische Geraden

Eine interaktive Version dieser Abbildung, die aus verschiedenen Richtungen betrachtet werden kann, enthält die Internetseite zu diesem Buch.

Im Folgenden werden Geraden mit nicht kollinearen Richtungsvektoren untersucht. Dabei sind die Fälle möglich, dass sich zwei Geraden schneiden oder dass sie keinen Schnittpunkt besitzen, man spricht dann von windschiefen Geraden.

Beispiel 4.8

Die Geraden g und h (siehe Abb. 4.6) mit den Parameterdarstellungen

$$g: \vec{p} = \vec{p}_0 + t \vec{a} = \begin{pmatrix} 4 \\ 2 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ -1 \\ 0,5 \end{pmatrix}, \quad h: \vec{p} = \vec{q}_0 + s \vec{b} = \begin{pmatrix} -2 \\ 4 \\ 1 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ -3 \\ 1 \end{pmatrix}$$

sind nicht parallel oder identisch, denn die Richtungsvektoren $\vec{a}$ und $\vec{b}$ sind nicht kollinear. Wenn es einen gemeinsamen Punkt P beider Geraden gäbe, so müssten Parameter t und s existieren, mit denen sich der Ortsvektor $\vec{p}$ von P aus den Parameterdarstellungen beider Geraden ergibt. Um zu prüfen, ob es derartige Parameter t, s gibt, werden die Parameterdarstellungen gleichgesetzt:

$$\begin{pmatrix} 4 \\ 2 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ -1 \\ 0,5 \end{pmatrix} = \begin{pmatrix} -2 \\ 4 \\ 1 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ -3 \\ 1 \end{pmatrix}$$

Dieser Vektorgleichung entspricht das lineare Gleichungssystem

$$\begin{aligned} -2t + s &= -6 \\ -t + 3s &= 2 \\ 0,5t - s &= 0, \end{aligned}$$

das die Lösung $t=4, s=2$ hat. Die beiden Geraden besitzen somit einen Schnittpunkt P , dessen Ortsvektor $\vec{p}$ sich durch Einsetzen der ermittelten Werte für s bzw. t in eine der beiden Parameterdarstellungen ergibt: $\vec{p} = \begin{pmatrix} -4 \\ -2 \\ 3 \end{pmatrix}$. ■

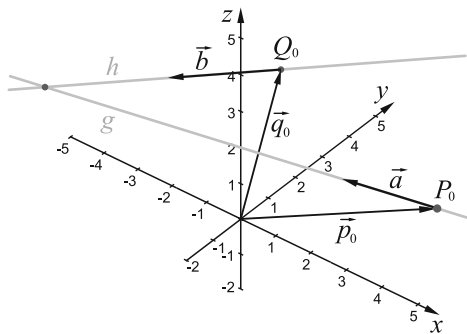


Abb. 4.6: Sich schneidende Geraden

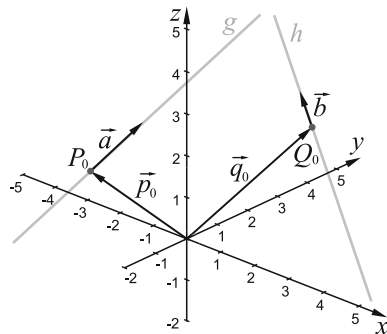


Abb. 4.7: Windschiefe Geraden

Interaktive Versionen dieser Abbildungen enthält die Internetseite zu diesem Buch.

Beispiel 4.9

Es wird die gegenseitige Lage der Geraden g und h (siehe Abb. 4.7) mit

$$g: \vec{p} = \vec{p}_0 + t \vec{a} = \begin{pmatrix} -1 \\ 3 \\ 2 \end{pmatrix} + t \cdot \begin{pmatrix} -4 \\ 6 \\ -2 \end{pmatrix}, \quad h: \vec{p} = \vec{q}_0 + s \vec{b} = \begin{pmatrix} 1 \\ 3 \\ 2 \end{pmatrix} + s \cdot \begin{pmatrix} -3 \\ 3 \\ -1 \end{pmatrix}$$

untersucht. Dabei wird wie in dem Beispiel 4.8 verfahren. Durch Gleichsetzen der Parameterdarstellungen ergibt sich das lineare Gleichungssystem

$$\begin{aligned} -4t + 3s &= 2 \\ 6t - 3s &= 5 \\ -2t + s &= 0. \end{aligned}$$

Dieses LGS ist nicht lösbar. Somit existiert kein gemeinsamer Punkt der Geraden g und h . Da die Richtungsvektoren $\vec{a}$ und $\vec{b}$ offensichtlich nicht kollinear sind, sind die Geraden g und h auch nicht parallel, sondern windschief. ■

4.1.4 Parameterdarstellungen als Funktionen, Beschreibung von Bewegungen

Durch Parameterdarstellungen von Geraden werden reellen Zahlen Vektoren $\vec{p}$ und damit Punkte P der Ebene oder des Raumes zugeordnet. Eine Parameterdarstellung lässt sich daher als Funktion φ auffassen mit

$$\varphi: \mathbb{R} \rightarrow \mathbb{R}^2 \text{ mit } \varphi: t \mapsto \vec{p}_0 + t\vec{a} \quad \text{bzw.} \quad \varphi: \mathbb{R} \rightarrow \mathbb{R}^3 \text{ mit } \varphi: t \mapsto \vec{p}_0 + t\vec{a}.$$

Der Parameter t hat hierbei die Funktion einer Variablen. Während aber Funktionswerte bei den aus der Schule bekannten Funktionen ebenfalls Zahlen sind, handelt es sich bei Funktionswerten von Parameterdarstellungen um n -Tupel (wir beschränken uns hier auf Paare und Tripel) reeller Zahlen.

Besonders bei physikalischen Anwendungen ist es sinnvoll, den Parameter t als Zeit aufzufassen. Durch eine Parameterdarstellung $\varphi: t \mapsto \vec{p} = \vec{p}_0 + t\vec{a}$ wird die zeitabhängige Position des durch den Ortsvektor $\vec{p}$ beschriebenen Punktes auf einer Geraden beschrieben. Parameterdarstellungen erhalten damit einen neuen Aspekt; neben der geometrischen Form einer Bewegungsbahn (im einfachsten Falle eine Gerade) wird auch die Art der Bewegung beschrieben.

Beispiel 4.10

Die drei Parameterdarstellungen

$$\begin{aligned} (1) \quad & \vec{p}(t) = \vec{p}_0 + t \cdot \vec{v}_1 \\ (2) \quad & \vec{p}(t) = \vec{p}_0 + t \cdot \vec{v}_2 \\ (3) \quad & \vec{p}(t) = \vec{p}_0 + t^2 \cdot \vec{a} \end{aligned} \quad \text{mit } t \in \mathbb{R}, t \geq 0, \vec{v}_1 = \begin{pmatrix} 1,5 \\ 0 \\ 0 \end{pmatrix}, \vec{v}_2 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \vec{a} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$$

stellen dieselbe Halbgerade dar. Wird der Parameter t als Zeit aufgefasst, so beschreibt (1) eine gleichförmige Bewegung, (2) eine gleichförmige Bewegung geringerer Geschwindigkeit und (3) eine gleichmäßig beschleunigte Bewegung. Die Richtungsvektoren $\vec{v}_1$ und $\vec{v}_2$ sind Geschwindigkeitsvektoren, $\vec{a}$ ist ein Beschleunigungsvektor. In Abb. 4.8 sind Positionen sich bewegendes Punkte bei diesen drei Bewegungsvorgängen innerhalb des Zeitintervalls $[0; 1]$ dargestellt; zwischen zwei benachbarten Punkten verstreicht jeweils gleich viel Zeit. ■

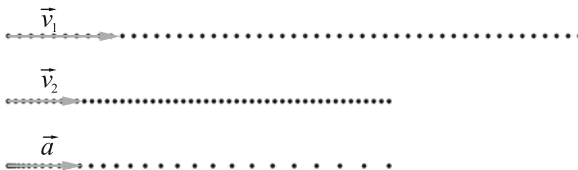


Abb. 4.8: Beschreibung einer geradlinigen Bahn durch drei Parameterdarstellungen

Ein Video, das die drei Bewegungsvorgänge zeigt, befindet sich auf der Internetseite zu diesem Buch.

Beispiel 4.11

Der *schräge Wurf* setzt sich aus einer gleichförmigen Bewegung mit der Abwurfgeschwindigkeit $\vec{v}_0$ und einer gleichmäßig beschleunigten Bewegung mit der Erdbeschleunigung $\vec{g}$ zusammen. Der gleichförmige Anteil beschreibt die (geradlinige) Bahn, die ein geworfener Gegenstand in der Schwerelosigkeit im Vakuum zurücklegen würde: $\vec{p} = \vec{p}_0 + \vec{v}_0 \cdot t$. Dazu addiert sich die Komponente des Fal-

**Abb. 4.9:** Schräger Wurf

Ein Video, das den schrägen Wurf mit seiner Bewegungsbahn zeigt, befindet sich auf der Internetseite zu diesem Buch.

lens, die senkrecht zur Erdoberfläche gerichtet ist und den Betrag $\frac{g}{2} t^2$ hat. Der schräge Wurf wird damit durch die vektorielle Gleichung

$$\vec{p} = \vec{p}_0 + \vec{v}_0 \cdot t + \frac{1}{2} \vec{g} \cdot t^2$$

beschrieben, siehe Abb. 4.9. ■

4.1.5 Exkurs: Parameterdarstellungen einiger Kurven

Bereits in dem Beispiel 4.11 zeigte sich, dass durch Parameterdarstellungen nicht nur Geraden beschrieben werden können, sondern auch Kurven wie die Wurfparabel. Im Folgenden werden – ausgehend von Kreisen – einige interessante Kurven betrachtet und durch Parameterdarstellungen beschrieben.

Die Sinus- und die Kosinusfunktion werden am Einheitskreis mit den in Abb. 4.10 verwendeten Bezeichnungen folgendermaßen definiert:

$$\sin \alpha = y_\alpha, \quad \cos \alpha = x_\alpha.$$

Betrachtet man umgekehrt die Koordinaten von Punkten des Kreises in Abhängigkeit von den entsprechenden Winkeln und multipliziert die trigonometrischen Terme mit einer positiven Zahl r , so erhält man die folgende Parameterdarstellung für einen Kreis in Mittelpunktslage mit beliebigem Radius r :

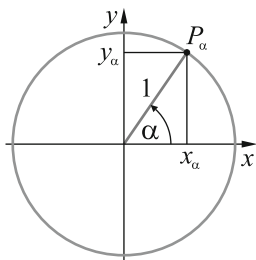
$$x(\alpha) = r \cdot \cos \alpha, \quad y(\alpha) = r \cdot \sin \alpha; \quad \alpha \in [0; 2\pi).$$

Da im Folgenden neue Kurven ausgehend vom Kreis entwickelt und darin weitere Größen in Abhängigkeit von dem Parameter ausgedrückt werden sollen, empfiehlt es sich, einen neuen Parameter t einzuführen, und α durch $2\pi t$ zu ersetzen, wobei t dann das Intervall $[0; 1)$ durchläuft. Daraus ergibt sich die folgende Parameterdarstellung des Kreises:

$$\begin{aligned} x(t) &= r \cdot \cos(2\pi t) \\ y(t) &= r \cdot \sin(2\pi t) \end{aligned} \tag{4.2}$$

Diese Beschreibung des Kreises ist äquivalent zu der Gleichung (2.7) des Kreises in Mittelpunktslage (siehe S. 58), denn es gilt

$$(x(t))^2 + (y(t))^2 = r^2 \cdot (\cos^2(2\pi t) + \sin^2(2\pi t)) = r^2.$$

**Abb. 4.10:** Sinus und Kosinus am Einheitskreis

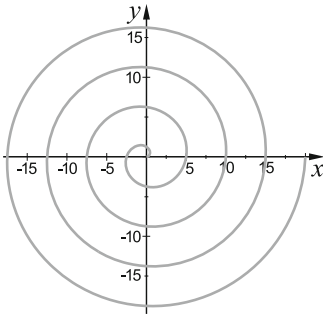


Abb. 4.11: Archimedische Spirale

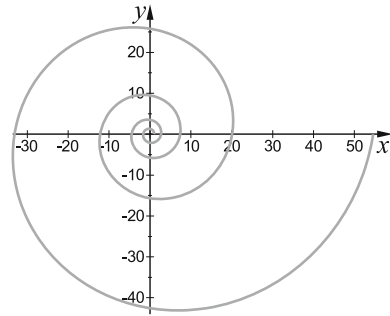


Abb. 4.12: Logarithmische Spirale

Ersetzt man in der Parameterdarstellung (4.2) eines Kreises den konstanten Radius durch Funktionen des Parameters t , so lassen sich damit interessante Kurven erzeugen.

Beispiel 4.12

Ersetzen von r durch $r \cdot t$ in der Parameterdarstellung (4.2) führt dazu, dass sich die Abstände der Punkte zum Mittelpunkt mit steigendem Winkel proportional zum Winkel vergrößern. Die entsprechende Parameterdarstellung

$$\begin{aligned} x(t) &= r t \cdot \cos(2\pi t) \\ y(t) &= r t \cdot \sin(2\pi t) \end{aligned} \quad (t \in [0; 4])$$

beschreibt eine *archimedische Spirale*, siehe Abb. 4.11. Da t das Intervall $[0; 4]$ durchläuft, hat diese Spirale vier Windungen. ■

Beispiel 4.13

Eine weitere bekannte Kurve ist die *logarithmische Spirale*, siehe Abb. 4.12. Auch sie geht aus der Parameterdarstellung (4.2) des Kreises dadurch hervor, dass der konstante Radius r durch eine Funktion des Parameters ersetzt wird. Da hierfür eine Exponentialfunktion verwendet wird, vergrößert sich der Abstand zwischen den Windungen sehr stark. Die in Abb. 4.12 dargestellte logarithmische Spirale hat die Parameterdarstellung

$$\begin{aligned} x(t) &= e^t \cdot \cos(2\pi t) \\ y(t) &= e^t \cdot \sin(2\pi t) \end{aligned} \quad (t \in [0; 4]).$$

■

Kurven im Raum

Ausgehend von Kreisen lassen sich neben ebenen Kurven auch Raumkurven untersuchen. Dazu wird zunächst ein Kreis in einer der Koordinatenebenen durch eine Parameterdarstellung beschrieben, z. B. der Kreis in der x - y -Ebene mit dem Radius r und dem Mittelpunkt im Koordinatenursprung durch

$$\begin{aligned} x(t) &= r \cdot \cos(2\pi t) \\ y(t) &= r \cdot \sin(2\pi t) \\ z(t) &= 0 \end{aligned} \quad (t \in [0; 1)). \quad (4.3)$$

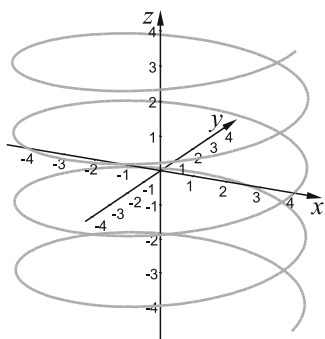


Abb. 4.13: Schraubenlinie

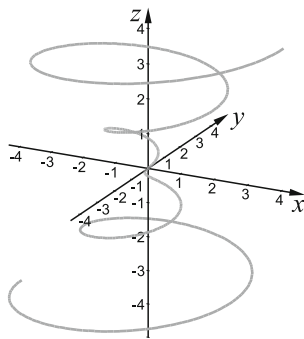


Abb. 4.14: Konische Spirale

Beispiel 4.14

Ersetzt man in der Parameterdarstellung (4.3) die konstante Koordinate z durch eine lineare Funktion y des Parameters t (z. B. $z(t) = h \cdot t$), so erhält man eine Kurve, die dem Gewinde einer Schraube ähnlich ist und als *Schraubenlinie* oder *zylindrische Spirale* oder *Helix* bezeichnet wird, siehe Abb. 4.13:

$$\begin{aligned} x(t) &= r \cdot \cos(2\pi t) \\ y(t) &= r \cdot \sin(2\pi t) \quad (t \in [-2; 2]) \\ z(t) &= h t \end{aligned}$$

Beispiel 4.15

Kombiniert man die Änderungen an der Parameterdarstellung des Kreises, die zur Archimedischen Spirale und zur Schraubenlinie geführt haben, so erhält man eine spiralförmige Kurve, die sich auf einem Kegel entlang windet (siehe Abb. 4.14). Diese Kurve wird als *konische Spirale* bezeichnet, und durch die folgende Parameterdarstellung beschrieben:

$$\begin{aligned} x(t) &= r t \cdot \cos(2\pi t) \\ y(t) &= r t \cdot \sin(2\pi t) \quad (t \in [-2; 2]) \\ z(t) &= h t \end{aligned}$$

Weitere Variationen der bisher betrachteten Kurven ergeben sich aus der Verwendung nichtlinearer Funktionsterme in t für die Höhe bzw. den Radius. Durch Versuche mit verschiedenen Funktionen lassen sich interessante und ästhetisch ansprechende Kurven erzeugen, siehe z. B. die Aufgabe 10 auf S. 150.

Auf der Internetseite zu diesem Buch steht eine Datei für das CAS Maxima zur Verfügung, mit der sich durch Variationen von Parameterdarstellungen leicht unterschiedlichste ebene und räumliche Kurven erzeugen lassen.

4.1.6 Aufgaben zu Abschnitt 4.1

- Überprüfen Sie, welche der Punkte $A(-9; 3; -3)$, $B(11; -5; 5; 9)$, $C(16; -7; 12)$, $D(-6; 5; 2; -1; 5)$ und $E(-16; 5; 6; -8; 5)$ auf der Geraden g mit der Parameterdarstellung $\vec{p} = \begin{pmatrix} 1 \\ -1 \\ 3 \end{pmatrix} + t \cdot \begin{pmatrix} -2 \\ 1 \\ -1 \end{pmatrix}$ liegen. Geben Sie für diejenigen dieser Punkte, die g angehören, die zugehörigen Werte des Parameters t an.

2. Es sind zwei Punkte $P(3; -4; 6)$ und $Q(8; 9; 10)$ einer Geraden g gegeben. Stellen Sie eine Parametergleichung für diese Gerade auf.

3. Eine Gerade in der Ebene ist durch die Koordinatengleichung $7x - 8y = 21$ gegeben. Geben Sie eine Parameterdarstellung dieser Geraden an.

4. Eine Gerade g in der Ebene ist durch folgende Parametergleichung gegeben:

$$\vec{p} = \begin{pmatrix} 6 \\ 11 \end{pmatrix} + t \cdot \begin{pmatrix} -3 \\ 7 \end{pmatrix}.$$

Geben Sie eine Koordinatengleichung von g an.

5. Eine Gerade im Raum ist durch das lineare Gleichungssystem

$$\begin{aligned} 3x + 4y - 2z &= 12 \\ 6x - 7y + z &= -2 \end{aligned}$$

gegeben. Geben Sie eine Parameterdarstellung von g an.

6. Eine Gerade g ist durch die Parameterdarstellung $\vec{p} = \begin{pmatrix} 2 \\ 5 \\ 3 \end{pmatrix} + t \cdot \begin{pmatrix} -4 \\ 3 \\ 2 \end{pmatrix}$ gegeben. Geben Sie ein lineares Gleichungssystem an, das g beschreibt.

7. Untersuchen Sie paarweise die gegenseitige Lage der Geraden g , h und k und bestimmen Sie, falls vorhanden, die jeweiligen Schnittpunkte.

$$g: \vec{p} = \begin{pmatrix} 4 \\ 5 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} 3 \\ -4 \\ 2 \end{pmatrix}, \quad h: \vec{p} = \begin{pmatrix} 6 \\ -6 \\ 23 \end{pmatrix} + s \cdot \begin{pmatrix} -2 \\ 1 \\ 3 \end{pmatrix}, \quad k: \vec{p} = \begin{pmatrix} -4 \\ 3 \\ -2 \end{pmatrix} + r \cdot \begin{pmatrix} 12 \\ -6 \\ -18 \end{pmatrix}$$

8. Beweisen Sie den folgenden Satz:

Zwei Geraden $g: \vec{p} = \vec{p}_0 + t\vec{a}$ und $h: \vec{q} = \vec{q}_0 + s\vec{b}$ mit Richtungsvektoren $\vec{a}$ und $\vec{b}$, die nicht kollinear sind, besitzen *genau dann* einen Schnittpunkt, wenn die Vektoren $\vec{a}$, $\vec{b}$ und $\vec{q}_0 - \vec{p}_0$ komplanar sind.

9. Stellen Sie eine Parameterdarstellung für eine Ellipse in Hauptachsenlage auf, die kein Kreis ist. Zeigen Sie, dass die von Ihnen aufgestellte Parameterdarstellung dieselbe Ellipse beschreibt wie die Gleichung (2.14) auf S. 77.

10. Die Kurve in Abb. 4.15 wurde durch eine Parameterdarstellung generiert. Geben Sie eine Parameterdarstellung an, die eine derartige Kurve erzeugt. Experimentieren Sie dazu auch mithilfe des Computers. (Sie können dazu eine auf der Internetseite zu diesem Buch vorhandene Datei zur Darstellung von Kurven in dem CAS Maxima benutzen.)

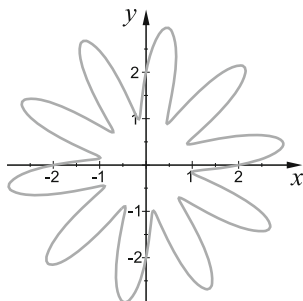


Abb. 4.15: Durch eine Parameterdarstellung beschriebene Kurve

4.2 Ebenen

4.2.1 Parameter- und Koordinatengleichungen von Ebenen

In dem Abschnitt 1.2.1 wurden Lösungsmengen linearer Gleichungen mit drei Variablen betrachtet. Es zeigte sich, dass diese i. Allg. als Ebenen des Raumes interpretiert und in der Form (1.14) beschrieben werden können, siehe S. 17.

Parameterdarstellungen von Ebenen

Jede Ebene ε im Raum lässt sich durch eine Gleichung der Form

$$\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b} \quad (4.4)$$

mit zwei nicht kollinearen Vektoren $\vec{a}$ und $\vec{b}$ beschreiben. Dabei sind $\vec{a}$ und $\vec{b}$ Richtungs- bzw. Spannvektoren und $\vec{p}_0$ ein Stützvektor der Ebene ε . Die Parameter s und t durchlaufen den gesamten Bereich der reellen Zahlen, wobei jedem Parameterpaar $(s; t)$ ein Punkt P der Ebene ε mit dem Ortsvektor $\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b}$ zugeordnet ist.

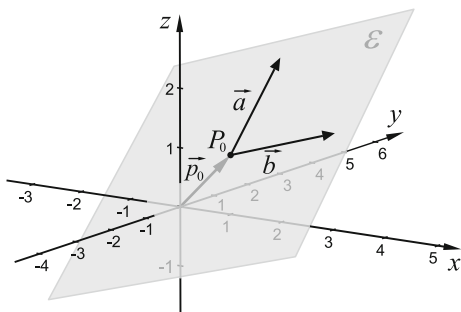


Abb. 4.16: Ebene mit Pfeilen, die einen Stützvektor und zwei Richtungsvektoren repräsentieren

Eine interaktive Version dieser Abbildung, die aus verschiedenen Richtungen betrachtet werden kann, enthält die Internetseite zu diesem Buch.

Beispiel 4.16

Durch die Parameterdarstellungen $\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b}$ und $\vec{p} = \vec{q}_0 + u\vec{c} + v\vec{d}$ mit $\vec{p}_0 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$, $\vec{a} = \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 3 \\ -1,5 \\ 1 \end{pmatrix}$ sowie $\vec{q}_0 = \begin{pmatrix} 4 \\ 6 \\ 1 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} -9 \\ 12 \\ 1 \end{pmatrix}$, $\vec{d} = \begin{pmatrix} 11 \\ -3 \\ 5 \end{pmatrix}$ wird dieselbe Ebene beschrieben (siehe Abb. 4.16).

Die Vektoren $\vec{c}$ und $\vec{d}$ sind als Linearkombinationen von $\vec{a}$ und $\vec{b}$ darstellbar:

$$\vec{c} = 3 \cdot \vec{a} - 2 \cdot \vec{b}, \quad \vec{d} = \vec{a} + 4 \cdot \vec{b}.$$

Außerdem liegt der Punkt Q_0 in der durch $\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b}$ beschriebenen Ebene. Dies lässt sich mittels einer Punktprobe überprüfen, indem $\vec{q}_0$ in diese Parameterdarstellung eingesetzt wird:

$$\begin{pmatrix} 4 \\ 6 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} 3 \\ -1,5 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{array}{rcl} -1s + 3t & = & 3 \\ 3s - 1,5t & = & 6 \\ s + t & = & 5. \end{array}$$

Dieses LGS ist lösbar, als Lösung erhält man $s=3, t=2$. Somit liegt Q_0 in der durch $\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b}$ beschriebenen Ebene. Da außerdem, wie bereits ausgeführt wurde, $\vec{c}, \vec{a}$ und $\vec{b}$ sowie $\vec{d}, \vec{a}$ und $\vec{b}$ jeweils komplanar sind, wird durch die beiden Parameterdarstellungen dieselbe Ebene beschrieben. ■

Wie bereits in dem Abschnitt 1.2.1 deutlich wurde, können Ebenen auch durch lineare Gleichungen mit drei Variablen beschrieben werden:

$$ax + by + cz = d.$$

Man nennt Gleichungen dieser Form auch *Koordinatengleichungen* oder *parameterfreie Ebenengleichungen*. Ist eine Ebene in dieser Form gegeben, so lässt sich in der bereits in den Beispielen 1.13-1.15 (siehe S. 15-17) beschriebenen Weise eine Parametergleichung aufstellen, siehe hierzu die Aufgabe 2 auf S. 155. Umgekehrt ist es auch möglich, parameterfreie Gleichungen für Ebenen aufzustellen, die durch Parameterdarstellungen gegeben sind.

Beispiel 4.17

Für die bereits in Beispiel 4.16 betrachtete Ebene mit der Parameterdarstellung

$$\vec{p} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} 3 \\ -1,5 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{array}{ll} \text{(I)} & x = -s + 3t + 1 \\ \text{(II)} & y = 3s - 1,5t \\ \text{(III)} & z = s + t + 1 \end{array}$$

wird eine Koordinatengleichung aufgestellt. Dazu wird zunächst aus jeweils zwei der Gleichungen einer der beiden Parameter eliminiert:

$$\begin{array}{ll} 3 \cdot \text{(I)} + \text{(II)} & 3x + y = 7,5t + 3 \\ \text{(I)} + \text{(III)} & x + z = 4t + 2. \end{array}$$

Nun lässt sich z. B. die zweite der entstandenen Gleichungen nach dem verbliebenen Parameter t umstellen: $t = \frac{x}{4} + \frac{z}{4} - \frac{1}{2}$. Durch Einsetzen in die erste der Gleichungen, aus denen s eliminiert wurde, ergibt sich:

$$3x + y = 7,5 \left(\frac{x}{4} + \frac{z}{4} - \frac{1}{2} \right) + 3 \quad \text{bzw.} \quad \frac{9}{8}x + y - \frac{15}{8}z = -\frac{3}{4}.$$

Durch diese Gleichung, die sich noch zu $9x + 8y - 15z = -6$ vereinfachen lässt, wird die gegebene Ebene beschrieben. ■

4.2.2 Lagebeziehungen und Schnittpunkte von Geraden und Ebenen sowie von Ebenen und Ebenen

Für eine Gerade g und eine Ebene ε im Raum bestehen drei Möglichkeiten der gegenseitigen Lage:

- Die Gerade g liegt in der Ebene ε , d. h. $g \subset \varepsilon$.
- g und ε haben keinen gemeinsamen Punkt: $g \cap \varepsilon = \{\}$, g und ε sind parallel.
- g und ε schneiden sich in genau einem Punkt S : $g \cap \varepsilon = \{S\}$.

Beispiel 4.18

Es wird die gegenseitige Lage der Geraden g und der Ebene ε (Abb. 4.17) mit

$$g: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 4 \\ -2 \\ -1 \end{pmatrix} + r \cdot \begin{pmatrix} -1 \\ -0,5 \\ 0,5 \end{pmatrix}, \quad \varepsilon: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -2 \\ 0 \\ 2 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} 2 \\ -2 \\ -1 \end{pmatrix}$$

untersucht. Durch Gleichsetzen der Parameterdarstellungen von g und ε lässt sich überprüfen, ob g und ε einen Schnittpunkt besitzen oder sogar alle Punkte von g in ε liegen:

$$\begin{pmatrix} 4 \\ -2 \\ -1 \end{pmatrix} + r \cdot \begin{pmatrix} -1 \\ -0,5 \\ 0,5 \end{pmatrix} = \begin{pmatrix} -2 \\ 0 \\ 2 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ 3 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} 2 \\ -2 \\ -1 \end{pmatrix} \quad \text{bzw.} \quad \begin{array}{rcl} -r + s - 2t & = & -6 \\ -0,5r - 3s + 2t & = & 2 \\ 0,5r - s + t & = & 3. \end{array}$$

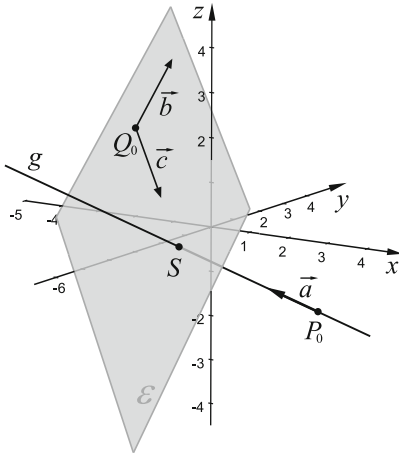


Abb. 4.17: Gerade und Ebene mit Schnittpunkt

Eine interaktive Version dieser Abbildung enthält die Internetseite zu diesem Buch.

Dieses lineare Gleichungssystem ist eindeutig lösbar; als Lösung erhält man $r = \frac{8}{3}$, $s = 0$, $t = \frac{5}{3}$. Somit schneiden sich g und ε in einem Punkt S , dessen Koordinaten sich durch Einsetzen dieser Werte in eine der Parameterdarstellungen ermitteln lassen: $S\left(\frac{4}{3}; -\frac{10}{3}; \frac{1}{3}\right)$. ■

Falls bei dem in dem Beispiel 4.18 praktizierten Vorgehen ein nicht lösbares LGS entsteht, so sind die betrachtete Gerade und die Ebene parallel; ergibt sich ein LGS mit unendlich vielen Lösungen, so ist $g \subset \varepsilon$. In welcher Lage sich eine Gerade und eine Ebene zueinander befinden, hängt von deren Richtungsvektoren ab.

Satz 4.1

Gegeben seien eine Gerade g und eine Ebene ε durch die Parameterdarstellungen $g: \vec{p} = \vec{p}_0 + r\vec{a}$ und $\varepsilon: \vec{p} = \vec{q}_0 + s\vec{b} + t\vec{c}$.

- Falls die Vektoren $\vec{a}$, $\vec{b}$ und $\vec{c}$ komplanar sind, so ist $g \subset \varepsilon$ oder $g \cap \varepsilon = \{\}$.
- Falls $\vec{a}$, $\vec{b}$ und $\vec{c}$ nicht komplanar sind, so existiert genau ein gemeinsamer Punkt von g und ε .

Beweis: Sind die Vektoren $\vec{a}$, $\vec{b}$ und $\vec{c}$ komplanar, so lässt sich, da $\vec{b}$ und $\vec{c}$ als Spannvektoren einer Ebene nicht kollinear sein können, $\vec{a}$ als Linearkombination von $\vec{b}$ und $\vec{c}$ darstellen: $\vec{a} = \lambda\vec{b} + \mu\vec{c}$. Wir zeigen, dass *alle* Punkte von g auch ε angehören, falls es *einen* gemeinsamen Punkt von g und ε gibt.

Es sei S (mit dem Ortsvektor $\vec{s}$) ein gemeinsamer Punkt von g und ε . Dann existieren $r_0, s_0, t_0 \in \mathbb{R}$ mit $\vec{s} = \vec{p}_0 + r_0\vec{a}$ und $\vec{s} = \vec{q}_0 + s_0\vec{b} + t_0\vec{c}$. Ist P ein beliebiger Punkt der Geraden g (mit dem Ortsvektor $\vec{p}$), so gilt

$$\begin{aligned}\vec{p} &= \vec{p}_0 + r\vec{a} = \vec{s} - r_0\vec{a} + r\vec{a} = \vec{s} + (r - r_0)\vec{a} \\ &= \vec{q}_0 + s_0\vec{b} + t_0\vec{c} + (r - r_0) \cdot (\lambda\vec{b} + \mu\vec{c}) \\ &= \vec{q}_0 + (s_0 + \lambda r - \lambda r_0) \cdot \vec{b} + (t_0 + \mu r - \mu r_0) \cdot \vec{c}.\end{aligned}$$

Somit ist $P \in \varepsilon$. Wenn also ein Punkt von g in ε liegt, so gehören alle Punkte von g auch der Ebene ε an. Unter der Voraussetzung, dass $\vec{a}$, $\vec{b}$ und $\vec{c}$ komplanar sind, kann also nur $g \cap \varepsilon = \{\}$ oder $g \subset \varepsilon$ sein.

Auf einen Beweis des Teils b) wird verzichtet (siehe hierzu aber die Aufgabe 5 auf S. 155). □

Lagebeziehungen und Schnittgeraden zweier Ebenen

Für zwei Ebenen ε_1 und ε_2 bestehen drei Möglichkeiten der gegenseitigen Lage:

- ε_1 und ε_2 sind identisch, siehe das Beispiel 4.16 auf S. 151.
- ε_1 und ε_2 sind parallel, haben also keinen gemeinsamen Punkt: $\varepsilon_1 \cap \varepsilon_2 = \{\}$.
- ε_1 und ε_2 haben eine Schnittgerade g , d. h. $\varepsilon_1 \cap \varepsilon_2 = g$.

Beispiel 4.19

Es wird die gegenseitige Lage der Ebenen ε_1 und ε_2 (siehe Abb. 4.18) untersucht:

$$\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 6 \\ 5 \\ 3 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ 7 \\ 10 \end{pmatrix} + t \cdot \begin{pmatrix} -3 \\ 0 \\ 3 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 \\ 9 \\ 7 \end{pmatrix} + u \cdot \begin{pmatrix} 2 \\ 11 \\ 8 \end{pmatrix} + v \cdot \begin{pmatrix} 2 \\ -4 \\ -1 \end{pmatrix}.$$

Durch Gleichsetzen der beiden Parameterdarstellungen ergibt sich das LGS:

$$\begin{aligned} s - 3t - 2u - 2v &= -5 \\ 7s &- 11u + 4v = 4 \\ 10s + 3t - 8u + v &= 4. \end{aligned}$$

Dieses lineare Gleichungssystem hat eine einparametrische Lösungsmenge:

$$s = r - 1, \quad t = 2 - r, \quad u = r - 1, \quad v = r.$$

Dies bedeutet, dass die Ebenen ε_1 und ε_2 eine Schnittgerade g besitzen. Eine Parameterdarstellung dieser Schnittgeraden erhält man durch Einsetzen der Lösungen für s, t bzw. u, v in eine der Parameterdarstellungen von ε_1 bzw. ε_2 (es ist sinnvoll, in beide Parametergleichungen einzusetzen, um das Ergebnis zu überprüfen):

$$g: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -1 \\ -2 \\ -1 \end{pmatrix} + r \cdot \begin{pmatrix} 4 \\ 7 \\ 7 \end{pmatrix}.$$

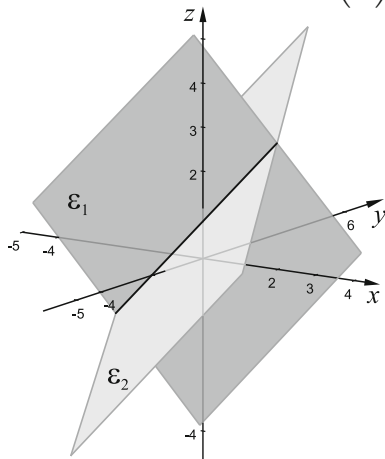


Abb. 4.18: Zwei Ebenen mit Schnittgerade

Eine interaktive Version dieser Abbildung enthält die Internetseite zu diesem Buch.

Ergibt sich bei dem in dem Beispiel 4.19 durchgeführten Verfahren ein nicht lösbares Gleichungssystem, so existieren keine Schnittpunkte, die gegebenen Ebenen sind in diesem Falle parallel (siehe die Aufgabe 7 auf S. 155).

Sind zwei Ebenen durch Koordinatengleichungen gegeben, so kann ihre gegenseitige Lage durch Lösen des aus den beiden Ebenengleichungen bestehenden linearen Gleichungssystems ermittelt werden, siehe das Beispiel 1.16 auf S. 19.

4.2.3 Aufgaben zu Abschnitt 4.2

1. Gegeben sind drei Punkte P , Q und R . Überprüfen Sie, ob durch diese Punkte eindeutig eine Ebene bestimmt wird und – falls ja – stellen Sie eine Parametergleichung dieser Ebene auf.

a) $P(-1; 1; 2)$, $Q(2; 3; -1)$, $R(3; 3; 3)$ b) $P(4; 5; 7)$, $Q(3; -2; -1)$, $R(5; 12; 15)$

2. Gegeben ist die folgende Gleichung einer Ebene ε : $3x - 5y - 2z = 8$.
Geben Sie eine Parameterdarstellung der Ebene ε an.

3. Gegeben sind zwei Ebenen ε_1 und ε_2 durch Parameterdarstellungen. Stellen Sie Koordinatengleichungen für diese Ebenen auf.

$$\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 \\ -2 \\ 0 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix} + t \cdot \begin{pmatrix} 3 \\ 3 \\ 1 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 \\ 1 \\ -3 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ 4 \\ 6 \end{pmatrix} + t \cdot \begin{pmatrix} 1 \\ -2 \\ 3 \end{pmatrix}$$

4. Untersuchen Sie die gegenseitige Lage der Geraden g und der Ebene ε . Geben Sie (falls vorhanden) den Schnittpunkt an.

a) $g: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -9 \\ 30 \\ 4 \end{pmatrix} + r \cdot \begin{pmatrix} -2 \\ -1 \\ 7 \end{pmatrix}, \quad \varepsilon: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 4 \\ 7 \\ -8 \end{pmatrix} + s \cdot \begin{pmatrix} -1 \\ -3 \\ 2 \end{pmatrix} + t \cdot \begin{pmatrix} -5 \\ 3 \\ 8 \end{pmatrix}$

b) $g: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 \\ -4 \\ 7 \end{pmatrix} + r \cdot \begin{pmatrix} 16 \\ -1 \\ -3 \end{pmatrix}, \quad \varepsilon: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 3 \\ 5 \\ 4 \end{pmatrix} + s \cdot \begin{pmatrix} -2 \\ -3 \\ 5 \end{pmatrix} + t \cdot \begin{pmatrix} 5 \\ -5 \\ 6 \end{pmatrix}$

c) $g: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 21 \\ 7 \\ 0 \end{pmatrix} + r \cdot \begin{pmatrix} -3 \\ -14 \\ -17 \end{pmatrix}, \quad \varepsilon: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 12 \\ 5 \\ 1 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ -2 \\ -3 \end{pmatrix} + t \cdot \begin{pmatrix} 7 \\ 6 \\ 5 \end{pmatrix}$

5. Gegeben seien eine Gerade $g: \vec{p} = \vec{p}_0 + r \vec{a}$ und eine Ebene $\varepsilon: \vec{p} = \vec{q}_0 + s \vec{b} + t \vec{c}$.
Weisen Sie nach: Falls $\vec{a}$, $\vec{b}$ und $\vec{c}$ nicht komplanar sind, so können g und ε nicht mehr als einen gemeinsamen Punkt besitzen.

6. Gegeben sind zwei Ebenen ε_1 und ε_2 .

$$\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 \\ 4 \\ 3 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 5 \\ 0 \\ 0 \end{pmatrix} + u \cdot \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + v \cdot \begin{pmatrix} 1 \\ 0 \\ 6 \end{pmatrix}$$

- a) Überprüfen Sie, ob sich ε_1 und ε_2 schneiden und bestimmen Sie eine Parameterdarstellung der Schnittgeraden.

- b) Geben Sie Parameterdarstellungen für ε_1 und ε_2 an, aus denen sofort ersichtlich wird, dass sich diese beiden Ebenen in der unter a) bestimmten Geraden schneiden.

7. Untersuchen Sie die gegenseitige Lage der Ebenen ε_1 und ε_2 . Geben Sie eine Parameterdarstellung der Schnittgeraden an, falls eine solche existiert.

a) $\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 5 \\ 2 \\ 0 \end{pmatrix} + s \cdot \begin{pmatrix} 2 \\ 3 \\ 6 \end{pmatrix} + t \cdot \begin{pmatrix} 4 \\ 3 \\ 8 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -7 \\ -7 \\ -24 \end{pmatrix} + u \cdot \begin{pmatrix} 6 \\ 6 \\ 14 \end{pmatrix} + v \cdot \begin{pmatrix} -2 \\ 0 \\ -2 \end{pmatrix}$

b) $\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + s \cdot \begin{pmatrix} 7 \\ -1 \\ 5 \end{pmatrix} + t \cdot \begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 3 \\ 9 \\ 3 \end{pmatrix} + u \cdot \begin{pmatrix} -1 \\ -11 \\ -7 \end{pmatrix} + v \cdot \begin{pmatrix} 15 \\ 9 \\ 17 \end{pmatrix}$

c) $\varepsilon_1: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 \\ 4 \\ 3 \end{pmatrix} + s \cdot \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix}, \quad \varepsilon_2: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 5 \\ 0 \\ 0 \end{pmatrix} + u \cdot \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + v \cdot \begin{pmatrix} 1 \\ 0 \\ 6 \end{pmatrix}$

4.3 Metrische Geometrie von Geraden und Ebenen

Bislang wurden Lagebeziehungen und Schnittpunkte bzw. -geraden von Geraden und Ebenen untersucht, wobei noch keinerlei Maße auftraten. Im Folgenden wird *metrische Geometrie* von Geraden und Ebenen, d. h. Geometrie *unter Verwendung von Längen- und Winkelmaßen* betrieben. Die entscheidende Grundlage hierfür ist das in dem Abschnitt 3.5 behandelte Skalarprodukt.

4.3.1 Normalengleichungen von Geraden und Ebenen

Normalenvektoren und -gleichungen von Geraden in der Ebene sowie von Ebenen im Raum ermöglichen es, auf recht einfache Weise verschiedene Winkelberechnungen durchzuführen.

Normalengleichungen von Geraden in der Ebene

Alle Punkte P der Ebene, deren Verbindungsvektoren $\overrightarrow{P_0P}$ mit einem Punkt P_0 zu einem Vektor $\vec{n}$ ($\vec{n} \neq \vec{o}$) orthogonal sind, liegen auf einer Geraden.

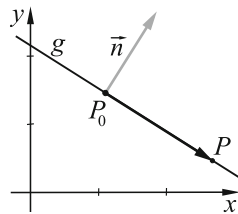


Abb. 4.19: Normalenvektor einer Geraden

Definition 4.1

- Ein zu einer Geraden g orthogonaler Vektor $\vec{n}$ (mit $\vec{n} \neq \vec{o}$) heißt *Normalenvektor* der Geraden g .
- Ist g eine Gerade, P_0 ein Punkt und $\vec{n}$ ein Normalenvektor von g , so ist die Gleichung $\vec{n} \cdot \overrightarrow{P_0P} = 0$ eine *Normalengleichung der Geraden* g . ♦

Beispiel 4.20

Eine Gerade g ist durch die Parametergleichung $g: \vec{x} = \begin{pmatrix} 5 \\ 6 \end{pmatrix} + t \begin{pmatrix} -3 \\ 4 \end{pmatrix}$ gegeben. Gesucht ist eine Normalengleichung von g .

Ein Normalenvektor $\vec{n} = \begin{pmatrix} x_n \\ y_n \end{pmatrix}$ von g muss orthogonal zu dem Richtungsvektor $\begin{pmatrix} -3 \\ 4 \end{pmatrix}$ sein, d. h. $\begin{pmatrix} -3 \\ 4 \end{pmatrix} \cdot \begin{pmatrix} x_n \\ y_n \end{pmatrix} = -3x_n + 4y_n = 0$. Somit ist z. B. $\vec{n} = \begin{pmatrix} 4 \\ 3 \end{pmatrix}$ ein Normalenvektor und $g: \begin{pmatrix} 4 \\ 3 \end{pmatrix} \cdot \begin{pmatrix} x-5 \\ y-6 \end{pmatrix} = 0$ eine Normalengleichung von g . Diese lässt sich auch in der Form $4 \cdot (x-5) + 3 \cdot (y-6) = 0$ bzw. $4x + 3y = 38$ schreiben. ■

Im Raum ist es nicht möglich, eine Gerade durch eine Normalengleichung der Form $\vec{n} \cdot \overrightarrow{P_0P} = 0$ zu beschreiben, denn die Menge aller Punkte P des Raumes, deren Verbindungsvektoren $\overrightarrow{P_0P}$ mit einem Punkt P_0 zu einem Vektor $\vec{n}$ ($\vec{n} \neq \vec{o}$) orthogonal sind, ist eine Ebene.

Normalengleichungen von Ebenen

Beispiel 4.21

Gegeben sind ein Punkt $P_0(2; 2; 1)$ und ein Vektor $\vec{n} = \begin{pmatrix} -1 \\ 3 \\ 2 \end{pmatrix}$. Wir betrachten die Menge aller Punkte P , für die der Verbindungsvektor $\overrightarrow{P_0P}$ zu $\vec{n}$ orthogonal ist. Für diese Punkte gilt $\vec{n} \cdot \overrightarrow{P_0P} = 0$, d. h.:

$$\begin{pmatrix} -1 \\ 3 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} x-2 \\ y-2 \\ z-1 \end{pmatrix} = 0 \quad \text{bzw.} \quad -(x-2) + 3(y-2) + 2(z-1) = 0.$$

Diese Bedingung ist für alle Punkte $P(x; y; z)$ mit $-x + 3y + 2z = 6$ ($x, y, z \in \mathbb{R}$) erfüllt. Da dies die Gleichung einer Ebene ist, bildet die Menge aller Punkte P mit $\overrightarrow{P_0P} \perp \vec{n}$ eine Ebene ε (siehe Abb. 4.20). Der Vektor $\vec{n}$ steht auf ε senkrecht.

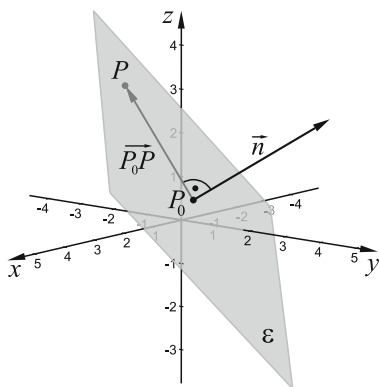


Abb. 4.20: Normalenvektor einer Ebene

Eine interaktive Version dieser Abbildung enthält die Internetseite zu diesem Buch.

Definition 4.2

Ein Vektor $\vec{n}$, der zu einer Ebene ε senkrecht ist, heißt *Normalenvektor* von ε .

Die Lage einer Ebene ε im Raum ist durch einen Punkt $P_0 \in \varepsilon$ und einen ihrer Normalenvektoren $\vec{n}$ eindeutig bestimmt. Ein Punkt P liegt also genau dann in ε , wenn der Vektor $\overrightarrow{P_0P}$ zu $\vec{n}$ orthogonal ist, also $\vec{n} \cdot \overrightarrow{P_0P} = 0$ gilt.

Definition 4.3

Ist ε eine Ebene, P_0 ein Punkt und $\vec{n}$ ein Normalenvektor von ε , so heißt die Gleichung $\vec{n} \cdot \overrightarrow{P_0P} = 0$ *Normalengleichung der Ebene* ε .

Wie das Beispiel 4.21 zeigt, lässt sich die Normalengleichung einer Ebene in die Form $ax + by + cz = d$ bringen; man erhält also eine Koordinatengleichung.

Satz 4.2

Ist $ax + by + cz = d$ eine Gleichung einer Ebene ε , so ist der Vektor $\vec{n} = \begin{pmatrix} a \\ b \\ c \end{pmatrix}$ ein Normalenvektor von ε .

Beweis: Sind $P_0(x_0; y_0; z_0)$, $P(x; y; z)$ Punkte von ε , so gilt $ax_0 + by_0 + cz_0 = d$ und $ax + by + cz = d$. Durch Subtraktion beider Gleichungen ergibt sich

$$a(x - x_0) + b(y - y_0) + c(z - z_0) = 0,$$

also ist

$$\begin{pmatrix} a \\ b \\ c \end{pmatrix} \cdot \begin{pmatrix} x-x_0 \\ y-y_0 \\ z-z_0 \end{pmatrix} = \begin{pmatrix} a \\ b \\ c \end{pmatrix} \cdot \overrightarrow{P_0P} = 0.$$

Da P beliebig gewählt wurde, sind alle Verbindungsvektoren von Punkten der Ebene ε zu $\begin{pmatrix} a \\ b \\ c \end{pmatrix}$ orthogonal; dieser Vektor ist somit Normalenvektor der Ebene. □

Beispiel 4.22

Bestimmung einer Normalengleichung aus einer Parametergleichung einer Ebene

Gegeben ist die Parametergleichung einer Ebene ε : $\vec{p} = \begin{pmatrix} 2 \\ 7 \\ 8 \end{pmatrix} + s \begin{pmatrix} 5 \\ 0 \\ 3 \end{pmatrix} + t \begin{pmatrix} -8 \\ -3 \\ 0 \end{pmatrix}$.

Ein Normalenvektor muss zu beiden Richtungsvektoren von ε orthogonal sein, er lässt sich durch das Vektorprodukt der Richtungsvektoren bestimmen:

$$\vec{n} = \begin{pmatrix} 5 \\ 0 \\ 3 \end{pmatrix} \times \begin{pmatrix} -8 \\ -3 \\ 0 \end{pmatrix} = \begin{pmatrix} 9 \\ -24 \\ -15 \end{pmatrix}.$$

Natürlich ist auch jedes Vielfache dieses Vektors ein Normalenvektor von ε , z. B. $\vec{n} = \begin{pmatrix} 3 \\ -8 \\ -5 \end{pmatrix}$. Da aus der Parametergleichung hervorgeht, dass $P_0(2; 7; 8)$ ein Punkt von ε ist, ergibt sich die Normalengleichung

$$\varepsilon: \vec{n} \cdot \overrightarrow{P_0P} = \begin{pmatrix} 3 \\ -8 \\ -5 \end{pmatrix} \cdot \begin{pmatrix} x-2 \\ y-7 \\ z-8 \end{pmatrix} = 0$$

oder in Koordinatenschreibweise $\varepsilon: 3x - 8y - 5z = -90$. ■

Beispiel 4.23

Bestimmung einer Parametergleichung aus der Normalengleichung einer Ebene

Gegeben ist die Normalengleichung einer Ebene ε : $\begin{pmatrix} -1 \\ 1 \\ 4 \end{pmatrix} \cdot \begin{pmatrix} x-4 \\ y+1 \\ z-2 \end{pmatrix} = 0$.

Um eine Parametergleichung $\vec{p} = \vec{p}_0 + t\vec{a} + s\vec{b}$ von ε zu ermitteln, sind zwei linear unabhängige Richtungsvektoren $\vec{a}$ und $\vec{b}$ der Ebene ε zu bestimmen, die zu $\vec{n}$ orthogonal sind. Gesucht sind also Vektoren $\vec{a}$ und $\vec{b}$ mit

$$\vec{a} \cdot \vec{n} = \begin{pmatrix} x_a \\ y_a \\ z_a \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 1 \\ 4 \end{pmatrix} = 0 \quad \text{und} \quad \vec{b} \cdot \vec{n} = \begin{pmatrix} x_b \\ y_b \\ z_b \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 1 \\ 4 \end{pmatrix} = 0.$$

Als Richtungsvektoren von ε können somit z. B. $\vec{a} = \begin{pmatrix} -1 \\ 3 \\ -1 \end{pmatrix}$ und $\vec{b} = \begin{pmatrix} 5 \\ 1 \\ 1 \end{pmatrix}$ verwendet werden. Da aus der Normalengleichung hervorgeht, dass $P_0(4; -1; 2)$ ein Punkt der Ebene ε ist, ergibt sich als eine Parametergleichung:

$$\vec{p} = \begin{pmatrix} 4 \\ -1 \\ 2 \end{pmatrix} + t \begin{pmatrix} -1 \\ 3 \\ -1 \end{pmatrix} + s \begin{pmatrix} 5 \\ 1 \\ 1 \end{pmatrix}. \quad \text{■}$$

Eine andere Möglichkeit, eine Parametergleichung einer Ebene aus einer Normalengleichung zu bestimmen, besteht darin, die Normalengleichung zunächst in die Form $ax + by + cz = d$ zu bringen und dann diese Gleichung – wie bereits in den Beispielen 1.13-1.15 auf S. 15-17 beschrieben – in eine Parametergleichung umzuwandeln.

4.3.2 Schnittwinkel zwischen Geraden und Ebenen

Schnittwinkel zweier Geraden

Zwei sich schneidende Geraden g und h schließen zwei Winkel miteinander ein, die sich zu 180° ergänzen. Derjenige dieser beiden Winkel, der nicht größer als 90° ist, heißt *Schnittwinkel der Geraden g und h* .

Beispiel 4.24

Der Schnittwinkel der Geraden $g: \vec{x} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 0 \\ 1 \\ \frac{1}{2} \end{pmatrix}$ und $h: \vec{x} = \begin{pmatrix} 4 \\ 1 \\ 0 \end{pmatrix} + s \begin{pmatrix} -\frac{3}{2} \\ 1 \\ 1 \end{pmatrix}$

(siehe Abb. 4.21 a) mit dem Schnittpunkt $S(1; 3; 2)$ ist gleich dem Winkel α zwischen ihren Richtungsvektoren (siehe S. 129):

$$\cos \alpha = \frac{0 \cdot (-1,5) + 1 \cdot 1 + 0,5 \cdot 1}{\sqrt{1 + 0,25} \cdot \sqrt{2,25 + 1 + 1}} \approx 0,651, \quad \angle(g, h) = \alpha \approx 49,4^\circ. \quad \blacksquare$$

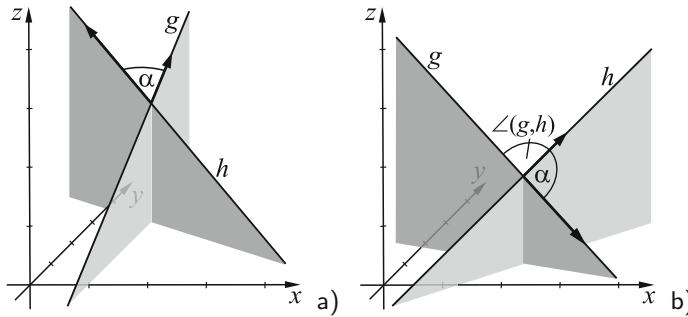


Abb. 4.21:
Schnittwinkel zweier Geraden

Beispiel 4.25

Für den Winkel zwischen den Richtungsvektoren der sich schneidenden Geraden

$$g: \vec{x} = \begin{pmatrix} -0,5 \\ 2 \\ 3,5 \end{pmatrix} + t \begin{pmatrix} 1,25 \\ -0,5 \\ -1 \end{pmatrix} \quad \text{und} \quad h: \vec{x} = \begin{pmatrix} 0,5 \\ -1 \\ 0 \end{pmatrix} + s \begin{pmatrix} 0,75 \\ 1 \\ 0,75 \end{pmatrix}$$

(siehe Abb. 4.21 b) ergibt sich:

$$\cos \alpha = \frac{1,25 \cdot 0,75 - 0,5 \cdot 1 - 1 \cdot 0,75}{\sqrt{1,25^2 + 0,5^2 + 1} \cdot \sqrt{0,75^2 + 1 + 0,75^2}} \approx -0,1278$$

und somit $\alpha \approx 97,3^\circ$. Der Schnittwinkel der Geraden g und h ist jedoch nicht größer als 90° . Er lässt sich als Differenz zwischen 180° und α berechnen:

$$\angle(g, h) \approx 180^\circ - 97,3^\circ = 82,7^\circ. \quad \blacksquare$$

Anstatt zunächst den Winkel zwischen den beiden Richtungsvektoren der Geraden g und h zu berechnen und diesen eventuell (wie in dem Beispiel 4.25) von 180° zu subtrahieren, kann auch der Absolutbetrag des Skalarproduktes für die Berechnung genutzt werden.

Satz 4.3

Für den Schnittwinkel $\angle(g, h)$ zweier sich schneidender Geraden $g: \vec{x} = \vec{p}_1 + t \vec{a}$

und $h: \vec{x} = \vec{p}_2 + s \vec{b}$ gilt: $\cos \angle(g, h) = \left| \cos \angle(\vec{a}, \vec{b}) \right| = \frac{|\vec{a} \cdot \vec{b}|}{|\vec{a}| \cdot |\vec{b}|}.$

Schnittwinkel zweier Ebenen

Unter dem Schnittwinkel zweier Ebenen ε_1 und ε_2 mit der Schnittgeraden g versteht man den Winkel zwischen zwei sich schneidenden Geraden g_1 und g_2 , die in ε_1 bzw. ε_2 liegen und auf der Schnittgeraden g senkrecht stehen (Abb. 4.22).

Sind $\vec{n}_1$ und $\vec{n}_2$ Normalenvektoren der Ebenen ε_1 bzw. ε_2 , so steht $\vec{n}_1$ auf ε_1 (und daher auch auf der in ε_1 liegenden Geraden g_1) senkrecht. Ebenso ist $\vec{n}_2$ senkrecht zu ε_2 und somit zu g_2 . Der Winkel zwischen den Normalenvektoren $\vec{n}_1$ und $\vec{n}_2$ ist daher gleich dem Winkel α zwischen ε_1 und ε_2 (Abb. 4.22) oder ergänzt sich mit diesem zu 180° (falls der Winkel zwischen $\vec{n}_1$ und $\vec{n}_2$ größer als 90° ist). In beiden Fällen lässt sich der Winkel der beiden Ebenen ε_1 und ε_2 als Winkel zweier auf diesen Ebenen senkrecht stehender Geraden mit $\vec{n}_1$ und $\vec{n}_2$ als Richtungsvektoren berechnen. Wegen des Satzes 4.3 gilt somit der folgende Satz.

Satz 4.4

Für den Schnittwinkel α zweier sich schneidender Ebenen ε_1 und ε_2 mit den Normalenvektoren $\vec{n}_1$ und $\vec{n}_2$ gilt: $\cos \alpha = \frac{|\vec{n}_1 \cdot \vec{n}_2|}{|\vec{n}_1| \cdot |\vec{n}_2|}$.

Schnittwinkel zwischen einer Geraden und einer Ebene

Unter dem Schnittwinkel einer Geraden g und einer Ebene ε , versteht man den Winkel α zwischen g und derjenigen Geraden h in der Ebene ε , die von allen in ε liegenden Geraden den kleinsten Winkel mit g einschließt (Abb. 4.23). Dieser Winkel α ergänzt sich mit dem Winkel β zwischen g und einer zu ε senkrechten Geraden zu 90° . Falls $\vec{n}$ ein Normalenvektor von ε (und somit Richtungsvektor einer zu ε senkrechten Geraden) und $\vec{a}$ ein Richtungsvektor von g ist, gilt nach dem Satz 4.3 $\cos \beta = \frac{|\vec{n} \cdot \vec{a}|}{|\vec{n}| \cdot |\vec{a}|}$. Da sich die Winkel α und β zu 90° ergänzen und $\cos(90^\circ - \alpha) = \sin \alpha$ ist, gilt der folgende Satz:

Satz 4.5

Falls eine Gerade g mit dem Richtungsvektor $\vec{a}$ eine Ebene ε mit dem Normalenvektor $\vec{n}$ schneidet, so gilt für den Schnittwinkel α zwischen g und ε :

$$\sin \alpha = \frac{|\vec{n} \cdot \vec{a}|}{|\vec{n}| \cdot |\vec{a}|}.$$

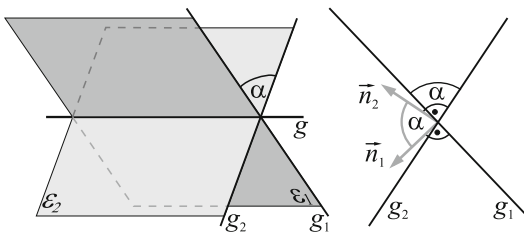


Abb. 4.22: Schnittwinkel zweier Ebenen

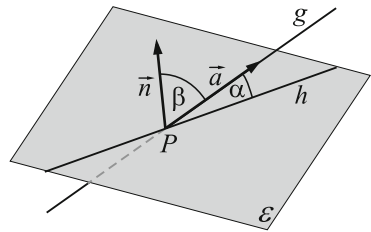


Abb. 4.23: Winkel Gerade-Ebene

4.3.3 Die Hessesche Normalform von Geraden- und Ebenengleichungen; Abstandsberechnungen

Viele Abstandsberechnungen lassen sich recht einfach durchführen, wenn besondere Normalengleichungen genutzt werden, bei denen der Normalenvektor ein *Einheitsvektor* ist.

Normaleneinheitsvektoren von Geraden und Ebenen

Ein Vektor, der den Betrag Eins hat, heißt *Einheitsvektor*. Zu jedem vom Nullvektor verschiedenen Vektor $\vec{a}$ existieren zwei kollineare Einheitsvektoren:

$$\vec{a}_{01} = \frac{1}{|\vec{a}|} \cdot \vec{a} \quad \text{und} \quad \vec{a}_{02} = -\frac{1}{|\vec{a}|} \cdot \vec{a}.$$

Definition 4.4

Ist g eine Gerade in der Ebene, $\vec{n}$ ein Normalenvektor von g und $\vec{n}_0$ ein zu $\vec{n}$ kollinearer Einheitsvektor, so heißt $\vec{n}_0$ *Normaleneinheitsvektor der Geraden g* .

Ist ε eine Ebene im Raum, $\vec{n}$ ein Normalenvektor von ε und $\vec{n}_0$ ein zu $\vec{n}$ kollinearer Einheitsvektor, so heißt $\vec{n}_0$ *Normaleneinheitsvektor der Ebene ε* . ♦

Jede Gerade in der Ebene und jede Ebene im Raum besitzen genau zwei Normaleneinheitsvektoren $\vec{n}_{01}$ und $\vec{n}_{02}$, für die $\vec{n}_{02} = -\vec{n}_{01}$ gilt.

Geraden- und Ebenengleichungen in Hessescher Normalform

Definition 4.5

Eine Normalengleichung einer Geraden g bzw. einer Ebene ε , in welcher der Normalenvektor den Betrag Eins hat (also ein Normaleneinheitsvektor ist), heißt *Geraden- bzw. Ebenengleichung in Hessescher Normalform*:

$$g: \vec{n}_0 \cdot \overrightarrow{P_0 P} = 0 \quad \text{bzw.} \quad \varepsilon: \vec{n}_0 \cdot \overrightarrow{P_0 P} = 0 \quad \text{mit} \quad |\vec{n}_0| = 1$$

(benannt nach Otto Ludwig Hesse, 1811-1874). ♦

Abstand eines Punktes von einer Geraden in der Ebene

Satz 4.6

Ein Punkt Q in der Ebene hat von einer durch eine Gleichung in Hessescher Normalform $\vec{n}_0 \cdot \overrightarrow{P_0 P} = 0$ mit $|\vec{n}_0| = 1$ gegebenen Geraden g den Abstand

$$d(Q, g) = \left| \vec{n}_0 \cdot \overrightarrow{P_0 Q} \right|.$$

Beweis: Der Abstand $d(Q, g)$ eines Punktes Q von einer Geraden g ist die Länge des Lotes von Q auf g . Ist L der Fußpunkt dieses Lotes, so gilt also

$$(1) \quad d(Q, g) = \left| \overrightarrow{LQ} \right|.$$

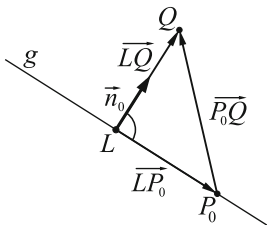


Abb. 4.24:

Abstand eines Punktes von einer Geraden in der Ebene

Der Vektor $\overrightarrow{LQ}$ lässt sich darstellen durch

$$(2) \quad \overrightarrow{LQ} = \overrightarrow{LP_0} + \overrightarrow{P_0Q}.$$

Da die Vektoren $\overrightarrow{LQ}$ und $\overrightarrow{LP_0}$ orthogonal sind, gilt für den Betrag von $\overrightarrow{LQ}$:

$$(3) \quad |\overrightarrow{LQ}|^2 = \overrightarrow{LQ} \cdot \overrightarrow{LQ} = \overrightarrow{LQ} \cdot (\overrightarrow{LP_0} + \overrightarrow{P_0Q}) = \overrightarrow{LQ} \cdot \overrightarrow{LP_0} + \overrightarrow{LQ} \cdot \overrightarrow{P_0Q} = \overrightarrow{LQ} \cdot \overrightarrow{P_0Q}.$$

Wegen $\overrightarrow{LQ} \perp g$ und $\vec{n}_0 \perp g$ sind $\overrightarrow{LQ}$ und $\vec{n}_0$ kollinear. Da außerdem $\vec{n}_0$ den Betrag Eins hat, ist $\vec{n}_0$ ein zu $\overrightarrow{LQ}$ kollinearer Einheitsvektor; es gilt also:

$$(4) \quad \vec{n}_0 = \frac{\overrightarrow{LQ}}{|\overrightarrow{LQ}|} \quad \text{oder} \quad \vec{n}_0 = -\frac{\overrightarrow{LQ}}{|\overrightarrow{LQ}|}.$$

Aus (3) und (4) ergibt sich

$$|\overrightarrow{LQ}| = \frac{\overrightarrow{LQ}}{|\overrightarrow{LQ}|} \cdot \overrightarrow{P_0Q} = \vec{n}_0 \cdot \overrightarrow{P_0Q} \quad \text{oder} \quad |\overrightarrow{LQ}| = \frac{\overrightarrow{LQ}}{|\overrightarrow{LQ}|} \cdot \overrightarrow{P_0Q} = -\vec{n}_0 \cdot \overrightarrow{P_0Q}.$$

In beiden Fällen muss, da $|\overrightarrow{LQ}|$ nicht negativ sein kann, $|\overrightarrow{LQ}| = |\vec{n}_0 \cdot \overrightarrow{P_0Q}|$ gelten, was wegen (1) gerade der Behauptung entspricht. $\square$

Abstand eines Punktes von einer Ebene im Raum

Satz 4.7

Ein Punkt Q hat von einer durch die Gleichung $\vec{n}_0 \cdot \overrightarrow{P_0P} = 0$ mit $|\vec{n}_0| = 1$ gegebenen Ebene ε den Abstand $d(Q, \varepsilon) = |\vec{n}_0 \cdot \overrightarrow{P_0Q}|$.

Der Beweis dieses Satzes kann in völliger Analogie zu dem des Satzes 4.6 geführt werden, siehe Abb. 4.25 sowie die Aufgabe 16 auf S. 166.

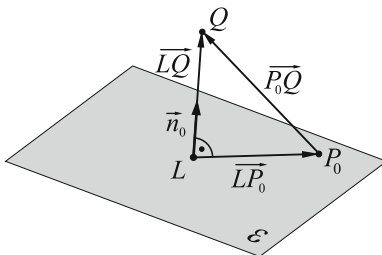


Abb. 4.25: Abstand eines Punktes von einer Ebene

Beispiel 4.26

Abstand des Punktes $Q(4; -7; 0)$ von der Ebene $\varepsilon: \vec{x} = \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 2 \\ 0 \\ -6 \end{pmatrix} + s \begin{pmatrix} -2 \\ 2 \\ 3 \end{pmatrix}$

Um den Abstand $d(Q, \varepsilon)$ mithilfe des Satzes 4.7 zu bestimmen, ist für ε eine Gleichung in Hessescher Normalform aufzustellen. Ein Normalenvektor von ε ist

$$\vec{n} = \begin{pmatrix} 2 \\ 0 \\ -6 \end{pmatrix} \times \begin{pmatrix} -2 \\ 2 \\ 3 \end{pmatrix} = \begin{pmatrix} 12 \\ 6 \\ 4 \end{pmatrix} \quad \text{und ein Normaleneinheitsvektor } \vec{n}_0 = \frac{1}{7} \begin{pmatrix} 6 \\ 3 \\ 2 \end{pmatrix}. \quad \text{Als}$$

Gleichung von ε in Hessescher Normalform ergibt sich: $\varepsilon: \frac{1}{7} \begin{pmatrix} 6 \\ 3 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} x-2 \\ y-1 \\ z-1 \end{pmatrix} = 0$.

Durch Einsetzen der Koordinaten des Punktes Q ergibt sich nach Satz 4.7 der

$$\text{Abstand } d(Q, \varepsilon) = |\vec{n}_0 \cdot \overrightarrow{P_0Q}| = \left| \frac{1}{7} \begin{pmatrix} 6 \\ 3 \\ 2 \end{pmatrix} \cdot \begin{pmatrix} 4-2 \\ -7-1 \\ 0-1 \end{pmatrix} \right| = \frac{|-14|}{7} = 2. \quad \blacksquare$$

Abstand zweier paralleler Ebenen

Sind ε_1 und ε_2 parallele Ebenen, so haben alle Punkte der Ebene ε_1 denselben Abstand von ε_2 ; umgekehrt haben alle Punkte der Ebene ε_2 denselben Abstand von ε_1 . Der Abstand der Ebenen ε_1 und ε_2 lässt sich also berechnen, indem der Abstand eines beliebigen Punktes $P \in \varepsilon_1$ von ε_2 oder der Abstand eines Punktes $Q \in \varepsilon_2$ von ε_1 ermittelt wird, siehe Abb. 4.26 und Aufgabe 20 auf S. 166.

Abstand einer Geraden von einer parallelen Ebene

Ist ε eine Ebene und g eine zu ε parallele Gerade, so haben alle Punkte der Geraden g denselben Abstand von der Ebene ε . Der Abstand einer Geraden g zu einer parallelen Ebene ε (siehe Abb. 4.27) ergibt sich somit als Abstand eines beliebigen Punktes $P \in g$ von ε , siehe Aufgabe 21 auf S. 166.

Abstand eines Punktes von einer Geraden im Raum

Die Berechnung des Abstandes eines Punktes von einer Geraden nach dem auf S. 161 für die Ebene behandelten Verfahren ist im Raum nicht möglich, da sich Geraden im Raum nicht durch Normalengleichungen beschreiben lassen.

Der Abstand $d(Q, g)$ eines Punktes Q von einer Geraden g ist die Länge $|QL|$ des Lotes von Q auf g . Um den Fußpunkt L des Lotes von Q auf g zu bestimmen, ermittelt man zuerst die Gleichung derjenigen Ebene ε , die Q enthält sowie zu g senkrecht ist (siehe Abb. 4.28), und erhält dann L als Schnittpunkt von ε und g . Der gesuchte Abstand $d(Q, g) = |QL|$ kann damit berechnet werden.

Beispiel 4.27

Abstand des Punktes $Q(4; 1; -2)$ von der Geraden $g: \vec{x} = \begin{pmatrix} 2 \\ 5 \\ 2 \end{pmatrix} + t \begin{pmatrix} 2 \\ -1 \\ -1 \end{pmatrix}$

1. Aufstellen einer Normalengleichung der Ebene ε , die zu g senkrecht ist und Q enthält: Wegen $g \perp \varepsilon$ ist der Richtungsvektor von g ein Normalenvektor der Ebene ε . Da weiterhin Q in ε liegen soll, erhält man $\varepsilon: \begin{pmatrix} 2 \\ -1 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} x-4 \\ y-1 \\ z+2 \end{pmatrix} = 0$.
2. Bestimmung des Fußpunktes L des Lotes von Q auf g : Der Fußpunkt des Lotes von Q auf g ist der Schnittpunkt von g und ε . Die Berechnung dieses Schnittpunktes ergibt $L(6; 3; 0)$.

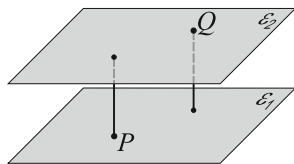


Abb. 4.26: Abstand zweier paralleler Ebenen

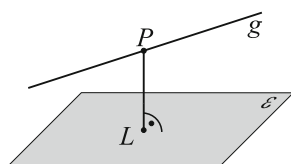


Abb. 4.27: Abstand einer Geraden von einer Ebene

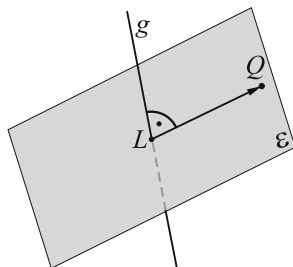


Abb. 4.28: Abstand eines Punktes von einer Geraden

3. Berechnung des gesuchten Abstandes von Q und g :

$$d(Q, g) = |QL| = |\vec{QL}| = \left| \begin{pmatrix} 6-4 \\ 3-1 \\ 0+2 \end{pmatrix} \right| = \sqrt{2^2 + 2^2 + 2^2} = \sqrt{12} \approx 3,464. \quad \blacksquare$$

Der Abstand zweier paralleler Geraden im Raum

Sind g und h parallele Geraden, so haben alle Punkte der Geraden g denselben Abstand von h ; umgekehrt haben alle Punkte von h denselben Abstand von g . Um den Abstand der Geraden g und h zu bestimmen, ist also nur der Abstand eines beliebigen Punktes $P \in g$ von h bzw. eines Punktes $Q \in h$ von g zu ermitteln: $d(g, h) = d(P, h) = d(Q, g)$, $P \in g, Q \in h$, siehe Aufgabe 22 auf S. 166.

Der Abstand zweier windschiefer Geraden

Die Abstandsberechnung windschiefer Geraden lässt sich auf die Berechnung des Abstandes paralleler Ebenen zurückführen. Sind zwei Geraden g und h durch

$$g: \vec{x} = \vec{p}_0 + t\vec{a} \quad \text{und} \quad h: \vec{x} = \vec{q}_0 + s\vec{b}$$

gegeben, so ist

$$\varepsilon_1: \vec{x} = \vec{p}_0 + t\vec{a} + s\vec{b}$$

eine Ebene, welche die Gerade g enthält und

$$\varepsilon_2: \vec{x} = \vec{q}_0 + t\vec{a} + s\vec{b}$$

eine Ebene, die h enthält (Abb. 4.29), wobei ε_1 und ε_2 parallel sind, da sie dieselben Richtungsvektoren haben. Der Abstand $d(g, h)$ der Geraden g und h ist gleich dem Abstand der Ebenen ε_1 und ε_2 und somit gleich dem Abstand des Punktes Q_0 von der Ebene ε_1 bzw. des Punktes P_0 von ε_2 .

Ist $\vec{n}_0$ ein Normaleneinheitsvektor von ε_1 (und damit auch von ε_2), so ergibt sich aus dem Satz 4.7 auf S. 162:

$$d(g, h) = d(Q_0, \varepsilon_1) = d(P_0, \varepsilon_2) = |\vec{n}_0 \cdot \overrightarrow{P_0 Q_0}| = |\vec{n}_0 \cdot (\vec{q}_0 - \vec{p}_0)|.$$

Satz 4.8

Sind $g: \vec{x} = \vec{p}_0 + t\vec{a}$ und $h: \vec{x} = \vec{q}_0 + s\vec{b}$ windschiefe Geraden, so gilt für den

$$\text{Abstand von } g \text{ und } h: d(g, h) = \left| \frac{\vec{a} \times \vec{b}}{|\vec{a} \times \vec{b}|} \cdot (\vec{q}_0 - \vec{p}_0) \right|.$$

Beweis: Da sich ein zu den Richtungsvektoren von g und h orthogonaler Einheitsvektor durch $\vec{n}_0 = \frac{\vec{a} \times \vec{b}}{|\vec{a} \times \vec{b}|}$ berechnen lässt, folgt die Behauptung aus der oberhalb des Satzes 4.8 hergeleiteten Formel. $\square$

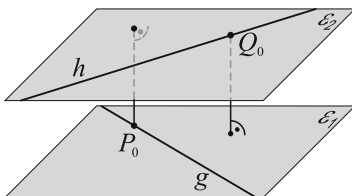


Abb. 4.29: Abstand zweier windschiefer Geraden

4.3.4 Aufgaben zu Abschnitt 4.3

1. Ermitteln Sie für die durch Parametergleichungen gegebenen Geraden einen Normalenvektor und geben Sie eine Normalengleichung an.

a) $g: \vec{p} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} 2 \\ 3 \end{pmatrix}$ b) $g: \vec{p} = \begin{pmatrix} 4 \\ 0 \end{pmatrix} + t \cdot \begin{pmatrix} 1 \\ 0 \end{pmatrix}$

2. Geben Sie eine Parameterdarstellung für g an.

a) $g: \begin{pmatrix} 6 \\ 8 \end{pmatrix} \cdot \begin{pmatrix} x+4 \\ y-1 \end{pmatrix} = 0$ b) $g: \begin{pmatrix} -11 \\ 9 \end{pmatrix} \cdot \begin{pmatrix} x \\ y+1 \end{pmatrix} = 0$

3. Bestimmen Sie eine Koordinatengleichung für die Ebene ε , die durch P_0 verläuft und für die $\vec{n}$ ein Normalenvektor ist.

a) $P_0(-7; 3; 6)$, $\vec{n} = \begin{pmatrix} -5 \\ 5 \\ 8 \end{pmatrix}$ b) $P_0(0; 9; 4)$, $\vec{n} = \begin{pmatrix} -8, 6 \\ -4, 4 \\ 7, 5 \end{pmatrix}$

4. Bestimmen Sie zu den durch Parametergleichungen gegebenen Ebenen einen Normalenvektor und geben Sie eine Normalengleichung an.

a) $\varepsilon: \vec{p} = \begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix} + t \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix} + s \begin{pmatrix} -4 \\ 4 \\ 4 \end{pmatrix}$ b) $\varepsilon: \vec{p} = \begin{pmatrix} 3 \\ 2 \\ 4 \end{pmatrix} + t \begin{pmatrix} -4 \\ -3 \\ 2 \end{pmatrix} + s \begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix}$

5. Bestimmen Sie eine Normalengleichung der Ebene ε , die A , B und C enthält.

a) $A(-5; 3; 4)$, $B(-7; 7; 7)$, $C(0; 3; 2)$ b) $A(6; -7; 9)$, $B(-3; 5; 4)$, $C(11; 6; 0)$

6. Ermitteln Sie eine Parametergleichung für die Ebene ε .

a) $\varepsilon: \begin{pmatrix} -1 \\ -7 \\ 8 \end{pmatrix} \cdot \begin{pmatrix} x-9 \\ y+8 \\ z-6 \end{pmatrix} = 0$ b) $\varepsilon: \begin{pmatrix} -1, 5 \\ -3 \\ -6 \end{pmatrix} \cdot \begin{pmatrix} x \\ y \\ z \end{pmatrix} = 0$ c) $\varepsilon: 33x - 7y - 15z = 22$

7. Berechnen Sie die Schnittwinkel zwischen den Ebenen ε_1 und ε_2 , ε_2 und ε_3 sowie zwischen ε_1 und ε_3 mit

$\varepsilon_1: \begin{pmatrix} -1 \\ -7 \\ 8 \end{pmatrix} \cdot \begin{pmatrix} x-9 \\ y+8 \\ z-6 \end{pmatrix} = 0$, $\varepsilon_2: \begin{pmatrix} -1, 5 \\ -3 \\ -6 \end{pmatrix} \cdot \begin{pmatrix} x \\ y \\ z \end{pmatrix} = 0$, $\varepsilon_3: \begin{pmatrix} -4 \\ 12 \\ 20 \end{pmatrix} \cdot \begin{pmatrix} x+5 \\ y+8 \\ z+9 \end{pmatrix} = 0$.

8. Bestimmen Sie Schnittgerade und Schnittwinkel der Ebenen ε_1 und ε_2 .

$\varepsilon_1: \vec{p} = \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix} + t \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + s \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix}$, $\varepsilon_2: \vec{p} = \begin{pmatrix} -1 \\ -1 \\ -1 \end{pmatrix} + u \begin{pmatrix} -2 \\ 0 \\ 5 \end{pmatrix} + v \begin{pmatrix} 7 \\ 7 \\ 7 \end{pmatrix}$

9. Ermitteln Sie die Koordinaten des Schnittpunktes und den Schnittwinkel der Geraden g mit der Ebene ε .

a) $g: \vec{p} = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}$, $\varepsilon: x + 2y - z = 1$

b) $g: \vec{p} = \begin{pmatrix} 2 \\ 2 \\ 1 \end{pmatrix} + t \begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix}$, $\varepsilon: \vec{p} = \begin{pmatrix} 1 \\ 1 \\ 5 \end{pmatrix} + r \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix} + s \begin{pmatrix} -1 \\ -1 \\ 3 \end{pmatrix}$

10. Berechnen Sie die Winkel zwischen der Ebene ε und den Koordinatenachsen.

a) $\varepsilon: -7x + 12y - z = 15$ b) $\varepsilon: 0,5x + y - 1,5z = 2,5$

11. Weisen Sie nach: Ist $\vec{a}$ ein beliebiger Vektor und $\vec{a}_0$ ein dazugehöriger kollinearer Einheitsvektor, so gilt $\vec{a} \cdot \vec{a}_0 = |\vec{a}|$ oder $\vec{a} \cdot \vec{a}_0 = -|\vec{a}|$.

12. Zeigen Sie, dass sich jede Geradengleichung in Hessescher Normalform in eine Koordinatengleichung $ax + by = c$ mit $a^2 + b^2 = 1$ umformen lässt und

dass umgekehrt jede Gleichung der Form $ax + by = c$ mit $a^2 + b^2 = 1$ eine Gerade mit dem Normaleneinheitsvektor $\vec{n} = \begin{pmatrix} a \\ b \end{pmatrix}$ beschreibt.

- 13.** Weisen Sie nach: Der Abstand eines Punktes $Q(x_Q; y_Q)$ von einer Geraden $g: y = mx + n$ lässt sich durch $d(Q, g) = \frac{1}{\sqrt{m^2 + 1}} \left| \begin{pmatrix} -m \\ 1 \end{pmatrix} \cdot \begin{pmatrix} x_Q \\ y_Q - n \end{pmatrix} \right|$ berechnen.

- 14.** Bestimmen Sie für die Ebene ε eine Gleichung in Hessescher Normalform.

a) $\varepsilon: \vec{p} = \begin{pmatrix} 2 \\ 7 \\ 8 \end{pmatrix} + t \begin{pmatrix} 5 \\ 0 \\ 3 \end{pmatrix} + s \begin{pmatrix} -8 \\ -3 \\ 0 \end{pmatrix}$ b) $\varepsilon: -2x + 6y - 3z = 14$

- 15.** Zeigen Sie, dass sich jede Ebenengleichung in Hessescher Normalform in die Form $ax + by + cz = d$ mit $a^2 + b^2 + c^2 = 1$ bringen lässt und umgekehrt jede Koordinatengleichung $ax + by + cz = d$ mit $a^2 + b^2 + c^2 = 1$ eine Ebene mit dem Normaleneinheitsvektor $\vec{n} = \begin{pmatrix} a \\ b \\ c \end{pmatrix}$ beschreibt.

- 16.** Beweisen Sie den Satz 4.7 auf S. 162.

- 17.** Berechnen Sie den Abstand des Punktes $Q(4; 1; -1)$ von der Ebene ε .

a) $\varepsilon: \begin{pmatrix} -3 \\ 5 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} x \\ y \\ z \end{pmatrix} = 0$ b) $\varepsilon: \vec{x} = \begin{pmatrix} 3 \\ 8 \\ 6 \end{pmatrix} + t \begin{pmatrix} 3 \\ 4 \\ 0 \end{pmatrix} + s \begin{pmatrix} 0 \\ 1 \\ 3 \end{pmatrix}.$

- 18.** Bestimmen Sie den Abstand des Punktes $Q(12; -3; 13)$ von der Ebene ε mit der Koordinatengleichung $-5x - 8y + 2z = 32$.

- 19.** Weisen Sie nach, dass ein Punkt $Q(x_Q; y_Q; z_Q)$ des Raumes von einer Ebene $\varepsilon: ax + by + cz = d$ den Abstand $d(Q, \varepsilon) = \left| \frac{ax_Q + by_Q + cz_Q - d}{\sqrt{a^2 + b^2 + c^2}} \right|$ hat.

- 20.** Bestimmen Sie den Abstand der zueinander parallelen Ebenen ε_1 und ε_2 .

a) $\varepsilon_1: \vec{p} = \begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix} + s \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + t \begin{pmatrix} 8 \\ -6 \\ 1 \end{pmatrix}, \quad \varepsilon_2: \vec{p} = \begin{pmatrix} 0 \\ 3 \\ 3 \end{pmatrix} + s \begin{pmatrix} 14 \\ -12 \\ 0 \end{pmatrix} + t \begin{pmatrix} 0 \\ 6 \\ 7 \end{pmatrix}$
 b) $\varepsilon_1: -3x - 4y + 5z = -20, \quad \varepsilon_2: 4,5x + 6y - 7,5z = 15$

- 21.** Zeigen Sie, dass die Gerade g zu der Ebene ε parallel ist und bestimmen Sie den Abstand von g zu ε .

a) $g: \vec{p} = \begin{pmatrix} 6 \\ 6 \\ 6 \end{pmatrix} + t \begin{pmatrix} 3 \\ 2 \\ -1 \end{pmatrix}, \quad \varepsilon: \begin{pmatrix} 1 \\ 1 \\ 5 \end{pmatrix} \cdot \begin{pmatrix} x-3 \\ y+1 \\ z-5 \end{pmatrix} = 0$
 b) $g: \vec{p} = \begin{pmatrix} 3 \\ -4 \\ 1 \end{pmatrix} + t \begin{pmatrix} 5 \\ -7 \\ 2 \end{pmatrix}, \quad \varepsilon: x + y + z = -68$

- 22.** Überprüfen Sie, ob die Geraden g und h parallel oder windschief sind und berechnen Sie ihren Abstand.

a) $g: \vec{p} = \begin{pmatrix} 3 \\ 2 \\ 0 \end{pmatrix} + s \begin{pmatrix} -2 \\ -1 \\ 2 \end{pmatrix}, \quad h: \vec{p} = \begin{pmatrix} -1 \\ -3 \\ 7 \end{pmatrix} + t \begin{pmatrix} -4 \\ -2 \\ 4 \end{pmatrix}$
 b) $g: \vec{p} = \begin{pmatrix} -1 \\ 0 \\ -1 \end{pmatrix} + t \begin{pmatrix} 2 \\ 1 \\ -2 \end{pmatrix}, \quad h: \vec{p} = \begin{pmatrix} 0 \\ 11 \\ 7 \end{pmatrix} + t \begin{pmatrix} 1 \\ 2 \\ 2 \end{pmatrix}$
 c) $g: \vec{p} = \begin{pmatrix} 3 \\ 2 \\ 5 \end{pmatrix} + t \begin{pmatrix} -3 \\ 1 \\ 1 \end{pmatrix}, \quad h: \vec{p} = \begin{pmatrix} 6 \\ 5 \\ 11 \end{pmatrix} + t \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix}$

5 Vektorräume

Übersicht

5.1	Der Begriff des Vektorraumes, Beispiele	168
5.2	Untervektorräume	172
5.3	Lineare Hüllen, Erzeugendensysteme, lineare Abhängigkeit	178
5.4	Basis und Dimension	187
5.5	Affine Punkträume	203
5.6	Euklidische Vektor- und Punkträume	219

In dem Kapitel 3 wurden Vektoren anhand von Beispielen (Pfeilklassen, n -Tupel) behandelt. Im Folgenden wird nun der allgemeine Begriff des Vektorraumes eingeführt. Während wir uns bisher auf zwei oder drei Dimensionen beschränken mussten, gelingt es mithilfe der Theorie der Vektor- und affinen Punkträume, auch vier- und mehrdimensionale Räume zu untersuchen. Wir werden den Bereich der unmittelbaren Anschauung verlassen und zu Erkenntnissen gelangen, die nur der Mathematik zugänglich sind. Beispielsweise wird sich in dem Abschnitt 5.5 herausstellen, dass sich zwei Ebenen im vierdimensionalen Raum in genau einem Punkt schneiden können.

Vektorräume sind die grundlegenden Strukturen der Linearen Algebra. Ihrem Verständnis kommt eine fundamentale Bedeutung für ein tieferes Eindringen in dieses Gebiet zu. Besonderer Wert wird auf gut nachvollziehbare Beweisführungen gelegt, da ein Verständnis der Beweise zentraler Sätze unverzichtbar für eine erfolgreiche Beschäftigung mit strukturellen Fragen der Linearen Algebra ist. Um das Verständnis nicht zu erschweren, werden in diesem Buch gewisse Einschränkungen hinsichtlich der Allgemeinheit des Vektorraumbegriffs gemacht. So erfolgt eine weitgehende Beschränkung auf endlichdimensionale Vektorräume, lediglich einige Beispiele verdeutlichen, dass es auch Vektorräume unendlicher Dimension gibt. Verallgemeinerungen sind jedoch recht leicht zu vollziehen, wenn die hier vollzogenen Gedankengänge verinnerlicht wurden.

Es wird sich herausstellen, dass Lösungsmengen linearer Gleichungssysteme fundamental für die Untersuchung von Vektor- und affinen Punkträumen sind. Gleichzeitig wird bei der Behandlung von Vektor- und affinen Räumen die in dem Kapitel 1 begonnene Theorie der linearen Gleichungssysteme weiterentwickelt. Der Leserin bzw. dem Leser sei empfohlen, beim Studium des vorliegenden Kapitels mitunter das Kapitel 1 (insbesondere Abschnitt 1.3) heranzuziehen.

5.1 Der Begriff des Vektorraumes, Beispiele

5.1.1 Definition des Begriffs Vektorraum

Definition 5.1

Eine nicht leere Menge V zusammen mit einer inneren Verknüpfung

$$+ : V \times V \rightarrow V, \quad (\vec{u}, \vec{v}) \mapsto \vec{u} + \vec{v}$$

und einer äußeren Verknüpfung

$$\cdot : \mathbb{R} \times V \rightarrow V, \quad (\lambda, \vec{u}) \mapsto \lambda \cdot \vec{u}$$

heißt *reeller Vektorraum* bzw. Vektorraum über dem Körper der reellen Zahlen,¹ falls folgende Bedingungen erfüllt sind:

A1. Für beliebige $\vec{u}, \vec{v} \in V$ gilt $\vec{u} + \vec{v} = \vec{v} + \vec{u}$ (*Kommutativität der Addition*).

A2. Für beliebige $\vec{u}, \vec{v}, \vec{w} \in V$ gilt $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$
(*Assoziativität der Addition*).

A3. Es existiert $\vec{o} \in V$, so dass für alle $\vec{u} \in V$ gilt: $\vec{u} + \vec{o} = \vec{u}$
(*Existenz eines Nullvektors*).

A4. Zu jedem $\vec{u} \in V$ existiert $-\vec{u} \in V$ mit $\vec{u} + (-\vec{u}) = \vec{o}$
(*Existenz eines Gegenvektors zu jedem Vektor*).

S1. Für beliebige $\vec{u} \in V$ gilt $1 \cdot \vec{u} = \vec{u}$.

S2. Für beliebige $\vec{u} \in V$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda \cdot \mu) \cdot \vec{u} = \lambda \cdot (\mu \cdot \vec{u})$
(*Assoziativität der Multiplikation von Vektoren mit reellen Zahlen*).

S3. Für beliebige $\vec{u}, \vec{v} \in V$ und beliebige $\lambda \in \mathbb{R}$ gilt $\lambda \cdot (\vec{u} + \vec{v}) = \lambda \cdot \vec{u} + \lambda \cdot \vec{v}$
(*1. Distributivgesetz*).

S4. Für beliebige $\vec{u} \in V$ und beliebige $\lambda, \mu \in \mathbb{R}$ gilt $(\lambda + \mu) \cdot \vec{u} = \lambda \cdot \vec{u} + \mu \cdot \vec{u}$
(*2. Distributivgesetz*). ◆

Bemerkungen:

- Die Bedingungen $+: V \times V \rightarrow V$ sowie $\cdot: \mathbb{R} \times V \rightarrow V$ beinhalten die *Abgeschlossenheit* der Menge V bezüglich der Verknüpfungen $+$ und $\cdot$, d. h. für beliebige $\vec{u}, \vec{v} \in V$ ist $\vec{u} + \vec{v} \in V$ und für beliebige $\lambda \in \mathbb{R}$, $\vec{u} \in V$ ist $\lambda \cdot \vec{u} \in V$.
- Die Eigenschaften A1-A4 lassen sich dahin gehend zusammenfassen, dass $(V, +)$, d. h. die Menge V zusammen mit der Verknüpfung $+$, eine *Abelsche* (kommutative) *Gruppe* bildet, siehe S. 94f.
- Der Satz 3.10 auf S. 111 wurde unter ausschließlicher Nutzung von Voraussetzungen bewiesen, die in dem Satz 3.9 (S. 110) formuliert wurden und somit auch durch die Definition 5.1 gegeben sind. Daher gelten die in dem Satz 3.10 aufgeführten Eigenschaften und Rechenregeln für beliebige Vektorräume.
- Mithilfe der Definition 5.1 kann der Begriff „Vektor“ nun allgemein definiert werden: Ist $(V, +, \cdot)$ ein Vektorraum, so heißen die Elemente von V Vektoren.

¹Die Elemente von $\mathbb{R}$ haben im Zusammenhang mit einem Vektorraum die Funktion von Skalaren, siehe S. 95. Es gibt auch Vektorräume über anderen Körpern, z. B. über dem Körper der komplexen Zahlen, die wir hier allerdings nicht betrachten.

5.1.2 Beispiele für Vektorräume

Beispiel 5.1

Die Mengen aller *Pfeilklassen der Ebene* sowie aller *Pfeilklassen des Raumes* bilden mit der in dem Abschnitt 3.1.3 definierten Pfeilklassenaddition sowie der Multiplikation von Pfeilklassen mit reellen Zahlen Vektorräume. Mit den Sätzen 3.3 auf S. 93 und 3.5 auf S. 96 wurde bereits nachgewiesen, dass diese Mengen alle in der Definition 5.1 geforderten Eigenschaften besitzen.

Auch die Mengen aller *Verschiebungen der Ebene bzw. des Raumes* (siehe S. 88) bilden Vektorräume, da Verschiebungen durch Pfeilklassen beschrieben und mit diesen identifiziert werden können. ■

Beispiel 5.2

Für jede natürliche Zahl $n > 0$ bildet die Menge $\mathbb{R}^n$ der *n -Tupel reeller Zahlen* mit der in dem Abschnitt 3.2.2 definierten Addition von n -Tupeln sowie der Multiplikation von n -Tupeln mit reellen Zahlen einen Vektorraum (siehe den Satz 3.6 auf S. 104). Unter anderem sind also $\mathbb{R}^2$ und $\mathbb{R}^3$ Vektorräume. ■

Bemerkung: Streng genommen sind die Mengen $\mathbb{R}^2$, $\mathbb{R}^3$ bzw. allgemein $\mathbb{R}^n$ für sich genommen keine Vektorräume, sondern nur im Zusammenhang mit den in dem Abschnitt 3.2.2 definierten Verknüpfungen. Es müsste also heißen $(\mathbb{R}^2, +, \cdot)$, $(\mathbb{R}^3, +, \cdot)$, $\dots$, $(\mathbb{R}^n, +, \cdot)$ sind Vektorräume. Dennoch wird oft die Kurzschreibweise $\mathbb{R}^2$, $\mathbb{R}^3$, $\dots$, $\mathbb{R}^n$ für diese Vektorräume verwendet; die Definition der Verknüpfungen wird dabei als gegeben betrachtet.

Beispiel 5.3

Folgen (a_n) reeller Zahlen besitzen eine Analogie zu n -Tupeln, wobei Zahlenfolgen jedoch unendlich viele Glieder $a_1, a_2, a_3, \dots$ besitzen, während n -Tupel nur aus endlich vielen (nämlich n) Komponenten bestehen. Die Addition von Zahlenfolgen und ihre Multiplikation mit reellen Zahlen lassen sich völlig analog zu der Addition und skalaren Multiplikation von n -Tupeln definieren. Sind $(a_n) := (a_1; a_2; a_3; \dots)$ und $(b_n) := (b_1; b_2; b_3; \dots)$ zwei beliebige Zahlenfolgen und $\lambda \in \mathbb{R}$, so legt man fest:

$$\begin{aligned}(a_n) + (b_n) &:= (a_1 + b_1; a_2 + b_2; a_3 + b_3; \dots), \\ \lambda \cdot (a_n) &:= (\lambda \cdot a_1; \lambda \cdot a_2; \lambda \cdot a_3; \dots).\end{aligned}$$

Es lässt sich (wie für den Vektorraum der n -Tupel reeller Zahlen) leicht zeigen, dass die Menge aller Folgen reeller Zahlen (a_n) mit den so definierten Verknüpfungen ein Vektorraum ist. Der Nullvektor ist dabei die Folge $(o_n) = (0; 0; 0; \dots)$; der Gegenvektor einer Folge $(a_n) := (a_1; a_2; a_3; \dots)$ ist die Folge $-(a_n) := (-a_1; -a_2; -a_3; \dots)$.

Gegenüber den bisher behandelten Vektorräumen weist der Vektorraum der Folgen reeller Zahlen eine Besonderheit auf. Ohne den Begriff der Dimension hier näher zu erläutern (siehe hierzu den Abschnitt 5.4.5) sei erwähnt, dass die Dimension dieses Vektorraumes unendlich ist, während die anderen bisher aufgeführten Vektorräume die Dimensionen 2, 3 bzw. n (mit $n \in \mathbb{N}$) haben. ■

Beispiel 5.4

Die Menge F_I der auf einem Intervall $I = [a; b]$ (mit $a, b \in \mathbb{R}$, $a < b$) definierten reellwertigen Funktionen ist mit den Verknüpfungen

$$+ : F_I \times F_I \rightarrow F_I, \quad (f+g)(x) := f(x) + g(x) \text{ f. a. } x \in I \quad \text{und}$$

$$\cdot : \mathbb{R} \times F_I \rightarrow F_I, \quad (\lambda \cdot f)(x) := \lambda \cdot f(x) \text{ f. a. } x \in I$$

(für beliebige Funktionen $f, g \in F_I$ und beliebige $\lambda \in \mathbb{R}$) ein Vektorraum. Um zu zeigen, dass dies zutrifft, vermerken wir zunächst, dass für beliebige $f, g \in F_I$ und $\lambda \in \mathbb{R}$ durch $f+g$ und $\lambda \cdot f$ ebenfalls auf dem Intervall I definierte Funktionen gegeben sind (Abgeschlossenheit). Auch die Gültigkeit der Eigenschaften A1-S4 lässt sich recht leicht nachweisen.

A1. Für alle $x \in I$ gilt $(f+g)(x) = f(x) + g(x) = g(x) + f(x) = (g+f)(x)$, also ist $f+g = g+f$ für beliebige $f, g \in F_I$.

A2. Wie die Kommutativität lässt sich auch die Assoziativität der Addition von Funktionen auf die Assoziativität der Addition reeller Zahlen (der Funktionswerte) zurückführen. Für alle $x \in I$ ist

$$((f+g)+h)(x) = (f(x)+g(x))+h(x) = f(x)+(g(x)+h(x)) = (f+(g+h))(x),$$

also $(f+g)+h = f+(g+h)$ für beliebige $f, g, h \in F_I$.

Auf analoge Weise lassen sich auch die Eigenschaften S1-S4 auf Rechenregeln reeller Zahlen zurückführen, worauf hier verzichtet wird.

A3. Die Funktion f_0 mit $f_0(x) = 0$ für alle $x \in I$ hat die Eigenschaft des Nullvektors, denn für beliebige $f \in F_I$ ist $(f+f_0)(x) = f(x)+f_0(x) = f(x)$ für alle $x \in I$ und somit $f+f_0 = f$.

A4. Ebenso lässt sich zeigen, dass für eine beliebige Funktion $f \in F_I$ die Funktion $-f$ mit $(-f)(x) = -f(x)$ die Bedingung $f + (-f) = f_0$ erfüllt. ■

Beispiel 5.5

Als Polynome n -ten Grades bezeichnet man Funktionen p in einer Variablen $x \in \mathbb{R}$, die sich folgendermaßen darstellen lassen:

$$p(x) = a_n x^n + \dots + a_2 x^2 + a_1 x + a_0 = \sum_{k=0}^n a_k x^k \quad \text{mit } a_0, \dots, a_n \in \mathbb{R}.$$

(Dabei muss für ein Polynom n -ten Grades $a_n \neq 0$ sein; der Grad eines Polynoms entspricht der höchsten auftretenden Potenz der Variablen x .)

Für jedes $n \in \mathbb{N}$ ist die Menge P_n der Polynome höchstens n -ten Grades, d. h.

$$P_n = \left\{ p \left| p(x) = \sum_{k=0}^n a_k x^k; \quad a_0, \dots, a_n \in \mathbb{R}; \quad x \in \mathbb{R} \right. \right\}$$

mit den durch

$$(p+q)(x) := p(x) + q(x) = \sum_{k=0}^n a_k x^k + \sum_{k=0}^n b_k x^k = \sum_{k=0}^n (a_k + b_k) x^k$$

$$(\lambda \cdot p)(x) := \lambda \cdot p(x) = \lambda \cdot \sum_{k=0}^n a_k x^k = \sum_{k=0}^n \lambda a_k x^k$$

gegebenen Verknüpfungen $+$ und $\cdot$ ein Vektorraum. Es gibt also einen Vektorraum P_1 der linearen Funktionen, einen Vektorraum P_2 der (höchstens) quadratischen Polynome (siehe Aufgabe 4, S. 171), P_3 der Polynome höchstens dritten

Grades usw. Dabei handelt es sich jeweils um Teilmengen des Raumes $F_{\mathbb{R}}$ aller über $\mathbb{R}$ definierten reellwertigen Funktionen (vgl. Beispiel 5.4); man spricht auch von linearen Unterräumen (siehe den folgenden Abschnitt 5.2). ■

Beispiel 5.6

In dem Abschnitt 1.3.1 wurden bereits Matrizen als „Zahlentabellen“ betrachtet, ausführlicher werden diese dann in dem Kapitel 6 behandelt. Es sei hier bereits vermerkt, dass die Menge aller Matrizen $\begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$ mit n Spalten, m Zeilen und $a_{ij} \in \mathbb{R}, 1 \leq i \leq m, 1 \leq j \leq n$ mit den folgendermaßen definierten Verknüpfungen $+$ und $\cdot$ ein Vektorraum ist:

$$\begin{pmatrix} a_{11} & \dots & a_{1n} \\ a_{21} & \dots & a_{2n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} + \begin{pmatrix} b_{11} & \dots & b_{1n} \\ b_{21} & \dots & b_{2n} \\ \vdots & & \vdots \\ b_{m1} & \dots & b_{mn} \end{pmatrix} := \begin{pmatrix} a_{11}+b_{11} & \dots & a_{1n}+b_{1n} \\ a_{21}+b_{21} & \dots & a_{2n}+b_{2n} \\ \vdots & & \vdots \\ a_{m1}+b_{m1} & \dots & a_{mn}+b_{mn} \end{pmatrix},$$

$$\lambda \cdot \begin{pmatrix} a_{11} & \dots & a_{1n} \\ a_{21} & \dots & a_{2n} \\ \vdots & & \vdots \\ a_{m1} & \dots & a_{mn} \end{pmatrix} := \begin{pmatrix} \lambda a_{11} & \dots & \lambda a_{1n} \\ \lambda a_{21} & \dots & \lambda a_{2n} \\ \vdots & & \vdots \\ \lambda a_{m1} & \dots & \lambda a_{mn} \end{pmatrix}.$$

Dass dieses Beispiel alle Bedingungen der Definition 5.1 erfüllt, lässt sich in derselben Weise zeigen wie in dem Abschnitt 3.2.2 für n -Tupel, denn eine $m \times n$ -Matrix lässt sich auch als $m \cdot n$ -Tupel reeller Zahlen auffassen. ■

5.1.3 Aufgaben zu Abschnitt 5.1

1. Weisen Sie nach, dass in jedem Vektorraum V für drei beliebige Vektoren $\vec{p}, \vec{x}, \vec{y} \in V$ gilt: $\vec{p} + \vec{x} = \vec{p} + \vec{y} \Rightarrow \vec{x} = \vec{y}$.
2. Weisen Sie nach, dass für alle Vektoren $\vec{p}, \vec{q}$ eines Vektorraumes V und alle Skalare $\lambda, \mu \in \mathbb{R}$ gilt:
 - a) Falls $\lambda \cdot \vec{p} = \lambda \cdot \vec{q}$ und $\lambda \neq 0$, so ist $\vec{p} = \vec{q}$.
 - b) Falls $\lambda \cdot \vec{p} = \mu \cdot \vec{p}$ und $\vec{p} \neq \vec{0}$, so ist $\lambda = \mu$.
3. Die Menge $\mathbb{R}^n$ aller n -Tupel reeller Zahlen mit den in dem Abschnitt 3.2.2 definierten Verknüpfungen $+$ und $\cdot$ ist, wie festgestellt wurde, ein Vektorraum. Welche der Eigenschaften A1-A4, S1-S4 der Definition 5.1 gelten auch dann, wenn die skalare Multiplikation durch $\lambda \vec{x} = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$ für alle $\vec{x} \in \mathbb{R}^n$ und alle $\lambda \in \mathbb{R}$ definiert wird? Welche der Eigenschaften werden verletzt?
4. Zeigen Sie, dass die Menge $P_2 = \{p \mid p(x) = a_2 x^2 + a_1 x + a_0; a_0, a_1, a_2 \in \mathbb{R}\}$ der Polynome höchstens 2. Grades mit den folgendermaßen definierten Verknüpfungen $+$ und $\cdot$ für beliebige $p, q \in P_2$ mit $p(x) = a_2 x^2 + a_1 x + a_0$ und $q(x) = b_2 x^2 + b_1 x + b_0$ sowie beliebige $\lambda \in \mathbb{R}$ ein Vektorraum ist:

$$(p+q)(x) := p(x) + q(x) = (a_2+b_2)x^2 + (a_1+b_1)x + (a_0+b_0),$$

$$(\lambda \cdot p)(x) := \lambda \cdot p(x) = \lambda a_2 x^2 + \lambda a_1 x + \lambda a_0.$$

5.2 Untervektorräume

5.2.1 Definition, Unterraumkriterium und Beispiele

Definition 5.2

Ist $(V, +, \cdot)$ ein Vektorraum und U eine nicht leere Teilmenge von V , die bezüglich der in V definierten Verknüpfungen $+$ und $\cdot$ selbst ein Vektorraum ist, so heißt U *Unterraum* von V . ♦

Bemerkungen:

- In der Literatur werden für Untervektorräume unterschiedliche Bezeichnungen wie *linearer Unterraum*, *Teilvektorraum* und *linearer Teilraum* verwendet. Wir nutzen im Folgenden die kurze Bezeichnung *Unterraum*.
- Jeder Vektorraum V besitzt zwei Teilmengen, bei denen sofort klar ist, dass es sich um Unterräume handelt. Dabei handelt es sich um V selbst und um die Teilmenge, die nur den Nullvektor von V enthält. Man nennt diese Unterräume *triviale Unterräume*.

Unterraumkriterium

Um für eine Menge M mit zwei Verknüpfungen $+$ und $\cdot$ zu überprüfen, ob es sich um einen Vektorraum handelt, muss die Abgeschlossenheit der Menge bezüglich der beiden Verknüpfungen untersucht und es muss geprüft werden, ob die acht Punkte A1-S4 der Definition 5.1 erfüllt sind. Derselbe Aufwand ist zu betreiben, wenn nach Definition 5.2 entschieden werden soll, ob eine Menge ein Unterraum eines Vektorraumes ist. Viele der in der Definition 5.1 geforderten Eigenschaften sind für eine Menge jedoch schon dadurch erfüllt, dass es sich um eine Teilmenge eines Vektorraumes handelt, so dass bereits wenige Bedingungen genügen, um zu überprüfen, ob eine Teilmenge eines Vektorraumes ein Unterraum ist. Diese Bedingungen sind in dem folgenden Unterraumkriterium zusammengefasst.

Satz 5.1

Es seien $(V, +, \cdot)$ ein Vektorraum und U eine nicht leere Teilmenge von V . U ist genau dann ein Unterraum von V , wenn folgende Bedingungen erfüllt sind:

- U1. U ist abgeschlossen bezüglich der Vektoraddition in V , d. h. für beliebige $\vec{u}, \vec{v} \in U$ ist auch $\vec{u} + \vec{v} \in U$.
- U2. U ist abgeschlossen bezüglich der skalaren Multiplikation in V , d. h. für beliebige $\vec{u} \in U$, $\lambda \in \mathbb{R}$ ist $\lambda \cdot \vec{u} \in U$.

Beweis:

- Ist U ein Unterraum von V , so gelten die Bedingungen U1 und U2 wegen der Abgeschlossenheit von V bezüglich $+$ und $\cdot$, siehe die entsprechende Bemerkung zu Definition 5.1.
- Es sei umgekehrt $U \subset V$, $U \neq \{\}$, und es seien die Bedingungen U1 und U2 erfüllt. Zu zeigen ist, dass $(U, +, \cdot)$ die Bedingungen A1-S4 der Definition 5.1 erfüllt. Für die Rechenregeln A1, A2 und S1-S4 ist dies dadurch gegeben, dass U eine Teilmenge von V ist und dieselben Verknüpfungen angewendet

werden; somit gelten diese Rechenregeln natürlich für alle Vektoren von U , da diese auch Vektoren von V sind. Es ist also nur noch zu zeigen, dass A3 und A4 in U erfüllt sind.

A3. Nach der Voraussetzung $U \neq \{\}$ existiert ein Vektor $\vec{u} \in U$, und wegen U2 gilt $\lambda \cdot \vec{u} \in U$ für beliebige $\lambda \in \mathbb{R}$. Mit $\lambda = 0$ ist also $0 \cdot \vec{u} \in U$. Wegen des Satzes 3.10, 1 (siehe S. 111) ist $0 \cdot \vec{u} = \vec{o}$ und somit $\vec{o} \in U$.

A4. Es sei $\vec{u}$ ein beliebiger Vektor von U . Dann ist nach U2 auch $(-1) \cdot \vec{u} \in U$. Wegen des Satzes 3.10, 5 gilt $(-1) \cdot \vec{u} = -(1 \cdot \vec{u}) = -\vec{u}$, also ist $-\vec{u} \in U$. $\square$

Beispiele für Unterräume

Beispiel 5.7

Ist $\vec{u}_0$ ein beliebiger Vektor eines Vektorraumes V , so ist die Menge aller Vektoren $\vec{u}$ mit $\vec{u} = r \cdot \vec{u}_0$, $r \in \mathbb{R}$, also die *Menge aller zu $\vec{u}_0$ kollinearen Vektoren*, ein Unterraum von V . Mitunter werden derartige Mengen als *Geraden durch den Koordinatenursprung* bezeichnet. Dies ist eine geometrisch sinnvolle Interpretation, da sich $\vec{u} = r \cdot \vec{u}_0 = \vec{o} + r \cdot \vec{u}_0$ als Parameterdarstellung einer Geraden mit dem Stützvektor $\vec{o}$ auffassen lässt (siehe Abschnitt 4.1). Jedoch sind nicht die Ursprungsgeraden selbst (bei denen es sich um Mengen von Punkten handelt), sondern die Mengen der zugehörigen Ortsvektoren Unterräume. $\blacksquare$

Beispiel 5.8

Es seien $\vec{u}$ und $\vec{v}$ zwei beliebige, nicht kollineare, Vektoren eines Vektorraumes V . Wir betrachten die *Menge M aller Vektoren $\vec{w}$, für die $\vec{u}$, $\vec{v}$ und $\vec{w}$ komplanar sind* (siehe Definition 3.8 auf S. 116):

$$M = \{\vec{w} \mid \vec{w} = r \vec{u} + s \vec{v}; r, s \in \mathbb{R}\}.$$

Unter Verwendung des Unterraumkriteriums (Satz 5.1) lässt sich leicht zeigen, dass M ein Unterraum von V ist:

- Die Vektoren $\vec{u}$ und $\vec{v}$ gehören zu M , somit ist $M \neq \{\}$.
- Sind $\vec{w}_1, \vec{w}_2 \in M$, so existieren $r_1, s_1, r_2, s_2 \in \mathbb{R}$ mit $\vec{w}_1 = r_1 \vec{u} + s_1 \vec{v}$ und $\vec{w}_2 = r_2 \vec{u} + s_2 \vec{v}$. Somit gilt

$$\vec{w}_1 + \vec{w}_2 = r_1 \vec{u} + s_1 \vec{v} + r_2 \vec{u} + s_2 \vec{v} = (r_1 + r_2) \vec{u} + (s_1 + s_2) \vec{v},$$

es ist also $\vec{w}_1 + \vec{w}_2 \in M$; die Bedingung U1 des Satzes 5.1 ist daher erfüllt.

- Für beliebige $\vec{w} \in M$ und $\lambda \in \mathbb{R}$ ist $\lambda \vec{w} = \lambda(r \vec{u} + s \vec{v}) = (\lambda r) \vec{u} + (\lambda s) \vec{v} \in M$; die Bedingung U2 ist also ebenfalls erfüllt.

Mengen aller zu zwei gegebenen (nicht kollinearen) Vektoren komplanaren Vektoren sind somit Unterräume. Geometrisch lassen sich diese als Ebenen interpretieren, die durch den Koordinatenursprung verlaufen (vgl. Abschnitt 4.2.1). $\blacksquare$

Beispiel 5.9

Lösungsmengen homogener linearer Gleichungssysteme in n Variablen (siehe Abschnitt 1.3.4) sind Unterräume von $\mathbb{R}^n$. Homogene LGS besitzen mit dem Null- n -Tupel stets eine triviale Lösung. Nach dem Satz 1.1 ist die Summe zweier Lösungen eines homogenen LGS wiederum eine Lösung, und ein beliebiges

reelles Vielfaches einer Lösung eines homogenen LGS ist ebenfalls eine Lösung dieses LGS. Somit sind die Bedingungen des Unterraumkriteriums erfüllt. ■

Beispiel 5.10

Die Menge S_I der auf einem Intervall $I = [a; b]$ (mit $a, b \in \mathbb{R}$, $a < b$) stetigen reellwertigen Funktionen ist ein Unterraum des Vektorraumes F_I der auf I definierten reellwertigen Funktionen (siehe das Beispiel 5.4 auf S. 170). Mithilfe des Unterraumkriteriums lässt sich dies leicht begründen:

- Es ist $S \neq \{\}$, da für jedes Intervall I eine auf I stetige Funktion existiert, z. B. die Funktion f_0 mit $f_0(x) = 0$ f. a. $x \in I$.
- Die Summe zweier auf einem Intervall I stetiger Funktionen ist auf I stetig.
- Das Produkt einer stetigen Funktion mit einer reellen Zahl ist ebenfalls stetig.

Auf völlig analoge Weise lässt sich begründen, dass die Menge der auf einem Intervall differenzierbaren Funktionen ein Unterraum ist. ■

Beispiel 5.11

Bereits in dem Beispiel 5.5 auf S. 170 wurde erwähnt, dass der Vektorraum P_n der Polynome höchstens n -ten Grades ein Unterraum des Raumes $F_{\mathbb{R}}$ aller über $\mathbb{R}$ definierten reellwertigen Funktionen ist. Zugleich ist P_n auch ein Unterraum des Raumes $S_{\mathbb{R}}$ aller über ganz $\mathbb{R}$ stetigen Funktionen sowie des Raumes aller über ganz $\mathbb{R}$ differenzierbaren Funktionen.

Betrachtet man für zwei natürliche Zahlen m, n mit $m < n$ den Vektorraum P_m der Polynome höchstens m -ten Grades sowie den Raum P_n der Polynome höchstens n -ten Grades, so lässt sich mithilfe des Unterraumkriteriums leicht zeigen, dass P_m ein Unterraum von P_n ist; zum Beispiel ist also der Raum der (höchstens) quadratischen Polynome ein Unterraum des Raumes der (höchstens) kubischen Polynome. ■

Beispiel 5.12

Unter einem *magischen Quadrat* der Kantenlänge n versteht man eine quadratische Anordnung der Zahlen $1, 2, \dots, n^2$, bei der die Summen der Zahlen aller Zeilen, Spalten und der beiden Diagonalen gleich sind. Das älteste bekannte magische Quadrat war bereits im 3. Jahrtausend v. Chr. in China bekannt; es hat die Kantenlänge 3 und ist durch folgende Matrix gegeben:

$$\begin{pmatrix} 4 & 9 & 2 \\ 3 & 5 & 7 \\ 8 & 1 & 6 \end{pmatrix}.$$

Verzichtet man auf die Bedingung, dass alle Elemente eines magischen Quadrates natürliche Zahlen sein müssen und lässt auch reelle Zahlen zu und verlangt nur, dass die Summen aller Zeilen, aller Spalten und der beiden Diagonalen gleich sind, so bildet die Menge aller magischen Quadrate der Kantenlänge 3 einen Unterraum des Vektorraumes der 3×3 -Matrizen (siehe das Beispiel 5.6 auf S. 171). Die Summenbedingungen an ein magisches Quadrat der Form

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$$

lassen sich folgendermaßen ausdrücken:

$$a_{11} + a_{12} + a_{13} = a_{21} + a_{22} + a_{23} = a_{31} + a_{32} + a_{33} =$$

$$a_{11} + a_{21} + a_{31} = a_{12} + a_{22} + a_{32} = a_{13} + a_{23} + a_{33} =$$

$$a_{11} + a_{22} + a_{33} = a_{13} + a_{22} + a_{31}.$$

Für die Summe

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} + \begin{pmatrix} b_{11} & b_{12} & b_{13} \\ b_{21} & b_{22} & b_{23} \\ b_{31} & b_{32} & b_{33} \end{pmatrix} = \begin{pmatrix} a_{11} + b_{11} & a_{12} + b_{12} & a_{13} + b_{13} \\ a_{21} + b_{21} & a_{22} + b_{22} & a_{23} + b_{23} \\ a_{31} + b_{31} & a_{32} + b_{32} & a_{33} + b_{33} \end{pmatrix}$$

zweier magischer Quadrate gelten wegen der obigen Gleichung und wegen

$$b_{11} + b_{12} + b_{13} = b_{21} + b_{22} + b_{23} = b_{31} + b_{32} + b_{33} =$$

$$b_{11} + b_{21} + b_{31} = b_{12} + b_{22} + b_{32} = b_{13} + b_{23} + b_{33} =$$

$$b_{11} + b_{22} + b_{33} = b_{13} + b_{22} + b_{31}$$

ebenfalls die Summenbedingungen:

$$(a_{11} + b_{11}) + (a_{12} + b_{12}) + (a_{13} + b_{13}) = (a_{21} + b_{21}) + (a_{22} + b_{22}) + (a_{23} + b_{23}) =$$

$$(a_{31} + b_{31}) + (a_{32} + b_{32}) + (a_{33} + b_{33}) = (a_{11} + b_{11}) + (a_{21} + b_{21}) + (a_{31} + b_{31}) =$$

$$(a_{12} + b_{12}) + (a_{22} + b_{22}) + (a_{32} + b_{32}) = (a_{13} + b_{13}) + (a_{23} + b_{23}) + (a_{33} + b_{33}) =$$

$$(a_{11} + b_{11}) + (a_{22} + b_{22}) + (a_{33} + b_{33}) = (a_{13} + b_{13}) + (a_{22} + b_{22}) + (a_{31} + b_{31}).$$

Somit ist also die Bedingung U1 des Unterraumkriteriums erfüllt. Auf dieselbe Weise lässt sich zeigen, dass auch jedes reelle Vielfache eines magischen Quadrats wiederum ein magisches Quadrat ist. Aufgrund des angegebenen Beispiels ist auch klar, dass mindestens ein magisches Quadrat existiert, die Menge aller magischen Quadrate der Kantenlänge 3 also nicht leer ist. Nach dem Satz 5.1 ist diese Menge somit ein Unterraum.

Magische Quadrate lassen sich auch als Lösungsmengen homogener linearer Gleichungssysteme betrachten. Die Bedingung, dass die Summen aller Zeilen, aller Spalten und der beiden Diagonalen gleich sein müssen, führt zu folgendem LGS mit den Variablen $a_{11}, a_{12}, \dots, a_{33}$ und s (der Summe jeder der Zeilen, Spalten und Diagonalen):

$$\begin{array}{ccccccc} a_{11} + a_{12} + a_{13} & & & & & & - s = 0 \\ & a_{21} + a_{22} + a_{23} & & & & & - s = 0 \\ & & a_{31} + a_{32} + a_{33} & & & & - s = 0 \\ a_{11} & & + a_{21} & & + a_{31} & & - s = 0 \\ & a_{12} & & + a_{22} & & + a_{32} & - s = 0 \\ & & a_{13} & & + a_{23} & & + a_{33} - s = 0 \\ a_{11} & & & + a_{22} & & & + a_{33} - s = 0 \\ & a_{13} & & + a_{22} & & + a_{31} & - s = 0 \end{array}$$

Die Lösungsmenge dieses LGS mit 10 Variablen und 8 Gleichungen ist ein linearer Unterraum von $\mathbb{R}^{10}$ (siehe das Beispiel 5.9 auf S. 173). ■

5.2.2 Der Durchschnitt und die Summe zweier Unterräume

Satz 5.2

Sind U_1, U_2 Unterräume eines Vektorraumes V , so ist der Durchschnitt $U_1 \cap U_2$ ebenfalls ein Unterraum von V .

Beweis: Da jeder Unterraum den Nullvektor enthält, ist $U_1 \cap U_2$ nicht leer. Es bleiben die Bedingungen U_1, U_2 des Unterraumkriteriums zu bestätigen. Dazu seien $v, w \in U_1 \cap U_2$ und $\lambda \in \mathbb{R}$. Da somit $v, w \in U_1$ sind, folgt durch Anwendung von U_1 und U_2 für U_1 , dass $v+w \in U_1$ und $\lambda v \in U_1$ ist. Ebenso gilt $v+w \in U_2$ und $\lambda v \in U_2$. Es ergibt sich also $v+w \in U_1 \cap U_2$ sowie $\lambda v \in U_1 \cap U_2$. $\square$

Beispiel 5.13

Die Lösungsmengen der einzelnen Gleichungen des homogenen LGS

$$\text{I} \quad 3x - 2y + 4z = 0$$

$$\text{II} \quad x + y - 3z = 0$$

sind Unterräume von $\mathbb{R}^3$ (siehe Beispiel 5.9 auf S. 173):

$$L_I = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \middle| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \lambda \cdot \begin{pmatrix} 2 \\ 3 \\ 0 \end{pmatrix} + \mu \cdot \begin{pmatrix} -4 \\ 0 \\ 3 \end{pmatrix}; \lambda, \mu \in \mathbb{R} \right\},$$

$$L_{II} = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \middle| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \lambda \cdot \begin{pmatrix} -1 \\ 1 \\ 0 \end{pmatrix} + \mu \cdot \begin{pmatrix} 3 \\ 0 \\ 1 \end{pmatrix}; \lambda, \mu \in \mathbb{R} \right\}.$$

Berechnet man durch Gleichsetzen dieser Parameterdarstellungen den Durchschnitt der beiden Unterräume (z. B. nach dem in dem Beispiel 4.19 auf S. 154 für Ebenen beschriebenen Vorgehen), so ergibt sich der Unterraum

$$L_I \cap L_{II} = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \middle| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \lambda \cdot \begin{pmatrix} 2 \\ 13 \\ 5 \end{pmatrix}; \lambda \in \mathbb{R} \right\}.$$

Dieser Unterraum ist die Lösungsmenge des gegebenen Gleichungssystems. Geometrisch lassen sich Lösungsmengen homogener Gleichungen mit drei Variablen als Ebenen interpretieren, die durch den Koordinatenursprung verlaufen. Der Durchschnitt zweier derartiger Ebenen ist (wenn diese nicht identisch sind) eine Ursprungsgerade. $\blacksquare$

Bemerkung: Die Vereinigungsmenge $U_1 \cup U_2$ zweier Unterräume U_1 und U_2 eines Vektorraumes ist i. Allg. kein Unterraum, siehe Aufgabe 6 auf S. 177.

Definition 5.3

Es seien U_1, U_2 Unterräume eines Vektorraumes V . Dann heißt die Menge

$$U_1 + U_2 := \{ \vec{u}_1 + \vec{u}_2 \mid \vec{u}_1 \in U_1, \vec{u}_2 \in U_2 \}$$

Summe der Unterräume U_1 und U_2 . $\blacklozenge$

Beispiel 5.14

Wir betrachten zwei Unterräume von $\mathbb{R}^3$:

$$U_1 = \left\{ \begin{pmatrix} r \\ 0 \\ 0 \end{pmatrix} \middle| r \in \mathbb{R} \right\}, \quad U_2 = \left\{ \begin{pmatrix} 0 \\ s \\ 0 \end{pmatrix} \middle| s \in \mathbb{R} \right\},$$

Die Summe dieser beiden Unterräume ist

$$U_1 + U_2 = \left\{ \begin{pmatrix} r \\ 0 \\ 0 \end{pmatrix} + \begin{pmatrix} 0 \\ s \\ 0 \end{pmatrix} \mid r, s \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} r \\ s \\ 0 \end{pmatrix} \mid r, s \in \mathbb{R} \right\}. \quad \blacksquare$$

Satz 5.3

Sind U_1, U_2 Unterräume eines Vektorraumes V , so ist die Summe $U_1 + U_2$ ebenfalls ein Unterraum von V .

Beweis: Es ist klar, dass $U_1 + U_2$ nicht die leere Menge ist, somit bleiben für $U_1 + U_2$ die Bedingungen U1 und U2 des Unterraumkriteriums zu bestätigen. Dazu seien $\vec{v}, \vec{w} \in U_1 + U_2$, d. h. es existieren $\vec{v}_1, \vec{w}_1 \in U_1$ und $\vec{v}_2, \vec{w}_2 \in U_2$ mit

$$\vec{v} = \vec{v}_1 + \vec{v}_2, \quad \vec{w} = \vec{w}_1 + \vec{w}_2.$$

Wegen der Assoziativität und der Kommutativität der Vektoraddition folgt

$$\vec{v} + \vec{w} = (\vec{v}_1 + \vec{v}_2) + (\vec{w}_1 + \vec{w}_2) = (\vec{v}_1 + \vec{w}_1) + (\vec{v}_2 + \vec{w}_2) \in U_1 + U_2.$$

Ist $\lambda \in \mathbb{R}$, so gilt

$$\lambda \vec{v} = \lambda (\vec{v}_1 + \vec{v}_2) = \lambda \vec{v}_1 + \lambda \vec{v}_2 \in U_1 + U_2.$$

Damit sind die Bedingungen U1, U2 erfüllt; $U_1 + U_2$ ist somit ein Unterraum. $\square$

5.2.3 Aufgaben zu Abschnitt 5.2

1. Gegeben sind ein Unterraum U eines Vektorraumes V und Vektoren $\vec{u}, \vec{v} \in V$. Welche der folgenden Aussagen sind richtig?

- a) Gehören $\vec{u}$ und $\vec{v}$ nicht zu U , so ist auch $\vec{u} + \vec{v} \notin U$.
- b) Gehören $\vec{u}$ und $\vec{v}$ nicht zu U , so ist $\vec{u} + \vec{v} \in U$.
- c) Gehört $\vec{u}$ zu U , nicht aber $\vec{v}$, so ist $\vec{u} + \vec{v} \notin U$.

Geben Sie Begründungen oder Gegenbeispiele an.

2. Geben Sie zu folgenden Teilmengen des Vektorraumes $\mathbb{R}^3$ an, ob sie Unterräume sind; begründen Sie Ihre Aussagen:

- a) $U_1 = \left\{ \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \in \mathbb{R}^3 \mid v_1 + v_2 = 2 \right\}$
- b) $U_2 = \left\{ \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \in \mathbb{R}^3 \mid v_1 + v_2 = v_3 \right\}$
- c) $U_3 = \left\{ \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \in \mathbb{R}^3 \mid v_1 \cdot v_2 = v_3 \right\}$

3. Ist die Menge aller konvergenten Zahlenfolgen ein Unterraum des Vektorraumes der Folgen (a_n) reeller Zahlen (siehe Beispiel 5.3 auf S. 169)? Begründen Sie Ihre Antwort.

4. Begründen Sie, dass Lösungsmengen inhomogener linearer Gleichungssysteme keine Unterräume sind.

5. Weisen Sie nach, dass die Menge aller 2×2 -Matrizen der Form $\begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix}$ mit $a, b \in \mathbb{R}$ mit den in dem Beispiel 5.6 auf S. 171 definierten Verknüpfungen $+$ und $\cdot$ ein Unterraum des Vektorraumes aller Matrizen $\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ mit 2 Zeilen und 2 Spalten ist.

6. Zeigen Sie, dass die Vereinigungsmenge $U_1 \cup U_2$ zweier Unterräume U_1 und U_2 eines Vektorraumes im Allgemeinen kein Unterraum ist.

5.3 Lineare Hüllen, Erzeugendensysteme, lineare Abhängigkeit

In dem Abschnitt 3.4 wurden *Linearkombinationen* von k Vektoren $\vec{u}_1, \dots, \vec{u}_k$ als Vektoren $\vec{x}$ definiert, die sich in der Form $\vec{x} = \sum_{i=1}^k \lambda_i \vec{u}_i$ mit $\lambda_1, \dots, \lambda_k \in \mathbb{R}$ darstellen lassen (siehe Definition 3.7 auf S. 113). Im Folgenden bilden Linearkombinationen die Grundlage für einige zentrale Begriffe der linearen Algebra.

5.3.1 Lineare Hüllen von Vektormengen

Definition 5.4

Als *lineare Hülle* von Vektoren $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ bezeichnet man die Menge aller Linearkombinationen dieser Vektoren:

$$\langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle = \left\{ \sum_{i=1}^k \lambda_i \cdot \vec{u}_i \mid \lambda_1, \dots, \lambda_k \in \mathbb{R} \right\}. \quad \blacklozenge$$

Beispiel 5.15

Die lineare Hülle der Vektoren $\vec{u}_1 = \begin{pmatrix} -3 \\ 2 \\ 4 \end{pmatrix}$ und $\vec{u}_2 = \begin{pmatrix} 5 \\ -1 \\ 7 \end{pmatrix}$ ist die Menge

$$\langle \vec{u}_1, \vec{u}_2 \rangle = \left\{ \lambda_1 \cdot \begin{pmatrix} -3 \\ 2 \\ 4 \end{pmatrix} + \lambda_2 \cdot \begin{pmatrix} 5 \\ -1 \\ 7 \end{pmatrix} \mid \lambda_1, \lambda_2 \in \mathbb{R} \right\}. \quad \blacksquare$$

Satz 5.4

Es seien $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ Vektoren eines Vektorraumes V . Dann ist die lineare Hülle $\langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle$ ein Unterraum von V .

Beweis: Da die Vektoren $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ selbst zu $\langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle$ gehören, ist die lineare Hülle nicht leer. Es genügt also, zu zeigen, dass die beiden Bedingungen U1 und U2 des Unterraumkriteriums (Satz 5.1, S. 172) erfüllt sind.

U1. Es seien $\vec{v}$ und $\vec{w}$ zwei Vektoren der linearen Hülle $\langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle$, d. h. es existieren $\mu_1, \dots, \mu_k \in \mathbb{R}$ sowie $\nu_1, \dots, \nu_k \in \mathbb{R}$ mit $\vec{v} = \sum_{i=1}^k \mu_i \vec{u}_i$ und $\vec{w} = \sum_{i=1}^k \nu_i \vec{u}_i$. Mit $\lambda_i = \mu_i + \nu_i$ (für $i = 1 \dots k$) ist $\vec{v} + \vec{w} = \sum_{i=1}^k \lambda_i \vec{u}_i$ und somit nach Definition 5.4 $\vec{v} + \vec{w} \in \langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle$.

U2. Ist $\vec{v} = \sum_{i=1}^k \mu_i \vec{u}_i$ und $r \in \mathbb{R}$, so ist $r\vec{v} = \sum_{i=1}^k r\mu_i \vec{u}_i$ und somit nach Definition 5.4 $r\vec{v} \in \langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle$. $\square$

Satz 5.5

Es seien $\vec{u}_1, \dots, \vec{u}_k$ Vektoren eines Vektorraumes V . Dann ist die lineare Hülle $\langle \vec{u}_1, \dots, \vec{u}_k \rangle$ der kleinste Unterraum von V , der $\vec{u}_1, \dots, \vec{u}_k$ enthält, d. h.: Ist U ein beliebiger Unterraum von V mit $\vec{u}_1 \in U, \dots, \vec{u}_k \in U$, so gilt $\langle \vec{u}_1, \vec{u}_2, \dots, \vec{u}_k \rangle \subseteq U$.

Beweis: Es sei U ein beliebiger Unterraum von V , der die Vektoren $\vec{u}_1, \dots, \vec{u}_k$ enthält. Es ist zu zeigen, dass jeder Vektor $\vec{v} \in \langle \vec{u}_1, \dots, \vec{u}_k \rangle$ zu U gehört. Sei dazu $\vec{v} = \sum_{i=1}^k \mu_i \vec{u}_i$ mit $\mu_1, \dots, \mu_k \in \mathbb{R}$. Da $\vec{u}_1 \in U, \vec{u}_2 \in U, \dots, \vec{u}_k \in U$, gilt wegen des Satzes 5.1, U2: $\mu_1 \vec{u}_1 \in U, \mu_2 \vec{u}_2 \in U, \dots, \mu_k \vec{u}_k \in U$. Nach Satz 5.1, U1 gehört zu zwei beliebigen Vektoren eines Unterraumes U auch deren Summe zu U . Durch mehrfache Anwendung dieser Bedingung folgt, dass auch die Summe der k Vektoren $\mu_1 \vec{u}_1, \mu_2 \vec{u}_2, \dots, \mu_k \vec{u}_k$ ein Vektor von U ist, also $\vec{v} \in U$. $\square$

5.3.2 Erzeugendensysteme

In den vorangegangenen Betrachtungen über lineare Hüllen von Vektormengen wurden Mengen von Vektoren untersucht, die aus gegebenen Vektoren als Linearkombinationen erzeugt werden können, und es wurde dabei festgestellt, dass die lineare Hülle einer beliebigen Menge von Vektoren stets ein Unterraum und somit ein Vektorraum ist. Die folgenden Überlegungen schließen daran an, indem Mengen von Vektoren betrachtet werden, die *jeden* Vektor eines Vektorraumes als Linearkombination erzeugen können.

Definition 5.5

Eine Teilmenge E eines Vektorraumes V heißt *Erzeugendensystem* von V , wenn sich jeder Vektor $\vec{v} \in V$ als Linearkombination von Vektoren aus E darstellen lässt. $\blacklozenge$

Beispiel 5.16

Die Vektormenge $E_1 = \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \end{pmatrix}; \begin{pmatrix} 1 \\ 2 \end{pmatrix} \right\}$ ist ein Erzeugendensystem von $\mathbb{R}^2$.

Um dies zu bestätigen, muss gezeigt werden, dass sich jeder Vektor $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ als Linearkombination der drei Vektoren von E_1 darstellen lässt, also $\lambda, \mu, \nu \in \mathbb{R}$ existieren mit $\begin{pmatrix} x \\ y \end{pmatrix} = \lambda \begin{pmatrix} 1 \\ 0 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ 1 \end{pmatrix} + \nu \begin{pmatrix} 1 \\ 2 \end{pmatrix}$. Dies ist gleichbedeutend mit der Lösbarkeit des linearen Gleichungssystems

$$\begin{aligned} \lambda + \mu + \nu &= x \\ \mu + 2\nu &= y. \end{aligned}$$

Dieses LGS ist für beliebige $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ lösbar (wobei für jeden Vektor $\begin{pmatrix} x \\ y \end{pmatrix}$ sogar unendlich viele Lösungen existieren). E_1 ist somit ein Erzeugendensystem von $\mathbb{R}^2$. Allerdings würden schon zwei der drei Vektoren von E_1 ausreichen, um $\mathbb{R}^2$ zu erzeugen. Auch $E_2 = \left\{ \begin{pmatrix} 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right\}$ ist ein Erzeugendensystem von $\mathbb{R}^2$. Dies ergibt sich daraus, dass das LGS

$$\begin{aligned} \lambda + \mu &= x \\ \mu &= y \end{aligned}$$

für beliebige $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ lösbar ist, jeder Vektor von $\mathbb{R}^2$ also als Linearkombination der Vektoren $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$ und $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$ dargestellt werden kann. $\blacksquare$

Das Beispiel zeigt, dass $\mathbb{R}^2$ ein Erzeugendensystem besitzt, das aus zwei Vektoren besteht; weitere derartige Erzeugendensysteme lassen sich leicht konstruieren. Analog dazu gibt es für $\mathbb{R}^3$ Erzeugendensysteme mit drei Vektoren sowie allgemein für $\mathbb{R}^n$ Erzeugendensysteme mit n Vektoren, z. B. $E = \{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\}$

mit $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$, $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}$, $\vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix}$, $\dots$, $\vec{e}_n = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{pmatrix}$. Jedoch gibt es kein Erzeugendensystem von $\mathbb{R}^n$, das aus weniger als n Vektoren besteht. Auf einen Beweis dieser Tatsache wird hier verzichtet, das folgende Beispiel verdeutlicht jedoch, dass Mengen von 2 Vektoren nur echte Unterräume von $\mathbb{R}^3$ erzeugen.

Beispiel 5.17

Zwei beliebige, nicht kollineare Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^3$ erzeugen stets den Unterraum U , der alle Vektoren $\vec{x}$ enthält, die mit $\vec{u}$ und $\vec{v}$ komplanar sind:

$$U = \langle \vec{u}, \vec{v} \rangle = \left\{ \vec{x} \mid \vec{x} = \lambda_1 \cdot \begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix} + \lambda_2 \cdot \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix}; \lambda_1, \lambda_2 \in \mathbb{R} \right\}.$$

Geometrisch lässt sich dieser Unterraum als Ebene interpretieren, die durch den Koordinatenursprung verläuft. In Abb. 5.1 sind für $\vec{u} = \begin{pmatrix} 4 \\ -2 \\ -2 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} -4 \\ 6 \\ 2 \end{pmatrix}$ einige Vektoren des durch $\vec{u}$ und $\vec{v}$ erzeugten Unterraumes durch Pfeile (mit Anfangspunkten im Koordinatenursprung) dargestellt. ■

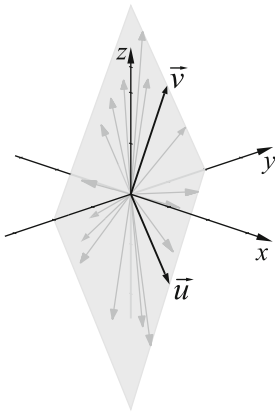


Abb. 5.1: Einige Vektoren des durch zwei Vektoren $\vec{u}$ und $\vec{v}$ erzeugten Unterraumes

Es ist leicht zu überlegen, dass zwei kollineare Vektoren $\vec{u}$ und $\vec{v}$ denselben Vektorraum erzeugen wie ein einziger dieser beiden Vektoren. In Verallgemeinerung dieser Tatsache darf aus einem Erzeugendensystem eines Vektorraumes V in dem zwei kollineare Vektoren enthalten sind, ein Vektor „herausgenommen“ werden, wobei die verbleibenden Vektoren immer noch ein Erzeugendensystem von V bilden. Ebenso kann aus einem Erzeugendensystem, das drei komplanare Vektoren enthält, ein Vektor entfernt werden. Der folgende Satz bringt diese Tatsachen in präziseren Formulierungen zum Ausdruck.

Satz 5.6

Es sei $E = \{\vec{e}_1; \vec{e}_2; \dots; \vec{e}_k\}$ ein Erzeugendensystem eines Vektorraumes V .

- Sind zwei Vektoren $\vec{e}_i, \vec{e}_j \in E$ ($i \neq j$) kollinear und ist $\vec{e}_i \neq \vec{o}$, so ist auch $E \setminus \{\vec{e}_j\} = \{\vec{e}_1; \dots; \vec{e}_{j-1}; \vec{e}_{j+1}; \dots; \vec{e}_k\}$ ein Erzeugendensystem von V .
- Sind drei Vektoren $\vec{e}_i, \vec{e}_j, \vec{e}_l \in E$ ($i \neq j, i \neq l, j \neq l$) komplanar und sind $\vec{e}_i$ und $\vec{e}_j$ nicht kollinear, so ist auch $E \setminus \{\vec{e}_l\} = \{\vec{e}_1; \dots; \vec{e}_{l-1}; \vec{e}_{l+1}; \dots; \vec{e}_k\}$ ein Erzeugendensystem von V .

Beweis: Auf den Beweis von Teil a wird hier verzichtet (siehe Aufgabe 4 auf S. 186). Um den Beweis von b einfacher aufschreiben zu können, wird zunächst angemerkt, dass die Reihenfolge, in der die Vektoren eines Erzeugendensystems nummeriert sind, unbedeutend ist. Ohne damit die Allgemeinheit der Beweisführung zu schmälern, darf daher eine Umnummerierung vorgenommen werden.

Wir nehmen deshalb im Folgenden an, dass die drei nach Voraussetzung komplanaren Vektoren $\vec{e}_1, \vec{e}_2$ und $\vec{e}_3$ seien und darunter die Vektoren $\vec{e}_2$ und $\vec{e}_3$ nicht kollinear sind. Dann lässt sich $\vec{e}_1$ als Linearkombination von $\vec{e}_2$ und $\vec{e}_3$ darstellen, es existieren also $\mu, \nu \in \mathbb{R}$ mit

$$\vec{e}_1 = \mu \vec{e}_2 + \nu \vec{e}_3.$$

Da $E = \{\vec{e}_1; \dots; \vec{e}_k\}$ ein Erzeugendensystem von V ist, existieren für jeden Vektor $\vec{x} \in V$ reelle Zahlen $\lambda_1, \dots, \lambda_k$ mit:

$$\begin{aligned} \vec{x} &= \sum_{i=1}^k \lambda_i \vec{e}_i = \lambda_1 \vec{e}_1 + \lambda_2 \vec{e}_2 + \lambda_3 \vec{e}_3 + \lambda_4 \vec{e}_4 + \dots + \lambda_k \vec{e}_k \\ &= \lambda_1 (\mu \vec{e}_2 + \nu \vec{e}_3) + \lambda_2 \vec{e}_2 + \lambda_3 \vec{e}_3 + \lambda_4 \vec{e}_4 + \dots + \lambda_k \vec{e}_k \\ &= (\lambda_1 \mu + \lambda_2) \vec{e}_2 + (\lambda_1 \nu + \lambda_3) \vec{e}_3 + \lambda_4 \vec{e}_4 + \dots + \lambda_k \vec{e}_k. \end{aligned}$$

Somit ist $\vec{x}$ als Linearkombination der Vektoren $\vec{e}_2, \dots, \vec{e}_k$ darstellbar; durch Einführung neuer Bezeichnungen $\lambda'_2 = \lambda_1 \mu + \lambda_2$, $\lambda'_3 = \lambda_1 \nu + \lambda_3$, $\lambda'_4 = \lambda_4, \dots, \lambda'_k = \lambda_k$ lässt sich dies besonders deutlich aufschreiben:

$$\vec{x} = \sum_{i=2}^k \lambda'_i \vec{e}_i.$$

Da für $\vec{x}$ ein beliebiger Vektor aus V betrachtet wurde, lässt sich jeder Vektor des Vektorraumes V als Linearkombination von $\vec{e}_2, \dots, \vec{e}_k$ ausdrücken; somit ist $\{\vec{e}_2; \dots; \vec{e}_k\} = E \setminus \{\vec{e}_1\}$ ein Erzeugendensystem von V . Aufgrund der eingangs gemachten Bemerkung zur Umnummerierung der Vektoren $\vec{e}_l, \vec{e}_i, \vec{e}_j$ in $\vec{e}_1, \vec{e}_2, \vec{e}_3$ ist der Teil b des Satzes damit bewiesen. $\square$

Bei der Formulierung des Satzes 5.6 wurde von einem Erzeugendensystem ausgegangen, das aus nur *endlich vielen Vektoren* $\vec{e}_1, \dots, \vec{e}_k$ besteht, obwohl der Satz auch für Erzeugendensysteme mit unendlich vielen Vektoren gilt. Um Schwierigkeiten mit Schreibweisen und Fallunterscheidungen zu vermeiden, verzichten wir darauf, unendliche Erzeugendensysteme zu betrachten. Es sei jedoch angemerkt, dass es Vektorräume gibt, die keine endlichen Erzeugendensysteme besitzen.

Beispiel 5.18

Während der Vektorraum P_n der Polynome höchstens n -ten Grades (siehe das Beispiel 5.5 auf S. 170) ein endliches Erzeugendensystem besitzt, nämlich z. B. $E = \{p_0; p_1; \dots; p_n\}$ mit $p_0(x) = 1$, $p_1(x) = x$, $p_2(x) = x^2$, $\dots$, $p_n(x) = x^n$, gibt es kein endliches Erzeugendensystem für den Vektorraum aller Polynome. Egal aus wie vielen Polynomen eine endliche Menge bestünde, würde man immer ein Polynom mit genügend hohem Grad finden, welches nicht als Linearkombination von Polynomen dieser Menge dargestellt werden kann. Ein Erzeugendensystem dieses Vektorraumes muss also unendlich viele Polynome enthalten, z. B. p_i mit $p_i(x) = x^i$ für alle $i \in \mathbb{N}$ (einschließlich Null). Dieses Erzeugendensystem ist gleichmächtig mit der Menge der natürlichen Zahlen, man nennt es *abzählbar unendlich*. Hingegen besitzt z. B. der Vektorraum F_I der auf einem Intervall I definierten reellwertigen Funktionen (siehe Beispiel 5.4 auf S. 170) nicht nur kein endliches, sondern auch kein abzählbar unendliches Erzeugendensystem. $\blacksquare$

5.3.3 Lineare Abhängigkeit und Unabhängigkeit

In dem Abschnitt 3.4.2 wurde die Komplanarität von Vektoren behandelt und dabei festgestellt, dass drei Vektoren genau dann komplanar sind, wenn sich der Nullvektor aus ihnen auf nicht triviale Weise als Linearkombination darstellen lässt (Satz 3.12 auf S. 117). Diese Grundidee wird nun für Mengen von beliebig vielen Vektoren verallgemeinert.

Definition 5.6

Eine Menge $\{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ von Vektoren eines Vektorraumes V heißt *linear abhängig*, wenn sich der Nullvektor auf nicht triviale Weise als Linearkombination der Vektoren $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ darstellen lässt, d. h. wenn $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{R}$ existieren, von denen für mindestens ein i (mit $1 \leq i \leq k$) $\lambda_i \neq 0$ ist, so dass gilt:

$$\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{0}.$$

Eine Vektormenge $\{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ heißt *linear unabhängig*, wenn sich der Nullvektor nur auf triviale Weise als Linearkombination der Vektoren $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_k$ darstellen lässt, d. h. wenn aus $\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{0}$ folgt $\lambda_1 = \lambda_2 = \dots = \lambda_k = 0$. ♦

Beispiel 5.19

Ein Vergleich der Definition 5.6 mit dem Satz 3.12 auf S. 117 ergibt, dass eine Menge von drei Vektoren genau dann linear abhängig ist, wenn die Vektoren komplanar sind. Somit müssten die drei in dem Beispiel 3.6 auf S. 116 betrachteten Vektoren $\vec{x} = \begin{pmatrix} 3,5 \\ 5 \\ 0,5 \end{pmatrix}$, $\vec{u} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$ und $\vec{w} = \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix}$ linear abhängig sein. Um dies anhand von Definition 5.6 zu überprüfen, wird die Vektorgleichung

$$\lambda \begin{pmatrix} 3,5 \\ 5 \\ 0,5 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \nu \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

bzw. das entsprechende lineare Gleichungssystem

$$\begin{aligned} 3,5\lambda + \mu + \nu &= 0 \\ 5\lambda + 2\mu + \nu &= 0 \\ 0,5\lambda + 3\mu - 2\nu &= 0. \end{aligned}$$

gelöst. Es ergibt sich eine einparametrische Lösungsmenge mit $\lambda = -\frac{1}{2}t$, $\mu = \frac{3}{4}t$ und $\nu = t$. Somit ist also z. B. durch

$$-2 \cdot \begin{pmatrix} 3,5 \\ 5 \\ 0,5 \end{pmatrix} + 3 \cdot \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + 4 \cdot \begin{pmatrix} 1 \\ 1 \\ -2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

eine nicht triviale Linearkombination des Nullvektors durch die Vektoren $\vec{x}$, $\vec{u}$ und $\vec{w}$ gegeben, die Vektoren sind also linear abhängig. ■

Beispiel 5.20

Um zu prüfen, ob die Vektoren $\vec{a} = \begin{pmatrix} 2 \\ 1 \\ 0 \\ -1 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 1 \\ 0 \\ 1 \\ -2 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} -1 \\ -2 \\ 0 \\ 1 \end{pmatrix}$, $\vec{d} = \begin{pmatrix} -1 \\ 0 \\ 1 \\ 0 \end{pmatrix}$

linear abhängig sind, wird das der Vektorgleichung $\lambda_1 \vec{a} + \lambda_2 \vec{b} + \lambda_3 \vec{c} + \lambda_4 \vec{d} = \vec{0}$ entsprechende lineare Gleichungssystem

$$\begin{array}{rclcl}
2\lambda_1 + \lambda_2 - \lambda_3 - \lambda_4 & = & 0 \\
\lambda_1 & - & 2\lambda_3 & = & 0 \\
& \lambda_2 & + & \lambda_4 & = & 0 \\
-\lambda_1 - 2\lambda_2 + \lambda_3 & = & 0
\end{array}$$

gelöst. Es ergibt sich die eindeutige Lösung $\lambda_1 = 0, \lambda_2 = 0, \lambda_3 = 0, \lambda_4 = 0$. Aus $\lambda_1 \vec{a} + \lambda_2 \vec{b} + \lambda_3 \vec{c} + \lambda_4 \vec{d} = \vec{0}$ folgt somit $\lambda_1 = \lambda_2 = \lambda_3 = \lambda_4 = 0$; nach Definition 5.6 sind die Vektoren $\vec{a}, \vec{b}, \vec{c}$ und $\vec{d}$ linear unabhängig. ■

Im Folgenden werden einige Eigenschaften von linear abhängigen und linear unabhängigen Mengen von Vektoren bewiesen, die teilweise selbstverständlich erscheinen mögen, aber für den folgenden Aufbau der Theorie der Vektorräume häufig gewissermaßen als „Handwerkszeug“ benötigt werden.

Satz 5.7

Enthält eine Menge M von Vektoren den Nullvektor, so ist M linear abhängig.

Beweis: Sei $M = \{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ und $\vec{u}_l = \vec{0}$ (mit $1 \leq l \leq k$). Es gilt

$$\begin{aligned}
& 0 \vec{u}_1 + \dots + 0 \vec{u}_{l-1} + 1 \vec{0} + 0 \vec{u}_{l+1} + \dots + 0 \vec{u}_k = \vec{0} \quad \text{bzw.} \\
& \sum_{i=1}^k \lambda_i \vec{u}_i = \lambda_1 \vec{u}_1 + \dots + \lambda_{l-1} \vec{u}_{l-1} + \lambda_l \vec{u}_l + \lambda_{l+1} \vec{u}_{l+1} + \dots + \lambda_k \vec{u}_k = \vec{0}
\end{aligned}$$

mit $\lambda_1 = \dots = \lambda_{l-1} = 0, \lambda_l = 1, \lambda_{l+1} = \dots = \lambda_k = 0$. Der Nullvektor lässt sich also auf nicht triviale Weise als Linearkombination der Vektormenge M darstellen, somit ist M linear abhängig. □

Wir merken an, dass insbesondere die Menge $M = \{\vec{0}\}$, die nur aus dem Nullvektor besteht, trivialerweise linear abhängig ist.

Satz 5.8

- a) *Für eine beliebige linear unabhängige Menge M von Vektoren ist jede Teilmenge $N \subset M$ ebenfalls linear unabhängig.*
b) *Sind M und N Mengen von Vektoren mit $M \subset N$ und ist M linear abhängig, so ist N ebenfalls linear abhängig.*

Beweis: Wir beschränken uns auf die Betrachtung endlicher Vektormengen.

- a) Es sei $M = \{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ und $N \subset M$. Da die Vektoren von M umnummeriert werden können, bedeutet es keine Einschränkung der Allgemeinheit, wenn wir schreiben $N = \{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_l\}$ mit $l < k$. Da M linear unabhängig ist, folgt aus

$$\sum_{i=1}^k \lambda_i \vec{u}_i = \sum_{i=1}^l \lambda_i \vec{u}_i + \sum_{i=l+1}^k \lambda_i \vec{u}_i = \vec{0},$$

dass $\lambda_1 = \lambda_2 = \dots = \lambda_l = \lambda_{l+1} = \dots = \lambda_k = 0$ sein muss. Dazu im Widerspruch stünde die Existenz von $(\lambda_1; \lambda_2; \dots; \lambda_l) \neq (0; 0; \dots; 0)$ mit $\sum_{i=1}^l \lambda_i \vec{u}_i = \vec{0}$, denn damit ließe sich $\lambda_{l+1} = \dots = \lambda_k = 0$ setzen und $\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{0}$ wäre mit diesen Koeffizienten eine nicht triviale Linearkombination des Nullvektors aus Vektoren der Menge M . Also folgt aus $\sum_{i=1}^l \lambda_i \vec{u}_i = \vec{0}$, dass $\lambda_1 = \dots = \lambda_l = 0$ gilt, also die lineare Unabhängigkeit von N .

- b) Dieser Teil des Beweises ist Gegenstand der Aufgabe 7 (S. 186). □

Satz 5.9

Eine Menge $M = \{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ von Vektoren ist genau dann linear abhängig, wenn ein Vektor in M existiert, der sich als Linearkombination der restlichen Vektoren von M darstellen lässt.

Beweis: Da es sich bei dem Satz 5.9 um eine „genau-dann-wenn-Aussage“ handelt, sind zwei Richtungen zu beweisen.

- Ist eine Menge $M = \{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\}$ linear abhängig, so existieren Koeffizienten $\lambda_1, \dots, \lambda_k$ mit $(\lambda_1; \lambda_2; \dots; \lambda_k) \neq (0; 0; \dots; 0)$ und $\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{o}$. Ist λ_l (mit $1 \leq l \leq k$) ein von Null verschiedener Koeffizient, so lässt sich diese Gleichung folgendermaßen umstellen:

$$\begin{aligned}\vec{o} &= \sum_{i=1}^{l-1} \lambda_i \vec{u}_i + \lambda_l \vec{u}_l + \sum_{i=l+1}^k \lambda_i \vec{u}_i \\ \lambda_l \vec{u}_l &= - \sum_{i=1}^{l-1} \lambda_i \vec{u}_i - \sum_{i=l+1}^k \lambda_i \vec{u}_i \\ \vec{u}_l &= \sum_{i=1}^{l-1} \left(-\frac{\lambda_i}{\lambda_l}\right) \vec{u}_i + \sum_{i=l+1}^k \left(-\frac{\lambda_i}{\lambda_l}\right) \vec{u}_i.\end{aligned}$$

Somit ist also der Vektor $\vec{u}_l$ als Linearkombination der Vektormenge $M \setminus \{\vec{u}_l\} = \{\vec{u}_1; \dots; \vec{u}_{l-1}; \vec{u}_{l+1}; \dots; \vec{u}_k\}$ darstellbar.

- Ist umgekehrt ein Vektor $\vec{u}_l \in M$ ($1 \leq l \leq k$) als Linearkombination der restlichen Vektoren von M darstellbar, d. h. existieren $\lambda_1, \dots, \lambda_{l-1}, \lambda_{l+1}, \dots, \lambda_k$ mit

$$\vec{u}_l = \sum_{i=1}^{l-1} \lambda_i \vec{u}_i + \sum_{i=l+1}^k \lambda_i \vec{u}_i,$$

so folgt daraus

$$\vec{o} = \sum_{i=1}^{l-1} \lambda_i \vec{u}_i + (-1) \vec{u}_l + \sum_{i=l+1}^k \lambda_i \vec{u}_i = \sum_{i=1}^k \mu_i \vec{u}_i$$

mit $\mu_i = \lambda_i$ für alle $i \neq l$ und $\mu_l = -1$, also $\mu_l \neq 0$. Der Nullvektor ist somit auf nicht triviale Weise als Linearkombination der Vektoren $\vec{u}_1, \dots, \vec{u}_k$ darstellbar; M ist daher linear abhängig. $\square$

Bemerkungen:

- Wenn sich ein Vektor $\vec{x}$ als Linearkombination von Vektoren $\vec{u}_1, \dots, \vec{u}_k$ darstellen lässt, sagt man, der Vektor $\vec{x}$ ist von den Vektoren $\vec{u}_1, \dots, \vec{u}_k$ linear abhängig. Ist dies nicht der Fall, so ist $\vec{x}$ von den Vektoren $\vec{u}_1, \dots, \vec{u}_k$ linear unabhängig. Satz 5.9 lässt sich damit auch folgendermaßen formulieren: Eine Menge M von Vektoren ist genau dann linear abhängig, wenn ein Vektor in M existiert, der von den restlichen Vektoren von M linear abhängig ist.
- Die Formulierung in dem Satz 5.9 „wenn ein Vektor in M existiert, der sich als Linearkombination der restlichen Vektoren von M darstellen lässt“ beinhaltet nicht, dass *nur* ein Vektor mit dieser Eigenschaft existieren muss.

Meist ist sogar die Mehrzahl der Vektoren einer linear abhängigen Vektormenge von den jeweils verbleibenden Vektoren linear abhängig. Genauer ausgedrückt lassen sich alle Vektoren $\vec{u}_i$, deren Koeffizienten λ_i bei der Darstellung $\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{o}$ von Null verschieden sind, als Linearkombinationen der jeweils restlichen Vektoren darstellen. Dies wird in dem ersten Teil des Beweises von Satz 5.9 deutlich: $\vec{u}_l$ kann jeder Vektor sein, für den $\lambda_l \neq 0$ ist.

Beispiel 5.21

Für die drei Vektoren $\vec{x}$, $\vec{u}$ und $\vec{w}$ in dem Beispiel 5.19 auf S. 182 wurde festgestellt, dass diese linear abhängig sind, da z. B. $-2\vec{x} + 3\vec{u} + 4\vec{w} = \vec{o}$ ist. Jeder dieser drei Vektoren lässt sich als Linearkombination der jeweils zwei anderen Vektoren darstellen: $\vec{x} = \frac{3}{2}\vec{u} + 2\vec{w}$, $\vec{u} = \frac{2}{3}\vec{x} - \frac{4}{3}\vec{w}$, $\vec{w} = \frac{1}{2}\vec{x} - \frac{3}{4}\vec{u}$. ■

Satz 5.10

Es seien $M = \{\vec{u}_1; \dots; \vec{u}_k\}$ eine linear unabhängige Menge von Vektoren eines Vektorraumes V und $\vec{v}$ ein beliebiger Vektor aus V . Dann folgt aus

$$\vec{v} = \sum_{i=1}^k \lambda_i \vec{u}_i \quad \text{und} \quad \vec{v} = \sum_{i=1}^k \mu_i \vec{u}_i$$

stets $\lambda_i = \mu_i$ für alle i mit $1 \leq i \leq k$.

Bemerkung: Der Satz 5.10 sagt aus, dass die Koeffizienten bei der Darstellung eines Vektors als Linearkombination einer Menge linear unabhängiger Vektoren eindeutig bestimmt sind.

Beweis: Aus $\vec{v} = \sum_{i=1}^k \lambda_i \vec{u}_i$ und $\vec{v} = \sum_{i=1}^k \mu_i \vec{u}_i$ folgt

$$\sum_{i=1}^k \lambda_i \vec{u}_i = \sum_{i=1}^k \mu_i \vec{u}_i \quad \text{bzw.} \quad \sum_{i=1}^k (\lambda_i - \mu_i) \vec{u}_i = \vec{o}.$$

Wegen der linearen Unabhängigkeit von $\{\vec{u}_1; \dots; \vec{u}_k\}$ folgt daraus $\lambda_i - \mu_i = 0$ für alle $i = 1 \dots k$ und somit die Behauptung. □

Das folgende Beispiel verdeutlicht die Notwendigkeit der Voraussetzung, dass eine Vektormenge M linear unabhängig ist, für die Eindeutigkeit der Darstellung eines Vektors als Linearkombination der Vektoren von M .

Beispiel 5.22

Die Vektoren $\vec{a} = \begin{pmatrix} 1 \\ 2 \\ 3 \\ 0 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 2 \\ 1 \\ 2 \\ -1 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} 3 \\ 1 \\ 2 \\ -1 \end{pmatrix}$, $\vec{d} = \begin{pmatrix} -4 \\ 1 \\ 0 \\ 3 \end{pmatrix}$ sind linear abhängig

(siehe Aufgabe 6 auf S. 186). Der Vektor $\vec{x} = \begin{pmatrix} 24 \\ 3 \\ 12 \\ -15 \end{pmatrix}$ lässt sich auf unterschied-

liche Arten als Linearkombination der Vektoren $\vec{a}$, $\vec{b}$, $\vec{c}$ und $\vec{d}$ darstellen, denn löst man das der Vektorgleichung $\lambda_1 \vec{a} + \lambda_2 \vec{b} + \lambda_3 \vec{c} + \lambda_4 \vec{d} = \vec{x}$ entsprechende LGS, so erhält man $\lambda_1 = -2t - 6$, $\lambda_2 = 3t + 15$, $\lambda_3 = 0$, $\lambda_4 = t$ (mit $t \in \mathbb{R}$). Somit ist z. B. $\vec{x} = -8\vec{a} + 18\vec{b} + 0\vec{c} + \vec{d}$ und $\vec{x} = -4\vec{a} + 12\vec{b} + 0\vec{c} - \vec{d}$. ■

5.3.4 Aufgaben zu Abschnitt 5.3

- Überprüfen Sie, ob die folgenden 2×2 -Matrizen als Linearkombinationen der Matrizen $\begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}$ und $\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$ darstellbar sind.
 - $A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$
 - $B = \begin{pmatrix} 5 & 3 \\ 2 & 5 \end{pmatrix}$
- Ermitteln Sie die lineare Hülle $\langle p_1, p_2, p_3 \rangle$ der Polynome p_1, p_2 und p_3 mit $p_1(x) = x + 1$, $p_2(x) = x^2 + x$ und $p_3(x) = x^2$.
- Zeigen Sie, dass die Vektormenge $E = \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$ ein Erzeugendensystem von $\mathbb{R}^3$ ist.
- Beweisen Sie den Satz 5.6 a auf S. 180.
- Weisen Sie nach: Sind U und V lineare Unterräume sowie $E_U = \{\vec{u}_1; \dots; \vec{u}_k\}$ und $E_V = \{\vec{v}_1; \dots; \vec{v}_m\}$ Erzeugendensysteme von U bzw. V , so ist die Summe der Unterräume U und V (siehe Definition 5.3) identisch mit der linearen Hülle der Vektormenge $E_U \cup E_V$, d. h. $U + V = \langle \vec{u}_1, \dots, \vec{u}_k, \vec{v}_1, \dots, \vec{v}_m \rangle$.
- Zeigen Sie, dass die Vektoren $\vec{a} = \begin{pmatrix} 1 \\ 2 \\ 3 \\ 0 \end{pmatrix}$, $\vec{b} = \begin{pmatrix} 2 \\ 1 \\ 2 \\ -1 \end{pmatrix}$, $\vec{c} = \begin{pmatrix} 3 \\ 1 \\ 2 \\ -1 \end{pmatrix}$, $\vec{d} = \begin{pmatrix} -4 \\ 1 \\ 0 \\ 3 \end{pmatrix}$ linear abhängig sind und überprüfen Sie, welche(r) der Vektoren sich als Linearkombination der jeweils drei anderen Vektoren darstellen lässt/lassen.
- Beweisen Sie Teil b des Satzes 5.8 auf S. 183.
- Folgt aus der linearen Unabhängigkeit von zwei Vektoren $\vec{u}$ und $\vec{v}$ auch die lineare Unabhängigkeit von $\vec{u} - \vec{v}$ und $\vec{u} + \vec{v}$?
- Folgt aus der linearen Unabhängigkeit dreier Vektoren $\vec{u}, \vec{v}, \vec{w}$ auch die lineare Unabhängigkeit der drei Vektoren $\vec{u} + \vec{v} + \vec{w}, \vec{u} + \vec{v}, \vec{v} + \vec{w}$?
- Gegeben sind die folgenden Unterräume von $\mathbb{R}^2$:
 $U_1 = \left\langle \begin{pmatrix} 1 \\ 2 \end{pmatrix} \right\rangle$, $U_2 = \left\langle \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -2 \end{pmatrix} \right\rangle$, $U_3 = \left\langle \begin{pmatrix} 1 \\ -3 \end{pmatrix} \right\rangle$.
 Welche der folgenden Aussagen ist (sind) richtig?
 - $\left\{ \begin{pmatrix} -2 \\ -4 \end{pmatrix} \right\}$ ist ein Erzeugendensystem von $U_1 \cap U_2$.
 - $\left\{ \begin{pmatrix} 1 \\ 4 \end{pmatrix} \right\}$ ist eine linear unabhängige Teilmenge von U_2 .
 - Es gilt $\langle U_1 \cup U_3 \rangle = \mathbb{R}^2$.
- Beweisen Sie den folgenden Satz: Ist eine Vektormenge $\{\vec{u}_1; \dots; \vec{u}_k\}$ linear unabhängig und ist ein Vektor $\vec{v}$ von $\{\vec{u}_1; \dots; \vec{u}_k\}$ linear unabhängig, so ist auch die Menge $\{\vec{u}_1; \dots; \vec{u}_k; \vec{v}\}$ linear unabhängig.
- Weisen Sie nach, dass eine Menge linear unabhängiger Vektoren des Vektorraumes $\mathbb{R}^n$ aus höchstens n Vektoren bestehen kann. Ziehen Sie dazu die in dem Abschnitt 1.3 gemachten Aussagen über Lösbarkeit und Lösungsmengen von linearen Gleichungssystemen heran.

5.4 Basis und Dimension

5.4.1 Der Begriff der Basis

In dem vorangegangenen Abschnitt wurde festgestellt, dass sich mit den Vektoren eines Erzeugendensystems alle Vektoren eines Vektorraumes erzeugen (d. h. als Linearkombinationen darstellen) lassen, es hierbei aber oft verschiedene Möglichkeiten gibt, die Darstellung also nicht eindeutig ist. Hingegen ist die Darstellung eines Vektors als Linearkombination linear unabhängiger Vektoren eindeutig. Jedoch ist es meistens nicht möglich, alle Vektoren eines Vektorraumes aus einer Menge linear unabhängiger Vektoren zu erzeugen. Um also die *eindeutige Erzeugung aller Vektoren eines Vektorraumes* zu ermöglichen, benötigt man linear unabhängige Erzeugendensysteme.

Definition 5.7

Eine Menge $B = \{\vec{b}_1; \vec{b}_2; \dots; \vec{b}_n\}$ von Vektoren eines Vektorraumes V heißt *Basis* von V , falls folgende Bedingungen erfüllt sind:

- B ist ein Erzeugendensystem von V ,
- B ist linear unabhängig.



Beispiel 5.23

- Die bekannteste Basis von $\mathbb{R}^3$ ist die *Standardbasis* bzw. *kanonische Basis* $B_0 = \{\vec{e}_1; \vec{e}_2; \vec{e}_3\}$ mit $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ und $\vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$, welche die Richtungen der Achsen eines kartesischen Koordinatensystems beschreiben. Dass es sich hierbei um ein Erzeugendensystem handelt, ist recht offensichtlich (siehe S. 179) und auch die lineare Unabhängigkeit liegt auf der Hand: nur mit $\lambda_1 = \lambda_2 = \lambda_3 = 0$ kann $\lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$ sein.
- Neben der Standardbasis gibt es unendlich viele weitere Basen von $\mathbb{R}^3$. Eine Variation der Standardbasis besteht darin, die Vektoren $\vec{e}_1, \vec{e}_2$ und $\vec{e}_3$ mit beliebigen (von Null verschiedenen) reellen Zahlen zu multiplizieren. Die dadurch entstehende Vektormenge $\left\{ \begin{pmatrix} r_1 \\ 0 \\ 0 \end{pmatrix}; \begin{pmatrix} 0 \\ r_2 \\ 0 \end{pmatrix}; \begin{pmatrix} 0 \\ 0 \\ r_3 \end{pmatrix} \right\}$ mit $r_1 \neq 0, r_2 \neq 0, r_3 \neq 0$ ist (wie sich leicht zeigen lässt) ebenfalls eine Basis.
- Als weiteres Beispiel einer Basis von $\mathbb{R}^3$ sei $B = \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$ angeführt (siehe Abb. 5.2). In der Aufgabe 3 (S. 186) wurde gezeigt, dass B ein

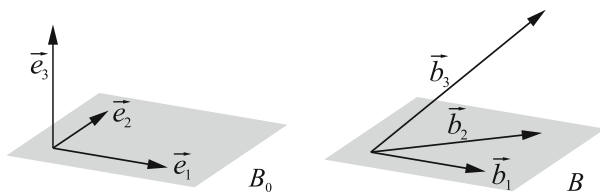


Abb. 5.2: Zwei Basen des Vektorraumes $\mathbb{R}^3$

Erzeugendensystem ist. Der Nachweis der linearen Unabhängigkeit ist ebenfalls sehr einfach; die Vektorgleichung $\lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$ ist nur mit $\lambda_1 = \lambda_2 = \lambda_3 = 0$ erfüllt. (In diesem Falle ist es nicht nötig, ein LGS aufzustellen und umzuformen. Durch Betrachtung der dritten Komponente folgt nämlich unmittelbar $\lambda_3 = 0$; betrachtet man danach die zweite Komponente, so folgt $\lambda_2 = 0$ und schließlich $\lambda_1 = 0$). ■

5.4.2 Koordinaten von Vektoren bezüglich Basen

Mithilfe des Begriffs der Basis lässt sich der bislang recht „intuitiv“ und eingeschränkt verwendete Koordinatenbegriff präzisieren und erweitern. Dies erfolgt hier zunächst für Koordinaten von Vektoren. Die angestellten Überlegungen bilden jedoch auch die Grundlage dafür, bei der Behandlung affiner Punkträume in Abschnitt 5.5 den Begriff „Koordinatensystem“ exakt zu definieren und Punkten Koordinaten bezüglich verschiedener Koordinatensysteme zuzuordnen.

Definition 5.8

Ist $B = \{\vec{b}_1; \vec{b}_2; \dots; \vec{b}_n\}$ eine Basis eines Vektorraumes V und $\vec{u}$ ein beliebiger Vektor aus V mit

$$\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i,$$

so heißen $\lambda_1, \lambda_2, \dots, \lambda_n \in \mathbb{R}$ *Koordinaten des Vektors $\vec{u}$ bezüglich der Basis B* . ♦

Mit dem Satz 5.10 auf S. 185 wurde bereits nachgewiesen, dass die Darstellung eines Vektors als Linearkombination einer linear unabhängigen Vektormenge *eindeutig* ist. Da eine Basis $B = \{\vec{b}_1; \vec{b}_2; \dots; \vec{b}_n\}$ eines Vektorraumes V zugleich ein Erzeugendensystem ist, *existiert* zudem für jeden Vektor $\vec{u} \in V$ eine Darstellung der Form $\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i$. Somit gilt also der folgende Satz.

Satz 5.11

Ist B eine Basis eines Vektorraumes V , so sind jedem Vektor $\vec{u} \in V$ eindeutig Koordinaten bezüglich B zugeordnet.

Beispiel 5.24

Es werden die Koordinaten des Vektors $\vec{u} = \begin{pmatrix} 4 \\ -5 \\ 6 \end{pmatrix}$ bezüglich der Standardbasis

$B_0 = \{\vec{e}_1; \vec{e}_2; \vec{e}_3\}$ sowie der Basis $B = \left\{ \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \right\}$ (siehe Beispiel 5.23)

bestimmt. Bezüglich der Standardbasis sind die Koordinaten bereits unmittelbar durch die Koeffizienten des Vektors $\vec{u}$ gegeben: $\vec{u} = 4\vec{e}_1 - 5\vec{e}_2 + 6\vec{e}_3$; $\vec{u}$ hat also bezüglich B_0 die Koordinaten $(4; -5; 6)$.

Um die Koordinaten von $\vec{u}$ bezüglich B zu bestimmen, ist die Vektorgleichung

$$\lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 4 \\ -5 \\ 6 \end{pmatrix}$$

zu lösen. Man erhält als Lösung $\lambda_1 = 9$, $\lambda_2 = -11$, $\lambda_3 = 6$; also hat $\vec{u}$ bezüglich B die Koordinaten $(9; -11; 6)$. ■

5.4.3 Weitere Beispiele für Basen

Beispiel 5.25

Basen des Vektorraumes der 2×2 -Matrizen (siehe Beispiel 5.6 auf S. 171) sind z. B. die Standardbasis $B_0 = \left\{ \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}; \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}; \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}; \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} \right\}$ sowie die Menge $B = \left\{ \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}; \begin{pmatrix} 1 & 2 \\ 0 & 0 \end{pmatrix}; \begin{pmatrix} 1 & 2 \\ 3 & 0 \end{pmatrix}; \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \right\}$.

Die Überprüfung, ob eine Menge von Matrizen eine Basis des Vektorraumes der 2×2 -Matrizen ist, kann analog zu Beispiel 5.23 erfolgen, da sich 2×2 -Matrizen als Quadrupel reeller Zahlen auffassen lassen. Um die lineare Unabhängigkeit der Menge B zu bestätigen, muss gezeigt werden, dass die Vektorgleichung

$$\lambda_1 \cdot \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} + \lambda_2 \cdot \begin{pmatrix} 1 & 2 \\ 0 & 0 \end{pmatrix} + \lambda_3 \cdot \begin{pmatrix} 1 & 2 \\ 3 & 0 \end{pmatrix} + \lambda_4 \cdot \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

nur die triviale Lösung besitzt. Durch Lösen des entsprechenden LGS

$$\begin{aligned} \lambda_1 + \lambda_2 + \lambda_3 + \lambda_4 &= 0 \\ 2\lambda_2 + 2\lambda_3 + 2\lambda_4 &= 0 \\ 3\lambda_3 + 3\lambda_4 &= 0 \\ 4\lambda_4 &= 0 \end{aligned}$$

ergibt sich tatsächlich $\lambda_1 = \lambda_2 = \lambda_3 = \lambda_4 = 0$; B ist somit linear unabhängig.

B ist ein Erzeugendensystem des Vektorraumes der 2×2 -Matrizen, wenn sich jede 2×2 -Matrix $\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ als Linearkombination der Matrizen von B darstellen lässt, d. h. wenn das Gleichungssystem

$$\begin{aligned} \lambda_1 + \lambda_2 + \lambda_3 + \lambda_4 &= a_{11} \\ 2\lambda_2 + 2\lambda_3 + 2\lambda_4 &= a_{12} \\ 3\lambda_3 + 3\lambda_4 &= a_{21} \\ 4\lambda_4 &= a_{22} \end{aligned}$$

für beliebige $a_{11}, a_{12}, a_{21}, a_{22} \in \mathbb{R}$ lösbar ist. Dies ist der Fall, durch Lösen des LGS erhält man $\lambda_1 = a_{11} - \frac{a_{12}}{2}$, $\lambda_2 = \frac{a_{12}}{2} - \frac{a_{21}}{3}$, $\lambda_3 = \frac{a_{21}}{3} - \frac{a_{22}}{4}$, $\lambda_4 = \frac{a_{22}}{4}$.

Somit ist B ein Erzeugendensystem und wegen der bereits gezeigten linearen Unabhängigkeit eine Basis des Vektorraumes der 2×2 -Matrizen. ■

Beispiel 5.26

Wir betrachten erneut den Vektorraum der Polynome höchstens n -ten Grades, siehe Beispiel 5.5 auf S. 170. In dem Beispiel 5.18 auf S. 181 wurde das Erzeugendensystem $E = \{p_0; p_1; p_2; \dots; p_n\}$ mit $p_0(x) = 1$, $p_1(x) = x$, $p_2(x) = x^2$, $\dots$, $p_n(x) = x^n$ für diesen Vektorraum angegeben. Dieses Erzeugendensystem ist eine Basis, denn das Nullpolynom $o(x) = 0 = 0x^n + \dots + 0x^2 + 0x + 0$ lässt sich nur auf triviale Weise als Linearkombination der Menge $\{p_0; p_1; \dots; p_n\}$ darstellen; aus $\lambda_n x^n + \dots + \lambda_2 x^2 + \lambda_1 x + \lambda_0 = o(x)$ folgt $\lambda_n = \dots = \lambda_2 = \lambda_1 = \lambda_0 = 0$.

Es fällt die Verwandtschaft der hier betrachteten Basis E mit der Standardbasis von $\mathbb{R}^{n+1}$ in das Auge. Allgemein lässt sich aus jeder Basis des Vektorraumes $\mathbb{R}^{n+1}$ eine Basis des Vektorraumes der Polynome höchstens n -ten Grades ableiten, indem die Komponenten der Basisvektoren von $\mathbb{R}^{n+1}$ als Koeffizienten der Basispolynome verwendet werden. ■

Wir haben bei den bisherigen Beispielen, bei denen zu entscheiden war, ob eine Menge M von Vektoren eines Vektorraumes V eine Basis bildet, entsprechend der Definition 5.7 jeweils überprüft, ob M linear unabhängig ist *und* ob sie ein Erzeugendensystem von V bildet. Mithilfe des folgenden Satzes kann in vielen Fällen auf die Prüfung einer der beiden Eigenschaften verzichtet werden; liegt eine Menge von n Vektoren des Vektorraumes $\mathbb{R}^n$ vor, so genügt es, zu prüfen, ob diese linear unabhängig ist *oder* ob sie ein Erzeugendensystem von $\mathbb{R}^n$ bildet.

Satz 5.12

Eine Menge $\{\vec{u}^{(1)}; \vec{u}^{(2)}; \dots; \vec{u}^{(n)}\}$ von n Vektoren des Vektorraumes $\mathbb{R}^n$ ist genau dann linear unabhängig, wenn sie ein Erzeugendensystem von $\mathbb{R}^n$ bildet.

Beweis: Die Menge $\{\vec{u}^{(1)}; \dots; \vec{u}^{(n)}\}$ ist genau dann linear unabhängig, wenn aus $\sum_{i=1}^n \lambda_i \vec{u}^{(i)} = \vec{0}$ folgt $\lambda_1 = \lambda_2 = \dots = \lambda_n = 0$. Dies bedeutet, dass das aus n Gleichungen bestehende und n Variablen enthaltende lineare Gleichungssystem

$$\begin{array}{ccccccc} u_1^{(1)} \lambda_1 + u_1^{(2)} \lambda_2 + \dots + u_1^{(n)} \lambda_n & = & 0 \\ u_2^{(1)} \lambda_1 + u_2^{(2)} \lambda_2 + \dots + u_2^{(n)} \lambda_n & = & 0 \\ \vdots & & \vdots \\ u_n^{(1)} \lambda_1 + u_n^{(2)} \lambda_2 + \dots + u_n^{(n)} \lambda_n & = & 0 \end{array} \quad \text{bzw.} \quad \left(\begin{array}{cccc|c} u_1^{(1)} & u_1^{(2)} & \dots & u_1^{(n)} & 0 \\ u_2^{(1)} & u_2^{(2)} & \dots & u_2^{(n)} & 0 \\ \vdots & \vdots & & \vdots & \vdots \\ u_n^{(1)} & u_n^{(2)} & \dots & u_n^{(n)} & 0 \end{array} \right)$$

nur die triviale Lösung besitzt, also eindeutig lösbar ist. Dies ist genau dann der Fall, wenn der Rang der einfachen Koeffizientenmatrix des LGS n ist, siehe Abschnitt 1.3.3 (speziell den Kasten auf S. 35). Genau dann, wenn der Rang der einfachen Koeffizientenmatrix des obigen LGS n ist, stimmt der Rang der einfachen Koeffizientenmatrix des folgenden LGS

$$\begin{array}{ccccccc} u_1^{(1)} \lambda_1 + u_1^{(2)} \lambda_2 + \dots + u_1^{(n)} \lambda_n & = & x_1 \\ u_2^{(1)} \lambda_1 + u_2^{(2)} \lambda_2 + \dots + u_2^{(n)} \lambda_n & = & x_2 \\ \vdots & & \vdots \\ u_n^{(1)} \lambda_1 + u_n^{(2)} \lambda_2 + \dots + u_n^{(n)} \lambda_n & = & x_n \end{array} \quad \text{bzw.} \quad \left(\begin{array}{cccc|c} u_1^{(1)} & u_1^{(2)} & \dots & u_1^{(n)} & x_1 \\ u_2^{(1)} & u_2^{(2)} & \dots & u_2^{(n)} & x_2 \\ \vdots & \vdots & & \vdots & \vdots \\ u_n^{(1)} & u_n^{(2)} & \dots & u_n^{(n)} & x_n \end{array} \right)$$

für beliebige $x_1, x_2, \dots, x_n \in \mathbb{R}$ mit dem Rang seiner erweiterten Koeffizientenmatrix überein. Dies wiederum ist gleichbedeutend damit, dass dieses LGS für beliebige $x_1, x_2, \dots, x_n \in \mathbb{R}$ lösbar und somit jeder Vektor $\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n$ als

Linearkombination der Vektoren $\vec{u}^{(1)}, \dots, \vec{u}^{(n)}$ darstellbar ist. Genau in diesem Falle ist $\{\vec{u}^{(1)}; \dots; \vec{u}^{(n)}\}$ ein Erzeugendensystem. Da alle in diesem Beweis gezogenen Schlüsse umkehrbar sind, wurde der Satz 5.12 damit bewiesen. $\square$

Beispiele für Basen von Unterräumen

Beispiel 5.27

Es wird eine Basis des von der Menge $E = \left\{ \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}; \begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix}; \begin{pmatrix} 7 \\ 8 \\ 9 \end{pmatrix} \right\}$ erzeugten Unterraumes $U = \langle E \rangle$ von $\mathbb{R}^3$ bestimmt. Wäre E linear unabhängig, so wäre E selbst eine Basis von U . Dies trifft jedoch nicht zu, denn die Vektoren sind linear

abhängig; z. B. ist $\begin{pmatrix} 7 \\ 8 \\ 9 \end{pmatrix} = -\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + 2\begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix}$. Jeder Vektor, der als Linearkombination von E darstellbar ist, lässt sich damit auch als Linearkombination von $E_1 = \left\{ \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}; \begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix} \right\} \subset E$ darstellen. Somit ist E_1 ein Erzeugendensystem von U . Da E_1 zudem linear unabhängig ist, handelt es sich um eine Basis von U . ■

Beispiel 5.28

Lösungsmengen homogener linearer Gleichungssysteme in n Variablen sind Unterräume von $\mathbb{R}^n$, siehe Beispiel 5.9 auf S. 173. Basen dieser Unterräume erhält man durch Lösen der LGS.

- Das aus nur einer Gleichung bestehende LGS $-3x + \frac{3}{2}y - 5z = 0$ hat die Lösungsmenge

$$L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mid \begin{pmatrix} x \\ y \\ z \end{pmatrix} = s \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + t \begin{pmatrix} -5 \\ 0 \\ 3 \end{pmatrix}; s, t \in \mathbb{R} \right\}.$$

Die Vektormenge $B = \left\{ \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix}; \begin{pmatrix} -5 \\ 0 \\ 3 \end{pmatrix} \right\}$ erzeugt somit den Unterraum (von $\mathbb{R}^3$) der Lösungen der gegebenen Gleichung. Da B zudem offensichtlich linear unabhängig ist, handelt es sich um eine Basis dieses Unterraumes.

- Bestimmt man die Lösungsmenge des homogenen Gleichungssystems

$$\begin{array}{rcl} 2x_1 + 7x_2 + 9x_3 + x_4 & = & 0 \\ 4x_1 + 14x_2 + 8x_3 + 3x_4 & = & 0 \\ x_1 + 3x_2 + 5x_3 - 3x_4 & = & 0 \end{array} \quad \text{bzw.} \quad \left(\begin{array}{cccc|c} 2 & 7 & 9 & 1 & 0 \\ 4 & 14 & 8 & 3 & 0 \\ 1 & 3 & 5 & -3 & 0 \end{array} \right),$$

so erhält man

$$L = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \mid \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = t \begin{pmatrix} 232 \\ -69 \\ 1 \\ 10 \end{pmatrix}; t \in \mathbb{R} \right\}.$$

Eine Basis des Unterraumes (von $\mathbb{R}^4$) der Lösungen des gegebenen LGS ist also die einelementige Menge $B = \left(\begin{pmatrix} 232 \\ -69 \\ 1 \\ 10 \end{pmatrix} \right)$. ■

5.4.4 Sätze über Basen von Vektorräumen

Satz 5.13

Eine Teilmenge $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eines Vektorraumes V ist genau dann eine Basis von V , wenn folgende Bedingungen erfüllt sind:

- B ist ein Erzeugendensystem von V .
- Es existiert keine echte Teilmenge von B , die Erzeugendensystem von V ist.

Bemerkung: Satz 5.13 sagt aus, dass jede Basis ein „minimales“ Erzeugendensystem ist.

Beweis: Falls B eine Basis ist, so ist B definitionsgemäß linear unabhängig. Damit kann keine echte Teilmenge von B ein Erzeugendensystem von V sein, denn wäre z. B. $B \setminus \{\vec{b}_k\} = \{\vec{b}_1; \dots; \vec{b}_{k-1}; \vec{b}_{k+1}; \dots; \vec{b}_n\}$ (mit $k \in \mathbb{N}, 1 \leq k \leq n$) ein

Erzeugendensystem von V , so wäre $\vec{b}_k$ als Linearkombination dieser Teilmenge von B darstellbar und somit von $\{\vec{b}_1; \dots; \vec{b}_{k-1}; \vec{b}_{k+1}; \dots; \vec{b}_n\}$ linear abhängig. Dies widerspricht aber der linearen Unabhängigkeit von $B = \{\vec{b}_1; \dots; \vec{b}_n\}$.

Wir kommen zum Beweis der Rückrichtung des Satzes. Hierbei ist zu zeigen, dass $\vec{b}_1, \dots, \vec{b}_n$ linear unabhängig sind. Wir führen den Beweis indirekt und nehmen an, $\vec{b}_1, \dots, \vec{b}_n$ seien linear abhängig. In diesem Falle würde ein Vektor in der Menge $B = \{\vec{b}_1, \dots, \vec{b}_n\}$ existieren, der von den übrigen $n-1$ Vektoren linear abhängig ist. Da die Vektoren $\vec{b}_1, \dots, \vec{b}_n$ keine besonderen Eigenschaften besitzen, die sie gegenüber den jeweils anderen Vektoren von B auszeichnen würden, ist es möglich, sie anders zu nummerieren. O. B. d. A. (ohne Beschränkung der Allgemeinheit) kann angenommen werden, dass der Vektor $\vec{b}_n$ von den Vektoren $\vec{b}_1, \dots, \vec{b}_{n-1}$ linear abhängig ist. Dann existieren $\lambda_1, \dots, \lambda_{n-1} \in \mathbb{R}$ mit

$$\vec{b}_n = \sum_{i=1}^{n-1} \lambda_i \vec{b}_i.$$

Wir zeigen, dass dann $\{\vec{b}_1, \dots, \vec{b}_{n-1}\}$ ein Erzeugendensystem von V ist, was der Voraussetzung widerspricht, dass kein Erzeugendensystem von V existiert, welches eine echte Teilmenge von B ist. Dazu sei $\vec{x}$ ein beliebiger Vektor von V . Da $B = \{\vec{b}_1, \dots, \vec{b}_n\}$ ein Erzeugendensystem ist, existieren $\mu_1, \dots, \mu_n \in \mathbb{R}$ mit

$$\begin{aligned} \vec{x} &= \sum_{i=1}^n \mu_i \vec{b}_i = \sum_{i=1}^{n-1} \mu_i \vec{b}_i + \mu_n \vec{b}_n = \sum_{i=1}^{n-1} \mu_i \vec{b}_i + \mu_n \sum_{i=1}^{n-1} \lambda_i \vec{b}_i \\ &= \sum_{i=1}^{n-1} \mu_i \vec{b}_i + \sum_{i=1}^{n-1} \lambda_i \mu_n \vec{b}_i = \sum_{i=1}^{n-1} (\mu_i + \lambda_i \mu_n) \vec{b}_i \end{aligned}$$

Setzt man $\nu_i = \mu_i + \lambda_i \mu_n$ (für $i = 1, \dots, n-1$), so ist

$$\vec{x} = \sum_{i=1}^{n-1} \nu_i \vec{b}_i.$$

Somit lässt sich jeder beliebige Vektor $\vec{x} \in V$ als Linearkombination der Vektoren $\vec{b}_1, \dots, \vec{b}_{n-1}$ darstellen; $\{\vec{b}_1, \dots, \vec{b}_{n-1}\}$ ist also ein Erzeugendensystem von V . Dies widerspricht der Voraussetzung, dass keine echte Teilmenge von B ein Erzeugendensystem von V ist. Die Annahme, dass $\vec{b}_1, \dots, \vec{b}_n$ linear abhängig sind, führt somit zum Widerspruch, und die lineare Unabhängigkeit von $\vec{b}_1, \dots, \vec{b}_n$ ist nachgewiesen. $\square$

Satz 5.14

Verkürzungssatz: *Es sei V ein Vektorraum und $E_m = \{\vec{e}_1; \dots; \vec{e}_m\}$ ein Erzeugendensystem von V . Dann ist entweder E eine Basis von V oder es existiert eine echte Teilmenge $B \subset E$, die eine Basis von V ist.*

Beweis: Wir gehen davon aus, dass E nicht den Nullvektor enthält, anderenfalls entfernen wir ihn zunächst, wobei natürlich die Eigenschaft, ein Erzeugendensystem zu sein, nicht verloren geht. Ist E linear unabhängig, so ist E eine Basis von V und die Behauptung ist gezeigt. Ist E linear abhängig, so lässt sich mindestens einer der Vektoren $\vec{e}_1, \dots, \vec{e}_m$ als Linearkombination der restlichen Vektoren darstellen, durch Umnummerierung lässt sich erreichen, dass dies

der Vektor $\vec{e}_m$ ist; es existieren also $\lambda_1, \dots, \lambda_{m-1} \in \mathbb{R}$ mit $\vec{e}_m = \sum_{i=1}^{m-1} \lambda_i \vec{e}_i$. Damit ist auch die Menge $E_{m-1} = E_m \setminus \{\vec{e}_m\} = \{\vec{e}_1; \dots; \vec{e}_{m-1}\}$ ein Erzeugendensystem von V , denn da sich jeder Vektor $\vec{x}$ von V als Linearkombination $\vec{x} = \sum_{i=1}^m \mu_i \vec{e}_i$ (mit $\mu_1, \dots, \mu_m \in \mathbb{R}$) darstellen lässt, gilt

$$\vec{x} = \sum_{i=1}^m \mu_i \vec{e}_i = \sum_{i=1}^{m-1} \mu_i \vec{e}_i + \mu_m \sum_{i=1}^{m-1} \lambda_i \vec{e}_i = \sum_{i=1}^{m-1} \nu_i \vec{e}_i$$

mit $\nu_i = \mu_i + \lambda_i \mu_m$ (für $i = 1, \dots, m-1$). Jeder Vektor $\vec{x} \in V$ lässt sich also auch aus E_{m-1} erzeugen.

E_{m-1} ist somit ein Erzeugendensystem von V . Ist E_{m-1} linear unabhängig, so ist E_{m-1} eine Basis von V , und die Behauptung ist gezeigt. Ist E_{m-1} linear abhängig, so lässt sich mindestens einer der Vektoren $\vec{e}_1, \dots, \vec{e}_{m-1}$ als Linearkombination der restlichen Vektoren darstellen, durch Umnummerierung lässt sich erreichen, dass dies der Vektor $\vec{e}_{m-1}$ ist. Auf dieselbe Weise wie zuvor für E_{m-1} , zeigt man, dass auch $E_{m-2} = E_m \setminus \{\vec{e}_{m-1}, \vec{e}_m\} = \{\vec{e}_1; \dots; \vec{e}_{m-2}\}$ ein Erzeugendensystem von V ist. Dieses Verfahren lässt sich weiter fortsetzen, wobei Erzeugendensysteme mit immer weniger Vektoren entstehen. Man erhält schließlich ein linear unabhängiges Erzeugendensystem, also eine Basis von V . Dies muss spätestens nach dem $m-1$ -ten Schritt der Fall sein, da dann nur noch ein einziger Vektor verbleibt, der (da E nicht den Nullvektor enthält) linear unabhängig sein muss. $\square$

Satz 5.15

Eine Teilmenge $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eines Vektorraumes V ist genau dann eine Basis von V , wenn folgende Bedingungen erfüllt sind:

- B ist eine linear unabhängige Teilmenge von V .
- Jede Menge $B' \subseteq V$, die B als echte Teilmenge enthält (für die also $B \subset B'$ gilt) ist linear abhängig.

Bemerkung: Der Satz 5.15 sagt aus, dass jede Basis ein „maximales“ System linear unabhängiger Vektoren ist.

Beweis: Ist B eine Basis, so ist B auch ein Erzeugendensystem von V ; jeder Vektor $\vec{x} \in V$ lässt sich somit als Linearkombination von B darstellen. Fügt man zu B also einen beliebigen Vektor $\vec{x} \in V$ hinzu, so entsteht eine linear abhängige Menge; somit ist jede Menge B' mit $B \subset B'$ linear abhängig.

Umgekehrt muss gezeigt werden, dass jede (im Sinne der Formulierung des Satzes) maximale Menge linear unabhängiger Vektoren eines Vektorraumes V ein Erzeugendensystem von V ist, sich also jeder Vektor $\vec{x} \in V$ als Linearkombination der Vektoren $\vec{b}_1; \dots; \vec{b}_n$ darstellen lässt. Wäre dies für einen Vektor $\vec{x}$ nicht der Fall, gäbe es also keine $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ mit $\vec{x} = \lambda_1 \vec{b}_1 + \dots + \lambda_n \vec{b}_n$, so wäre die Menge $B' = \{\vec{b}_1; \dots; \vec{b}_n; \vec{x}\}$ linear unabhängig. Wegen $B \subset B'$ steht dies im Widerspruch zur Voraussetzung. Somit lässt sich jeder Vektor $\vec{x} \in V$ als Linearkombination der Vektoren $\vec{b}_1; \dots; \vec{b}_n$ darstellen, B ist somit ein Erzeugendensystem von V . $\square$

Satz 5.16

Basisergänzungssatz: *Es seien V ein Vektorraum, $U = \{\vec{u}_1; \dots; \vec{u}_k\}$ eine linear unabhängige Teilmenge und $E = \{\vec{e}_1; \dots; \vec{e}_l\}$ ein Erzeugendensystem von V . Dann lässt sich U durch Elemente von E zu einer Basis von V ergänzen.*

Satz 5.17

Es seien $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis und $U = \{\vec{u}_1; \dots; \vec{u}_k\}$ eine linear unabhängige Teilmenge eines Vektorraumes V . Dann gilt $k \leq n$.

Beweis: Da B eine Basis ist, existieren $\lambda_{ij} \in \mathbb{R}$ ($1 \leq i \leq k, 1 \leq j \leq n$) mit

$$\begin{aligned}
 \vec{u}_1 &= \sum_{j=1}^n \lambda_{1j} \vec{b}_j = \lambda_{11} \vec{b}_1 + \dots + \lambda_{1n} \vec{b}_n \\
 (*) \quad \vec{u}_2 &= \sum_{j=1}^n \lambda_{2j} \vec{b}_j = \lambda_{21} \vec{b}_1 + \dots + \lambda_{2n} \vec{b}_n \\
 &\vdots \quad \quad \quad = \quad \quad \quad \vdots \\
 \vec{u}_k &= \sum_{j=1}^n \lambda_{kj} \vec{b}_j = \lambda_{k1} \vec{b}_1 + \dots + \lambda_{kn} \vec{b}_n.
 \end{aligned}$$

Nach Voraussetzung ist $U = \{\vec{u}_1; \dots; \vec{u}_k\}$ linear unabhängig, die Vektorgleichung

$$(**) \quad \mu_1 \vec{u}_1 + \mu_2 \vec{u}_2 + \dots + \mu_k \vec{u}_k = \vec{o}.$$

darf also nur die triviale Lösung $(\mu_1, \dots, \mu_k) = (0, \dots, 0)$ besitzen. Wir formen $(**)$ um und zeigen, dass daraus $k \leq n$ folgt. Einsetzen von $(*)$ in $(**)$ ergibt

$$\mu_1 \sum_{j=1}^n \lambda_{1j} \vec{b}_j + \dots + \mu_k \sum_{j=1}^n \lambda_{kj} \vec{b}_j = \vec{o}$$

bzw. in ausführlicher Schreibweise und sortiert nach den Vektoren $\vec{b}_1, \dots, \vec{b}_n$:

$$\begin{aligned}
 &(\mu_1 \lambda_{11} + \mu_2 \lambda_{21} + \dots + \mu_k \lambda_{k1}) \cdot \vec{b}_1 \\
 &+ (\mu_1 \lambda_{12} + \mu_2 \lambda_{22} + \dots + \mu_k \lambda_{k2}) \cdot \vec{b}_2 \\
 &\vdots \quad \quad \quad \vdots \\
 &+ (\mu_1 \lambda_{1n} + \mu_2 \lambda_{2n} + \dots + \mu_k \lambda_{kn}) \cdot \vec{b}_n = \vec{o}.
 \end{aligned}$$

Da $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis und damit linear unabhängig ist, ist diese Vektorgleichung genau dann erfüllt, wenn jeder der Koeffizienten von $\vec{b}_1, \dots, \vec{b}_n$ Null ist. Diese Bedingung ist gleichbedeutend mit dem linearen Gleichungssystem

$$\begin{aligned}
 \mu_1 \lambda_{11} + \mu_2 \lambda_{21} + \dots + \mu_k \lambda_{k1} &= 0 \\
 \mu_1 \lambda_{12} + \mu_2 \lambda_{22} + \dots + \mu_k \lambda_{k2} &= 0 \\
 \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\
 \mu_1 \lambda_{1n} + \mu_2 \lambda_{2n} + \dots + \mu_k \lambda_{kn} &= 0.
 \end{aligned}$$

Dabei handelt es sich um ein homogenes LGS mit n Gleichungen und k Variablen $\mu_1, \dots, \mu_k$ (die Koeffizienten λ_{ij} sind durch $(*)$ bestimmt). Nur wenn dieses LGS *eindeutig* lösbar ist, d. h. $(\mu_1, \dots, \mu_k) = (0, \dots, 0)$, ist wegen $(**)$ die Voraussetzung der linearen Unabhängigkeit von $U = \{\vec{u}_1; \dots; \vec{u}_k\}$ erfüllt. Somit muss das LGS eine Lösungsmenge ohne Parameter besitzen. Da ein lineares Gleichungssystem in k Variablen mit dem Rang r der einfachen und der erweiterten Koeffizientenmatrix eine $k-r$ -parametrische Lösungsmenge besitzt (siehe Abschnitt 1.3.3, speziell den Kasten auf S. 35), muss $k-r = 0$ sein. Da zudem der Rang eines LGS stets kleiner oder gleich der Anzahl der Gleichungen ist, also $r \leq n$, gilt $k-n \leq 0$ und somit die Behauptung $k \leq n$. $\square$

Satz 5.18

Alle Basen eines Vektorraumes V bestehen aus gleich vielen Vektoren, d. h.: sind $B_1 = \{\vec{b}_1; \dots; \vec{b}_n\}$ und $B_2 = \{\vec{c}_1; \dots; \vec{c}_m\}$ Basen von V , so ist $n = m$.

Beweis: Da jede Basis von V eine linear unabhängige Teilmenge von V ist, lässt sich der Satz 5.17 zweifach anwenden. Es ergibt sich $m \leq n$ sowie $n \leq m$ und somit $n = m$. $\square$

Folgerung: Ist $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis von V und $U = \{\vec{u}_1; \dots; \vec{u}_n\}$ eine linear unabhängige Teilmenge von V , die ebenso viele Elemente enthält wie B , so ist U ebenfalls eine Basis von V .

Diese Folgerung kann z. B. mithilfe der Sätze 5.16 und 5.18 begründet werden, siehe die Aufgabe 8 auf S. 201.

Die Sätze 5.13-5.18 stellen wesentliche Zusammenhänge zwischen Erzeugendensystemen, Mengen linear unabhängiger Vektoren und Basen endlich erzeugbarer Vektorräume her, die im Folgenden zusammengefasst werden.

- Jede Basis ist ein „minimales“ Erzeugendensystem eines Vektorraumes.
- Jedes Erzeugendensystem eines Vektorraumes lässt sich zu einer Basis reduzieren.
- Jede Basis ist ein „maximales“ System linear unabhängiger Vektoren.
- Jedes System linear unabhängiger Vektoren lässt sich zu einer Basis erweitern.
- Jede Basis eines endlich erzeugbaren Vektorraumes besteht aus gleich vielen Vektoren wie jede andere Basis dieses Vektorraumes.

Der Steinitzsche Austauschsatz

Häufig ist es z. B. bei Beweisführungen nötig, Vektoren einer Basis gegen andere Vektoren auszutauschen. Wir zeigen mit dem folgenden Satz zunächst, dass sich einzelne Vektoren einer Basis austauschen lassen und in Verallgemeinerung dessen dann, dass der Austausch beliebiger Teilmengen von Basen gegen Mengen linear unabhängiger Vektoren möglich ist (Steinitzscher Austauschsatz).

Satz 5.19

Es seien V ein Vektorraum, $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis von V und $\vec{b} \in V$ ein Vektor mit $\vec{b} = \sum_{j=1}^n \mu_j \vec{b}_j$. Falls $\mu_k \neq 0$ ist (für ein $k \in \mathbb{N}$, $1 \leq k \leq n$), so ist auch die Menge $B' = \{\vec{b}_1; \dots; \vec{b}_{k-1}; \vec{b}; \vec{b}_{k+1}; \dots; \vec{b}_n\}$ eine Basis von V .

Beweis: Da $\vec{b}_1, \dots, \vec{b}_n$ unnummeriert werden können, lässt sich o. B. d. A. $k=1$ setzen. Es ist damit zu zeigen, dass $B' = \{\vec{b}; \vec{b}_2; \dots; \vec{b}_n\}$ eine Basis von V ist.

- B' ist ein Erzeugendensystem von V

Es ist $\vec{b} = \mu_1 \vec{b}_1 + \mu_2 \vec{b}_2 + \dots + \mu_n \vec{b}_n$ mit $\mu_1 \neq 0$. Somit existiert $\frac{1}{\mu_1}$ und es gilt

$$\vec{b}_1 = \frac{1}{\mu_1} \vec{b} - \frac{\mu_2}{\mu_1} \vec{b}_2 - \dots - \frac{\mu_n}{\mu_1} \vec{b}_n.$$

Somit kann jeder Vektor, der als Linearkombination von B darstellbar ist, auch als Linearkombination von B' dargestellt werden. Also lassen sich alle Vektoren von V als Linearkombination von B' darstellen, und B' ist somit ein Erzeugendensystem von V .

■ *B' ist linear unabhängig*

Zu zeigen ist, dass aus $\lambda \vec{b}_1 + \lambda_2 \vec{b}_2 + \dots + \lambda_n \vec{b}_n = \vec{0}$ folgt $\lambda = \lambda_2 = \dots = \lambda_n = 0$. Dazu ersetzt man $\vec{b}$ durch $\sum_{j=1}^n \mu_j \vec{b}_j$ und erhält

$$\begin{aligned} \lambda (\mu_1 \vec{b}_1 + \mu_2 \vec{b}_2 + \dots + \mu_n \vec{b}_n) + \lambda_2 \vec{b}_2 + \dots + \lambda_n \vec{b}_n &= \vec{0} \quad \text{bzw.} \\ (\lambda \mu_1) \vec{b}_1 + (\lambda_2 + \lambda \mu_2) \vec{b}_2 + \dots + (\lambda_n + \lambda \mu_n) \vec{b}_n &= \vec{0}. \end{aligned}$$

Da die Vektoren $\vec{b}_1, \dots, \vec{b}_n$ linear unabhängig sind, folgt daraus

$$\lambda \mu_1 = 0, \quad \lambda_2 + \lambda \mu_2 = 0, \quad \dots, \quad \lambda_n + \lambda \mu_n = 0.$$

Wegen $\mu_1 \neq 0$, folgt $\lambda = 0$ und somit auch $\lambda_2 = \dots = \lambda_n = 0$.

Die Menge B' ist somit linear unabhängig. $\square$

Satz 5.20

Austauschsatz von Steinitz: Es seien V ein Vektorraum, $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis und $U = \{\vec{u}_1; \dots; \vec{u}_k\}$ eine linear unabhängige Teilmenge von V . Dann lassen sich k Vektoren von B gegen die Vektoren $\vec{u}_1, \dots, \vec{u}_k$ so austauschen, dass die entstehende Menge auch eine Basis ist, d. h. bei geeigneter Nummerierung der Vektoren $\vec{b}_1, \dots, \vec{b}_n$ ist $B' = \{\vec{u}_1; \dots; \vec{u}_k, \vec{b}_{k+1}; \dots; \vec{b}_n\}$ eine Basis von V .

Beweis: Wegen des Satzes 5.17 ist $k \leq n$. Im Falle $k = n$ müssen zu U keine Vektoren von B hinzugefügt werden, denn U ist nach der Folgerung aus dem Satz 5.18 auf S. 195 bereits eine Basis, also $B' = U$. Für alle $k < n$ beweisen wir den Satz mittels vollständiger Induktion.

Induktionsanfang: Für $k = 1$ gilt Satz 5.20 unmittelbar wegen des Satzes 5.19.

Induktionsvoraussetzung: Ist $U_{k-1} = \{\vec{u}_1; \dots; \vec{u}_{k-1}\} \subset V$ linear unabhängig, so ist bei geeigneter Nummerierung der Vektoren $\vec{b}_1, \dots, \vec{b}_n$ die Menge $B' = \{\vec{u}_1; \dots; \vec{u}_{k-1}, \vec{b}_k; \dots; \vec{b}_n\}$ eine Basis von V .

Induktionsbehauptung: Ist $U_k = \{\vec{u}_1; \dots; \vec{u}_k\}$ linear unabhängig, so ist bei geeigneter Nummerierung $B'' = \{\vec{u}_1; \dots; \vec{u}_k, \vec{b}_{k+1}; \dots; \vec{b}_n\}$ eine Basis von V .

Induktionsbeweis: Ist $U_k = \{\vec{u}_1; \dots; \vec{u}_k\}$ linear unabhängig, so ist wegen der Induktionsvoraussetzung $B' = \{\vec{u}_1; \dots; \vec{u}_{k-1}, \vec{b}_k; \dots; \vec{b}_n\}$ eine Basis von V .

Es muss also nur noch einer der Vektoren $\vec{b}_k, \dots, \vec{b}_n$ durch $\vec{u}_k$ ersetzt werden. Um den Satz 5.19 dafür anzuwenden, ist zu zeigen, dass in der Darstellung $\vec{u}_k = \sum_{i=1}^{k-1} \lambda_i \vec{u}_i + \sum_{i=k}^n \lambda_i \vec{b}_i$ durch die Basisvektoren von B' einer der Koeffizienten $\lambda_k, \dots, \lambda_n$ von Null verschieden ist. Dies ist der Fall, da $\vec{u}_k$ als Element einer linear unabhängigen Vektormenge nicht der Nullvektor sein kann und zudem wegen der linearen Unabhängigkeit von U_k der Vektor $\vec{u}_k$ keine Linearkombination von $\vec{u}_1, \dots, \vec{u}_{k-1}$ ist. Nach Satz 5.19 kann somit bei geeigneter Nummerierung $\vec{b}_k$ gegen $\vec{u}_k$ ausgetauscht werden, und $B'' = \{\vec{u}_1; \dots; \vec{u}_k, \vec{b}_{k+1}; \dots; \vec{b}_n\}$ ist eine Basis von V . $\square$

5.4.5 Der Begriff der Dimension

Mithilfe der in dem vorherigen Abschnitt herausgearbeiteten Eigenschaften von Basen können wir nun einen Begriff exakt fassen, der in der Umgangssprache von Bedeutung ist und der auch in diesem Buch bereits auftrat, dabei aber stets in einem intuitiven, nicht genau bestimmbareren Sinne verwendet wurde.

Mit dem Satz 5.18 wurde gezeigt, dass alle Basen eines Vektorraumes aus gleich vielen Vektoren bestehen. Die Anzahl der Vektoren einer Basis ist somit unabhängig von der Wahl dieser Basis; sie ist ein Merkmal, das einen Vektorraum kennzeichnet – wie wir noch sehen werden, das wichtigste Merkmal.

Definition 5.9

Ist V ein Vektorraum mit einer endlichen Basis $B = \{\vec{b}_1; \dots; \vec{b}_n\}$, so heißt die Anzahl der Basisvektoren *Dimension* von V :

$$\dim V = n.$$

Hat ein Vektorraum keine endliche Basis, so heißt er *unendlichdimensional*. Der Vektorraum $V = \{\vec{0}\}$, der nur den Nullvektor enthält, heißt *nulldimensional*. ♦

Beispiele

Die Dimensionen einiger der bislang betrachteten Vektorräume lassen sich aufgrund der bereits angestellten Überlegungen zu Basen sofort angeben:

- Die Dimension des Vektorraumes $\mathbb{R}^n$ ist n , siehe Satz 5.12 auf S. 190.
- Wegen der Verwandtschaft mit $\mathbb{R}^{n+1}$ ist die Dimension des Vektorraumes der Polynome höchstens n -ten Grades $n+1$, siehe Beispiel 5.26 auf S. 189.
- Der Vektorraum aller *Pfeilklassen der Ebene* hat die Dimension 2, der Vektorraum aller *Pfeilklassen des Raumes* die Dimension 3. Dies folgt nicht nur aus anschaulichen Überlegungen, sondern aus der in dem Abschnitt 3.3.2 festgestellten Strukturgleichheit (Isomorphie) dieser Vektorräume mit $\mathbb{R}^2$ bzw. $\mathbb{R}^3$. Wegen dieser Isomorphie lassen sich Basen der Vektorräume der Pfeilklassen auf Basen von $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ abbilden und umgekehrt.
- Der Unterraum der Lösungen eines homogenen linearen Gleichungssystems in n Variablen vom Rang r hat $n-r$ Basisvektoren, seine Dimension ist daher $n-r$, siehe Beispiel 5.28 auf S. 191 sowie Abschnitt 1.3.
- Der Vektorraum der 2×2 -Matrizen hat die Dimension 4 (siehe Beispiel 5.25 auf S. 189); allgemein hat der Vektorraum der $m \times n$ -Matrizen (Beispiel 5.6 auf S. 171) die Dimension $m \cdot n$.

Nach diesen recht allgemeinen Beispielen betrachten wir nun einen Vektorraum, dessen Dimension nicht unmittelbar ersichtlich ist.

Beispiel 5.29

In dem Beispiel 5.12 auf S. 174 wurde der Vektorraum aller magischen Quadrate der Kantenlänge 3 betrachtet. Dabei handelt es sich um einen Unterraum des Vektorraumes der 3×3 -Matrizen und gleichzeitig um die Lösungsmenge des auf

S. 175 angegebenen linearen Gleichungssystems. Um eine Basis und damit die Dimension des Vektorraumes der magischen 3×3 -Quadrate zu bestimmen, muss dieses aus acht Gleichungen bestehende LGS mit 10 Variablen (9 Matrixelemente sowie die Zeilen-, Spalten- und Diagonalsumme s) gelöst werden. Man erhält eine Lösungsmenge folgender Gestalt:

$$L = \left\{ \left(\begin{pmatrix} a_{11} \\ a_{12} \\ a_{13} \\ a_{21} \\ a_{22} \\ a_{23} \\ a_{31} \\ a_{32} \\ a_{33} \end{pmatrix} \middle| \begin{pmatrix} a_{11} \\ a_{12} \\ a_{13} \\ a_{21} \\ a_{22} \\ a_{23} \\ a_{31} \\ a_{32} \\ a_{33} \end{pmatrix} = \lambda_1 \begin{pmatrix} -1 \\ 1 \\ 0 \\ 1 \\ 0 \\ -1 \\ 0 \\ -1 \\ 1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ -1 \\ 1 \\ 1 \\ 0 \\ -1 \\ -1 \\ 1 \\ 0 \end{pmatrix} + s \begin{pmatrix} \frac{2}{3} \\ \frac{1}{3} \\ 0 \\ -\frac{1}{3} \\ \frac{1}{3} \\ 1 \\ \frac{2}{3} \\ \frac{1}{3} \\ 0 \end{pmatrix} ; \lambda_1, \lambda_2, s \in \mathbb{R} \right\}.$$

Mithilfe dieser Darstellung lassen sich leicht magische Quadrate mit frei wählbarer Zeilen-, Spalten- und Diagonalsumme s konstruieren. Kehrt man zur Matrizenscheibweise zurück, so erhält man die folgende Basis des Vektorraumes der magischen Quadrate der Kantenlänge 3:

$$B = \left\{ \begin{pmatrix} -1 & 1 & 0 \\ 1 & 0 & -1 \\ 0 & -1 & 1 \end{pmatrix}; \begin{pmatrix} 0 & -1 & 1 \\ 1 & 0 & -1 \\ -1 & 1 & 0 \end{pmatrix}; \begin{pmatrix} \frac{2}{3} & \frac{1}{3} & 0 \\ -\frac{1}{3} & \frac{1}{3} & 1 \\ \frac{2}{3} & \frac{1}{3} & 0 \end{pmatrix} \right\}.$$

Die Dimension dieses Vektorraumes ist also 3.

Eine alternative Möglichkeit, die Dimension des Vektorraumes der magischen Quadrate der Kantenlänge 3 zu bestimmen (ohne eine Basis zu ermitteln) besteht in der Berechnung des Rangs der Koeffizientenmatrix des auf S. 175 angegebenen LGS (darauf wird im nächsten Kapitel noch näher eingegangen). Man erhält den Rang 7; somit ist die Dimension des Lösungsraumes wegen der Anzahl von $n=10$ Variablen $n-r=10-7=3$ (siehe Abschnitt 1.3). ■

Bislang wurden nur Vektorräume endlicher Dimension betrachtet. Jedoch gibt es auch Vektorräume, die keine endlichen Basen besitzen.

Beispiel 5.30

Der Vektorraum aller *Folgen* (a_n) *reeller Zahlen* (siehe Beispiel 5.3 auf S. 169) ist unendlichdimensional. Wäre seine Dimension nämlich endlich, so müsste eine Zahl $k \in \mathbb{N}$ existieren, so dass der Vektorraum eine Basis $B = \{\vec{b}_1; \dots; \vec{b}_k\}$ besitzt. Nach Satz 5.17 kann es dann keine Mengen linear unabhängiger Vektoren geben, die mehr als k Elemente enthalten. Jedoch gelingt es – wie groß k auch immer gewählt sein möge – mühelos, $k+1$ linear unabhängige Vektoren des Vektorraumes der reellen Zahlenfolgen anzugeben, nämlich die Zahlenfolgen:

$$(a_n)^{(1)} = (1; 0; 0; \dots; 0; 0; 0; \dots), \dots, (a_n)^{(2)} = (0; 1; 0; \dots; 0; 0; 0; \dots), \dots, \\ (a_n)^{(k)} = (0; 0; 0; \dots; 1; 0; 0; \dots), \quad (a_n)^{(k+1)} = (0; 0; 0; \dots; 0; 1; 0; \dots).$$

$\underset{1}{}, \dots, \underset{2}{}, \dots, \underset{k}{}, \underset{k+1}{}$

Somit gibt es also keine Begrenzung der Anzahl linear unabhängiger Vektoren und daher keine endliche Basis des Vektorraumes aller reellen Zahlenfolgen. Der Vektorraum ist daher unendlichdimensional. ■

Zu dem Beispiel 5.30 sei angemerkt, dass der Vektorraum der reellen Zahlenfolgen vielfältige Unterräume besitzt, die teilweise von endlicher Dimension sind, siehe die Aufgabe 10 auf S. 201.

5.4.6 Die Dimensionsformel für lineare Unterräume

Der folgende Satz gehört zu den zentralen Strukturaussagen der linearen Algebra. Es ist sinnvoll, diesen Satz zunächst anhand eines Beispiels selbst zu erarbeiten. Der Leserin bzw. dem Leser sei daher empfohlen, vor dem Weiterlesen zunächst die Aufgabe 13 auf S. 202 zu bearbeiten.

Satz 5.21

Es seien V ein Vektorraum und U_1, U_2 zwei Unterräume von V . Dann gilt die Dimensionsformel

$$\dim U_1 + \dim U_2 = \dim(U_1 \cap U_2) + \dim(U_1 + U_2)$$

Beweis: Es sei k die Dimension und $B_\cap = \{\vec{b}_1; \dots; \vec{b}_k\}$ eine Basis der Schnittmenge $U_1 \cap U_2$, von der nach Satz 5.2 (S. 176) bekannt ist, dass sie ein linearer Unterraum von V ist. Die Grundidee des Beweises besteht darin, in Basen von U_1 und U_2 jeweils k Basisvektoren durch $\vec{b}_1, \dots, \vec{b}_k$ zu ersetzen, alle Vektoren der so entstehenden Basen zu betrachten und nachzuweisen, dass diese eine Basis der Summe $U_1 + U_2$ der Unterräume U_1 und U_2 (siehe Abschnitt 5.2.2) bilden.

Es seien m und l die Dimensionen der Unterräume U_1 bzw. U_2 , wobei wegen $U_1 \cap U_2 \subseteq U_1$ und $U_1 \cap U_2 \subseteq U_2$ gilt $k \leq m$ und $k \leq l$. In zwei Basen

$$B_1 = \{\vec{c}_1; \dots; \vec{c}_m\}, \quad B_2 = \{\vec{d}_1; \dots; \vec{d}_l\}$$

der Unterräume U_1 bzw. U_2 lassen sich nach dem Steinitzschen Austauschsatz (Satz 5.20) jeweils k Vektoren gegen $\vec{b}_1, \dots, \vec{b}_k$ ersetzen. Da $\vec{b}_1, \dots, \vec{b}_k$ zu $U_1 \cap U_2$ gehören, sind sie sowohl in U_1 als auch in U_2 enthalten, und da diese Vektoren eine Basis von $U_1 \cap U_2$ bilden, sind sie linear unabhängig – die Voraussetzungen des Steinitzschen Austauschsatzes sind also erfüllt. Bei geeigneter Nummerierung entstehen folgende Basen B'_1 und B'_2 von U_1 bzw. U_2 :

$$B'_1 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{c}_{k+1}; \dots; \vec{c}_m\}, \quad B'_2 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{d}_{k+1}; \dots; \vec{d}_l\}.$$

Wir vereinigen nun alle Vektoren dieser beiden Basen zu einer Menge

$$B_+ = \{\vec{b}_1; \dots; \vec{b}_k; \vec{c}_{k+1}; \dots; \vec{c}_m; \vec{d}_{k+1}; \dots; \vec{d}_l\}.$$

und weisen nach, dass B_+ eine Basis von $U_1 + U_2$ ist.

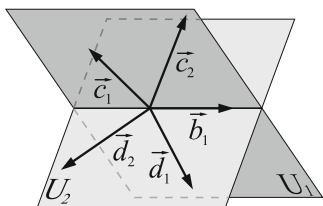


Abb. 5.3: Dimensionsformel für lineare Unterräume

Die zweidimensionalen Unterräume U_1 und U_2 werden durch Ebenen dargestellt. Es ist in diesem Beispiel $\dim(U_1 \cap U_2) = 1$, $B_\cap = \{\vec{b}_1\}$, $B_1 = \{\vec{c}_1; \vec{c}_2\}$, $B'_1 = \{\vec{b}_1; \vec{c}_2\}$ sowie $B_2 = \{\vec{d}_1; \vec{d}_2\}$, $B'_2 = \{\vec{b}_1; \vec{d}_2\}$.

- B_+ ist ein Erzeugendensystem von $U_1 + U_2$.

Jeder Vektor $\vec{u} \in U_1 + U_2$ ist die Summe zweier Vektoren $\vec{u}_1 \in U_1$ und $\vec{u}_2 \in U_2$ (vgl. Definition 5.3). Da B'_1 eine Basis von U_1 und B'_2 eine Basis von U_2 ist, existieren Koeffizienten $\lambda_1, \dots, \lambda_m$ und $\mu_1, \dots, \mu_l$ mit

$$\vec{u}_1 = \sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^m \lambda_i \vec{c}_i, \quad \vec{u}_2 = \sum_{i=1}^k \mu_i \vec{b}_i + \sum_{i=k+1}^l \mu_i \vec{d}_i.$$

Somit ist

$$\begin{aligned} \vec{u} &= \vec{u}_1 + \vec{u}_2 = \sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^m \lambda_i \vec{c}_i + \sum_{i=1}^k \mu_i \vec{b}_i + \sum_{i=k+1}^l \mu_i \vec{d}_i \\ &= \sum_{i=1}^k (\lambda_i + \mu_i) \vec{b}_i + \sum_{i=k+1}^m \lambda_i \vec{c}_i + \sum_{i=k+1}^l \mu_i \vec{d}_i. \end{aligned}$$

Jeder Vektor $\vec{u} \in U_1 + U_2$ ist somit als Linearkombination der Vektoren von $B_+ = \{\vec{b}_1; \dots; \vec{b}_k; \vec{c}_{k+1}; \dots; \vec{c}_m; \vec{d}_{k+1}; \dots; \vec{d}_l\}$ darstellbar; B_+ ist daher ein Erzeugendensystem von $U_1 + U_2$.

- B_+ ist linear unabhängig.

Wir setzen

$$\sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^m \mu_i \vec{c}_i + \sum_{i=k+1}^l \nu_i \vec{d}_i = \vec{o}$$

bzw., äquivalent dazu,

$$\sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^m \mu_i \vec{c}_i = - \sum_{i=k+1}^l \nu_i \vec{d}_i =: \vec{x}$$

und zeigen, dass dann $\lambda_1 = \dots = \lambda_k = \mu_{k+1} = \dots = \mu_m = \nu_{k+1} = \dots = \nu_l = 0$ sein muss. Der oben definierte Vektor $\vec{x}$ ist eine Linearkombination von $B'_1 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{c}_{k+1}; \dots; \vec{c}_m\}$ und zugleich eine Linearkombination von $B'_2 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{d}_{k+1}; \dots; \vec{d}_l\}$ (wobei die zu $\vec{b}_1, \dots, \vec{b}_k$ gehörenden Koeffizienten Null sind). Somit gilt $\vec{x} \in U_1$ und $\vec{x} \in U_2$, also $\vec{x} \in U_1 \cap U_2$. Damit muss es eine Darstellung $\vec{x} = \sum_{i=1}^k \lambda_i \vec{b}_i$ geben, woraus folgt:

$$- \sum_{i=k+1}^l \nu_i \vec{d}_i = \sum_{i=1}^k \lambda_i \vec{b}_i \quad \text{bzw.} \quad \sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^l \nu_i \vec{d}_i = \vec{o}.$$

Wegen der linearen Unabhängigkeit von $B'_2 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{d}_{k+1}; \dots; \vec{d}_l\}$ gilt daher $\lambda_1 = \dots = \lambda_k = \nu_{k+1} = \dots = \nu_l = 0$. Also ist $\vec{x} = \vec{o}$ und somit

$$\sum_{i=1}^k \lambda_i \vec{b}_i + \sum_{i=k+1}^m \mu_i \vec{c}_i = \vec{o}.$$

Wegen der linearen Unabhängigkeit von $B'_1 = \{\vec{b}_1; \dots; \vec{b}_k; \vec{c}_{k+1}; \dots; \vec{c}_m\}$ folgt daraus $\mu_{k+1} = \dots = \mu_m = 0$. Der Nullvektor lässt sich also nur auf triviale Weise als Linearkombination von $\vec{b}_1, \dots, \vec{b}_k, \vec{c}_{k+1}, \dots, \vec{c}_m, \vec{d}_{k+1}, \dots, \vec{d}_l$ darstellen; die lineare Unabhängigkeit von B_+ ist damit gezeigt.

B_+ ist somit eine Basis. Aufgrund der Anzahl ihrer Elemente gilt $\dim(U_1 + U_2) = k + (m - k) + (l - k)$. Wegen $\dim(U_1 \cap U_2) = k$, $\dim U_1 = m$ und $\dim U_2 = l$ folgt

$$\begin{aligned} \dim U_1 + \dim U_2 &= m + l = k + k + (m - k) + (l - k) \\ &= \dim(U_1 \cap U_2) + \dim(U_1 + U_2). \end{aligned}$$

□

5.4.7 Aufgaben zu Abschnitt 5.4

1. Zeigen Sie, dass $B = \left\{ \begin{pmatrix} 1 \\ 3 \\ -2 \end{pmatrix}; \begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix}; \begin{pmatrix} 3 \\ -2 \\ 1 \end{pmatrix} \right\}$ eine Basis von $\mathbb{R}^3$ ist.
2. Zeigen Sie, dass $B = \left\{ \begin{pmatrix} 1 \\ 4 \\ 3 \\ 0 \end{pmatrix}; \begin{pmatrix} 2 \\ 0 \\ 1 \\ 1 \end{pmatrix}; \begin{pmatrix} 1 \\ 0 \\ 0 \\ -1 \end{pmatrix}; \begin{pmatrix} 0 \\ 2 \\ 3 \\ 1 \end{pmatrix} \right\}$ eine Basis von $\mathbb{R}^4$ ist.
3. Bestimmen Sie die Koordinaten des Vektors $\vec{x}$ bezüglich der Basis B .

$$\vec{x} = \begin{pmatrix} 19 \\ 5 \\ -17 \end{pmatrix}; \quad B = \left\{ \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}; \begin{pmatrix} -4 \\ -5 \\ -6 \end{pmatrix}; \begin{pmatrix} 7 \\ 8 \\ 7 \end{pmatrix} \right\}$$
4. Bestimmen Sie eine Basis des von der Menge

$$X = \left\{ \begin{pmatrix} 0 \\ 1 \\ 0 \\ -1 \end{pmatrix}; \begin{pmatrix} 1 \\ 0 \\ 1 \\ -2 \end{pmatrix}; \begin{pmatrix} -1 \\ -2 \\ 0 \\ 1 \end{pmatrix}; \begin{pmatrix} -1 \\ 0 \\ 1 \\ 0 \end{pmatrix}; \begin{pmatrix} 1 \\ 0 \\ -1 \\ -1 \end{pmatrix}; \begin{pmatrix} 2 \\ 0 \\ -1 \\ 0 \end{pmatrix} \right\}$$
erzeugten Unterraumes $U = \langle X \rangle$ von $\mathbb{R}^4$.
5. Geben Sie Basen der folgenden Unterräume von $\mathbb{R}^3$ bzw. von $\mathbb{R}^4$ an:
 - a) $U_1 = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathbb{R}^3 \mid x = z \right\}$
 - b) $U_2 = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \in \mathbb{R}^4 \mid x_1 + 3x_2 + 2x_4 = 0; 2x_1 + x_2 + x_3 = 0 \right\}$
6. Zeigen Sie, dass 3 Vektoren der Form $\begin{pmatrix} 1 \\ x \\ x^2 \end{pmatrix}, \begin{pmatrix} 1 \\ y \\ y^2 \end{pmatrix}, \begin{pmatrix} 1 \\ z \\ z^2 \end{pmatrix}$ mit $x, y, z \in \mathbb{R}$,
 $x \neq y, x \neq z, y \neq z$ stets eine Basis von $\mathbb{R}^3$ bilden.
7. Beweisen Sie den Basisergänzungssatz (Satz 5.16 auf S. 194).
8. Begründen Sie die Folgerung aus dem Satz 5.18 auf S. 195.
9. Begründen Sie: Ist V ein n -dimensionaler Vektorraum und E ein Erzeugendensystem von V , das aus n Vektoren besteht, so ist E eine Basis von V .
10. Leonardo von Pisa (auch Fibonacci genannt) beschrieb das Wachstum einer Kaninchenpopulation durch eine später nach ihm benannte Zahlenfolge.
 - Zu Beginn gibt es ein Paar neugeborener Kaninchen.
 - Jedes neugeborene Kaninchenpaar wirft nach 2 Monaten ein neues Paar.
 - Anschließend wirft jedes Kaninchenpaar jeden Monat ein weiteres Paar.
 - Kaninchen leben ewig und haben einen unbegrenzten Lebensraum.

Jeden Monat kommt also zu der Anzahl der Paare, die im letzten Monat gelebt haben, die Anzahl derjenigen Paare hinzu, die bereits im vorletzten Monat gelebt haben, da diese geschlechtsreif sind und sich vermehren. Dem entspricht die Rekursionsformel $f_{n+2} = f_n + f_{n+1}$ mit $f_1 = 1, f_2 = 1$. Im Folgenden wird auf diese Anfangsbedingungen verzichtet, und es wird die Menge $\mathcal{F}$ der verallgemeinerten Fibonacci-Folgen betrachtet, d. h. die Menge aller Folgen (f_n) reeller Zahlen, für die gilt: $f_{n+2} = f_n + f_{n+1}$ (für alle $n \in \mathbb{N}, n \geq 1$).

- a) Weisen Sie nach, dass $\mathcal{F}$ ein Unterraum des Vektorraumes aller Folgen reeller Zahlen ist.
- b) Bestimmen Sie eine Basis von $\mathcal{F}$; geben Sie die Dimension von $\mathcal{F}$ an.
11. a) Bestimmen Sie die Koordinaten des magischen Quadrats $\begin{pmatrix} 4 & 9 & 2 \\ 3 & 5 & 7 \\ 8 & 1 & 6 \end{pmatrix}$ bezüglich der in dem Beispiel 5.29 auf S. 197f. ermittelten Basis.
- b) Geben Sie eine Basis des Vektorraumes der magischen Quadrate der Kantenlänge 3 an, in der das in a) aufgeführte magische Quadrat auftritt.
12. *Rechenmauern* sind ein beliebtes Übungsformat in der Grundschule. Es wird eine Grundreihe von k „Steinen“ (z. B. $k = 3, 4, 5$ oder 6) vorgegeben, und darüber werden dann Steine versetzt angeordnet, wobei jede Reihe einen Stein weniger enthält als die darunter liegende Reihe. Bei additiven Rechenmauern enthält jeder Stein diejenige Zahl, welche die Summe der Zahlen der beiden darunter liegenden Steine ist, siehe Abb. 5.4.
- a) Aus wie vielen Steinen besteht eine Rechenmauer insgesamt, wenn ihre Grundreihe k Steine enthält?
- b) Weisen Sie nach: Wenn in den Steinen der Grundreihen von Rechenmauern beliebige reelle Zahlen stehen können, so ist die Menge aller Rechenmauern mit vier Grundsteinen ein Unterraum von $\mathbb{R}^{10}$.
- c) Geben Sie eine Basis und die Dimension dieses Unterraumes an.
13. Gegeben sind die folgenden Unterräume von $\mathbb{R}^4$:

$$U_1 = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \mid x_1 = x_2 ; x_3 = x_4 \right\}$$

$$U_2 = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \mid 2x_1 + x_2 + x_3 = 0 ; 6x_1 + 3x_2 + 2x_3 + x_4 = 0 \right\}$$

Bestimmen Sie die Dimensionen der vier Unterräume U_1 , U_2 , $U_1 \cap U_2$ und $U_1 + U_2$. Welchen Zusammenhang erkennen Sie zwischen den Dimensionen?

Hinweise: – U_1 und U_2 sind Lösungsmengen homogener linearer Gleichungssysteme. Wie lässt sich dann $U_1 \cap U_2$ ausdrücken?

- Um $\dim(U_1 + U_2)$ zu bestimmen, ermitteln Sie Basen von U_1 und U_2 ; die Vereinigungsmenge dieser Basen ist ein Erzeugendensystem von $U_1 + U_2$ (warum?). Überprüfen Sie, ob dieses Erzeugendensystem linear unabhängig ist und – falls dies nicht zutrifft – reduzieren Sie es zu einer Basis von $U_1 + U_2$.

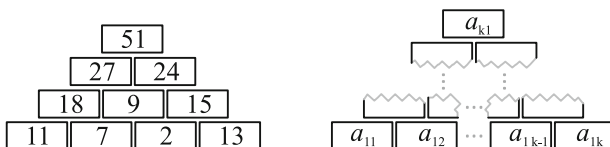


Abb. 5.4:
Rechenmauern

5.5 Affine Punkträume

Mithilfe affiner Räume lässt sich die in den vorangegangenen Abschnitten entwickelte Theorie der Vektorräume für die Geometrie nutzen. Dadurch ist es möglich, die bereits in den Kapiteln 2 und 4 behandelte analytische Geometrie zu verallgemeinern und Geometrie in beliebigen Dimensionen zu betreiben. Außerdem werden wir dadurch den Begriff des Koordinatensystems exakt definieren und eine Vielzahl von Koordinatensystemen verwenden können.

5.5.1 Der Begriff des affinen Raumes

Definition 5.10

Es sei V ein Vektorraum über dem Körper der reellen Zahlen. Eine nicht leere Menge A heißt *affiner Punktraum* bzw. kurz *affiner Raum* über dem Vektorraum V , falls eine Verknüpfung $+: A \times V \rightarrow A$ mit folgenden Eigenschaften existiert:

- (i) Für beliebige $P \in A$ gilt $P + \vec{o} = P$.
- (ii) Für beliebige $P, Q \in A$ existiert genau ein Vektor $\vec{v} \in V$ mit $Q = P + \vec{v}$.
- (iii) Für beliebige $P \in A$ und $\vec{u}, \vec{v} \in V$ gilt $P + (\vec{u} + \vec{v}) = (P + \vec{u}) + \vec{v}$.

Die Elemente von A werden als *Punkte* bezeichnet. ◆

Bemerkungen:

- Die Forderung nach der Existenz einer Verknüpfung $+: A \times V \rightarrow A$ beinhaltet, dass es zu jedem Punkt $P \in A$ und zu jedem Vektor $\vec{v} \in V$ einen Punkt $Q \in A$ mit $Q = P + \vec{v}$ gibt.
- Das Zeichen „+“ tritt nunmehr in drei verschiedenen Bedeutungen auf: Es wird damit die Summe zweier Zahlen, die Summe zweier Vektoren und jetzt auch die additive Verknüpfung eines Punktes mit einem Vektor bezeichnet. In welchem Sinne „+“ jeweils zu verstehen ist, wird aus den Operanden sichtbar, zwischen denen es steht. So steht in $P + (\vec{u} + \vec{v}) = (P + \vec{u}) + \vec{v}$ das „+“ in $(\vec{u} + \vec{v})$ für die Vektoraddition, die anderen drei „+“-Zeichen stehen für die Verknüpfung eines Punktes mit einem Vektor.
- Der Begriff „Punkt“ ist in der Elementargeometrie ein nicht definierbarer Grundbegriff. In der analytischen Geometrie auf Grundlage des Vektorraumbegriffs lässt sich dieser Begriff nun wie in der Definition 5.10 exakt definieren.
- Streng genommen dürfte nicht die Menge A , sondern es müsste das Tripel $(A, V, +)$ als affiner Raum bezeichnet werden, da der Vektorraum V und die Verknüpfung „+“ die Struktur eines affinen Raumes bestimmen. Wir nutzen die kürzere Bezeichnung und setzen den Bezug zu V und „+“ voraus.
- Statt $Q = P + \vec{v}$ schreiben wir künftig auch $\vec{v} = \overrightarrow{PQ}$ und verwenden damit eine bereits bekannte Schreibweise. Die Eigenschaft (ii) in der Definition 5.10 sagt aus, dass zu zwei beliebigen Punkten $P, Q \in A$ ein eindeutig bestimmter Vektor $\vec{v} \in V$ mit $\vec{v} = \overrightarrow{PQ}$ existiert. Mit dieser Schreibweise lässt sich an Stelle von (i) in Definition 5.10 schreiben: Für beliebige $P \in A$ gilt $\overrightarrow{PP} = \vec{o}$.

Beispiele für affine Punkträume

- Die Menge aller *Punkte des dreidimensionalen Anschauungsraumes* bildet mit den dort eingeführten Pfeilklassen als Vektoren (siehe Abschnitt 3.1) einen affinen Punktraum. Durch die Theorie der Vektorräume und Definition 5.10 wurden die in dem Abschnitt 3.1 angestellten Betrachtungen und die darauf basierende vektorielle Raumgeometrie nun strukturell fundiert.
- Die Menge $\mathbb{R}^n$ aller *n-Tupel reeller Zahlen* ist ein affiner Punktraum. Dies mag zunächst merkwürdig erscheinen, da $\mathbb{R}^n$ hierbei zugleich als Vektorraum auftritt, jedoch können sowohl Punkte als auch Vektoren als *n-Tupel* interpretiert werden. In dem Kapitel 4 wurden Punkte vielfach durch Ortsvektoren beschrieben, die als Zahlentripel angegeben wurden. Diese lassen sich auch direkt als Punkte, d. h. als Elemente des affinen Punktraumes $\mathbb{R}^n$ ansehen. Jedoch ist bei der Betrachtung des affinen Punktraumes $\mathbb{R}^n$ mit dem Vektorraum $\mathbb{R}^n$ stets zu beachten, welche *n-Tupel* für Punkte und welche *n-Tupel* für Vektoren stehen. Die in dem Kapitel 4 praktizierte Verwendung von Ortsvektoren an Stelle von Punkten ist daher weiterhin sinnvoll, um Fehlinterpretationen zu vermeiden.
- Lösungen (lösbarer) inhomogener linearer Gleichungssysteme lassen sich als Punkte, Lösungen der zugehörigen homogenen LGS als Vektoren interpretieren. Die Grundlage hierfür bildet der durch den Satz 1.3 (S. 38) ausgedrückte Zusammenhang zwischen Lösungen inhomogener und zugehöriger homogener LGS. Von dem Beispiel 5.9 auf S. 173 ist bekannt, dass Lösungsmengen homogener LGS in *n* Variablen lineare Unterräume von $\mathbb{R}^n$ sind. In dem Abschnitt 5.5.4 wird dann gezeigt, dass Lösungsmengen inhomogener LGS in *n* Variablen affine Unterräume von $\mathbb{R}^n$ sind.

Einige Folgerungen aus der Definition 5.10

Satz 5.22

Es seien *A* ein affiner Punktraum und *V* der zugehörige Vektorraum. Für beliebige Punkte $P, Q, R, S \in A$ und Vektoren $\vec{u}, \vec{v} \in V$ gilt:

1. $P + \vec{u} = P + \vec{v} \Rightarrow \vec{u} = \vec{v}$,
2. $P + \vec{v} = Q + \vec{v} \Rightarrow P = Q$,
3. $P + \vec{v} = Q \Rightarrow P = Q + (-\vec{v})$,
4. $\overrightarrow{PR} = \overrightarrow{PQ} + \overrightarrow{QR}$ (siehe Abb. 5.5 a),
5. $\overrightarrow{P_1P_2} + \overrightarrow{P_2P_3} + \dots + \overrightarrow{P_{k-1}P_k} = \sum_{i=1}^{k-1} \overrightarrow{P_iP_{i+1}} = \overrightarrow{P_1P_k}$ (siehe Abb. 5.5 b),
6. $\overrightarrow{PR} + \overrightarrow{RP} = \vec{0}$ bzw. $\overrightarrow{PR} = -\overrightarrow{RP}$,
7. $\overrightarrow{PQ} = \overrightarrow{RS} \Rightarrow \overrightarrow{PR} = \overrightarrow{QS}$ (siehe Abb. 5.5 c).

Die Aussagen dieses Satzes sind bereits bekannt und wurden in den Kapiteln 3 und 4 verwendet. Jedoch standen sie dort für spezielle Vektorräume und zugehörige Punktmengen zur Verfügung, während hier nun ihre Gültigkeit in beliebigen affinen Punkträumen mit zugehörigen Vektorräumen konstatiert wird.

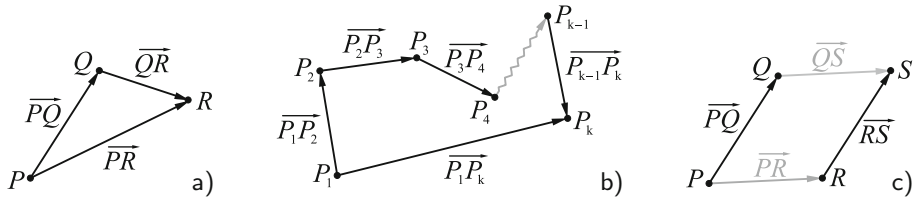


Abb. 5.5: Illustrationen zu Satz 5.22

Beweis: Wir beschränken uns darauf, einige Behauptungen des Satzes 5.22 zu beweisen, zu Nachweisen weiterer Teile dieses Satzes siehe Aufgabe 1 auf S. 218.

1. Nach Definition 5.10 (ii) existiert für beliebige $P, Q \in A$ genau ein Vektor $\vec{v} \in V$ mit $Q = P + \vec{v}$. Aus $P + \vec{u} = Q = P + \vec{v}$ folgt unmittelbar $\vec{u} = \vec{v}$.

4. Nach Definition 5.10 (ii) und der letzten Bemerkung dazu (siehe S. 203) existieren eindeutig bestimmte Vektoren $\overrightarrow{PR}$, $\overrightarrow{PQ}$ und $\overrightarrow{QR}$ mit

$$(1) \ P + \overrightarrow{PR} = R, \quad (2) \ P + \overrightarrow{PQ} = Q, \quad (3) \ Q + \overrightarrow{QR} = R.$$

Durch Einsetzen von (2) in (3) und Gleichsetzen mit (1) ergibt sich

$$(P + \overrightarrow{PQ}) + \overrightarrow{QR} = R = P + \overrightarrow{PR}.$$

Wegen der Definition 5.10 (iii) lässt sich dies zu $P + (\overrightarrow{PQ} + \overrightarrow{QR}) = P + \overrightarrow{PR}$ umformen, und wegen Satz 5.22, 1 folgt die Behauptung $\overrightarrow{PQ} + \overrightarrow{QR} = \overrightarrow{PR}$.

7. Aus $\overrightarrow{PQ} = \overrightarrow{RS}$ folgt wegen der Kommutativität der Vektoraddition zunächst

$$\overrightarrow{PQ} + \overrightarrow{QR} = \overrightarrow{RS} + \overrightarrow{QR} = \overrightarrow{QR} + \overrightarrow{RS}.$$

Wegen 4. ist $\overrightarrow{PQ} + \overrightarrow{QR} = \overrightarrow{PR}$ und $\overrightarrow{QR} + \overrightarrow{RS} = \overrightarrow{QS}$, woraus sich die Behauptung $\overrightarrow{PR} = \overrightarrow{QS}$ ergibt. $\square$

5.5.2 Koordinatensysteme

Definition 5.11

Es seien A ein affiner Raum und V der dazugehörige Vektorraum. Weiterhin seien O ein beliebiger Punkt von A und $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis von V . Dann heißt $K = \{O; \vec{b}_1; \dots; \vec{b}_n\}$ *Koordinatensystem* des affinen Raumes A . Der Punkt O wird als *Koordinatenursprung* des Koordinatensystems K bezeichnet.

Ist $P \in A$ ein Punkt mit $P = O + \sum_{i=1}^n \lambda_i \vec{b}_i$, so heißen $\lambda_1, \dots, \lambda_n$ die *Koordinaten* des Punktes P bezüglich des Koordinatensystems K . $\blacklozenge$

Satz 5.23

Jedem Punkt P eines affinen Raumes sind bezüglich eines beliebigen Koordinatensystems von A eindeutig Koordinaten zugeordnet.

Beweis: Es sei $K = \{O; \vec{b}_1; \dots; \vec{b}_n\}$ ein Koordinatensystem. Nach Definition 5.10 (ii) existiert ein eindeutig bestimmter Vektor $\vec{v}$ mit $P = O + \vec{v}$. Wegen Satz 5.11 sind $\vec{v}$ eindeutig Koordinaten $\lambda_1, \dots, \lambda_n$ bezüglich $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ zuzuordnen. Somit gibt es eine eindeutig bestimmte Darstellung $P = O + \vec{v} = O + \sum_{i=1}^n \lambda_i \vec{b}_i$ des Punktes P bezüglich des Koordinatensystems K . $\square$

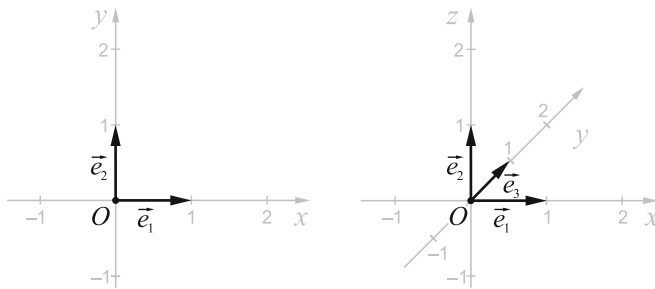


Abb. 5.6: Standardkoordinatensysteme von $\mathbb{R}^2$ und $\mathbb{R}^3$

Beispiel 5.31

Die bekannten kartesischen Koordinatensysteme der Ebene und des dreidimensionalen Anschauungsraumes sind Koordinatensysteme im Sinne der Definition 5.11 (siehe Abb. 5.6). Ihnen lassen sich (mittels der in dem Abschnitt 3.3 beschriebenen Zusammenhänge zwischen Pfeilklassen und n -Tupeln) Koordinatensysteme von $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ zuordnen, welche aus einer Standardbasis (vgl. Beispiel 5.23 auf S. 187) und dem Koordinatenursprung $O = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ bzw. $O = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$ bestehen. Analog dazu gibt es für beliebige $n \in \mathbb{N}$ „Standardkoordinatensysteme“ von $\mathbb{R}^n$. Die Koordinaten $(x_1; \dots; x_n)$ eines Punktes bezüglich eines Standardkoordinatensystems entsprechen seinen Komponenten als n -Tupel $\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$. ■

Komponenten- und Koordinatenschreibweise

Punkte des affinen Raumes $\mathbb{R}^n$ schreiben wir – ebenso wie Vektoren des Vektorraumes $\mathbb{R}^n$ – in Komponentenform in der Spaltenschreibweise, z. B. $P = \begin{pmatrix} 2 \\ -1 \\ 4 \end{pmatrix}$ (da es sich bei dem Punkt P um ein Element von $\mathbb{R}^3$ handelt, ist er das angegebene Tripel). Gleichzeitig *hat* er bezüglich des Standardkoordinatensystems die Koordinaten 2, -1 und 4. Dies schreiben wir in der Form $P(2; -1; 4)$. Wenn mehrere Koordinatensysteme auftreten, werden die Koordinaten so gekennzeichnet, dass Klarheit darüber besteht, bezüglich welchen Koordinatensystems die Koordinaten eines Punktes angegeben sind, z. B. $P(2; -1; 4)_K$.

Beispiel 5.32

Wir betrachten in $\mathbb{R}^2$ das Koordinatensystem $K = \{O_K; \vec{b}_1; \vec{b}_2\}$ mit dem Koordinatenursprung $O_K = \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix}$ und den Basisvektoren $b_1 = \begin{pmatrix} 1 \\ -\frac{1}{2} \end{pmatrix}$, $b_2 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$, siehe Abb. 5.7, und berechnen die Koordinaten der Punkte $P = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ und $Q = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$ bezüglich dieses Koordinatensystems.

Gesucht sind für P Koordinaten λ_1, λ_2 mit $P = O + \lambda_1 \vec{b}_1 + \lambda_2 \vec{b}_2$, also mit

$$\begin{pmatrix} 1 \\ 2 \end{pmatrix} = \begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ -\frac{1}{2} \end{pmatrix} + \lambda_2 \begin{pmatrix} 2 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{array}{rr} \lambda_1 + 2\lambda_2 &= -\frac{1}{2} \\ -\frac{1}{2}\lambda_1 + \lambda_2 &= 1. \end{array}$$

Als Lösung dieses Gleichungssystems erhält man $\lambda_1 = -\frac{5}{4}$, $\lambda_2 = \frac{3}{8}$. Somit ist die Darstellung des Punktes P bezüglich des Koordinatensystems K : $P(-\frac{5}{4}; \frac{3}{8})_K$.

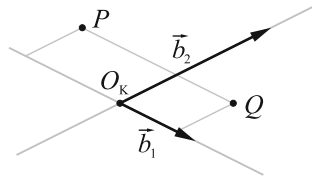
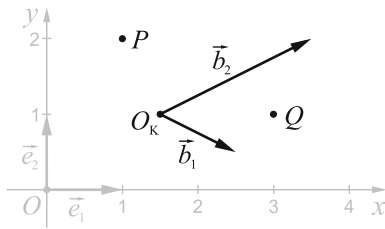


Abb. 5.7:
Koordinatensystem von $\mathbb{R}^2$

Um die Koordinaten von Q bezüglich K zu bestimmen, ist die Gleichung $Q = O + \lambda_1 \vec{b}_1 + \lambda_2 \vec{b}_2$ zu lösen; es ergibt sich $\lambda_1 = \frac{3}{4}$, $\lambda_2 = \frac{3}{8}$, also $Q(\frac{3}{4}; \frac{3}{8})_K$.

P und Q haben dieselbe Koordinate $\lambda_2 = \frac{3}{8}$ bezüglich des Basisvektors $\vec{b}_2$. Dies ist anhand von Abb. 5.7 (rechts) einleuchtend: beide Punkte liegen auf einer Parallelen zu der durch $\vec{b}_1$ gegebenen Koordinatenachse, sie unterscheiden sich also nur durch ihre Koordinate bezüglich $\vec{b}_1$.

Es lassen sich auch für beliebige Punkte $P = \begin{pmatrix} x \\ y \end{pmatrix}$ die Koordinaten λ_1, λ_2 bezüglich K ermitteln. Dazu ist die Vektorgleichung

$$\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \frac{3}{2} \\ 1 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ -\frac{1}{2} \end{pmatrix} + \lambda_2 \begin{pmatrix} 2 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{aligned} \lambda_1 + 2\lambda_2 &= x - \frac{3}{2} \\ -\frac{1}{2}\lambda_1 + \lambda_2 &= y - 1. \end{aligned}$$

nach λ_1, λ_2 zu lösen. Man erhält $\lambda_1 = -y + \frac{1}{2}x + \frac{1}{4}$, $\lambda_2 = \frac{1}{2}y + \frac{1}{4}x - \frac{7}{8}$. Durch diese Gleichungen werden die Koordinaten x, y beliebiger Punkte bezüglich des Standardkoordinatensystems in Koordinaten λ_1, λ_2 bezüglich des Koordinatensystems K überführt, man spricht von einer *Koordinatentransformation*. ■

Beispiel 5.33

In dem Beispiel 5.23 auf S. 187 wurde bereits gezeigt, dass $B = \{\vec{b}_1; \vec{b}_2; \vec{b}_3\}$ mit $\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$ und $\vec{b}_3 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$ eine Basis des Vektorraumes $\mathbb{R}^3$ ist. Durch Hinzunahme eines Punktes, z. B. $O' = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$, als Koordinatenursprung wird B zu einem Koordinatensystem des affinen Punktraumes $\mathbb{R}^3$: $K = \{O'; \vec{b}_1; \vec{b}_2; \vec{b}_3\}$.

Es werden im Folgenden die Koordinaten $\lambda_1, \lambda_2, \lambda_3$ des Punktes $P = \begin{pmatrix} 2 \\ 4 \\ \frac{7}{2} \end{pmatrix}$ bezüglich des Koordinatensystems K bestimmt:

$$\begin{pmatrix} 2 \\ 4 \\ \frac{7}{2} \end{pmatrix} = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{aligned} \lambda_1 + \lambda_2 + \lambda_3 &= 1 \\ \lambda_2 + \lambda_3 &= 2 \\ \lambda_3 &= \frac{1}{2}. \end{aligned}$$

Als Lösung dieses LGS erhält man $\lambda_1 = -1$, $\lambda_2 = \frac{3}{2}$, $\lambda_3 = \frac{1}{2}$ (siehe Abb. 5.8) bzw. $P(-1; \frac{3}{2}; \frac{1}{2})_K$. ■

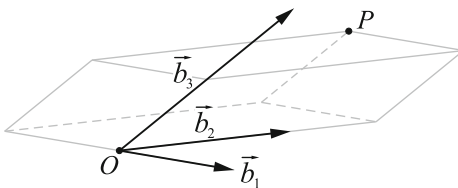


Abb. 5.8:
Koordinatensystem von $\mathbb{R}^3$

Beispiel 5.34

Wir betrachten die durch die Gleichung

$$y = 2x^2 - 12x + 16$$

(in Koordinaten bezüglich des Standardkoordinatensystems von $\mathbb{R}^2$) gegebene Parabel (Abb. 5.9) und suchen nach einem Koordinatensystem $K = \{O; \vec{b}_1; \vec{b}_2\}$, bezüglich dessen diese durch die Gleichung $\lambda_2 = \lambda_1^2$ beschrieben wird. Zunächst formen wir die Gleichung der Parabel mittels quadratischer Ergänzung um:

$$y = 2(x-3)^2 - 2.$$

Der Scheitelpunkt der Parabel wird als Koordinatenursprung des neuen Koordinatensystems festgelegt: $O' = \begin{pmatrix} 3 \\ -2 \end{pmatrix}$. Nun lässt sich die Gleichung der Parabel in Koordinaten x', y' bezüglich des Koordinatensystems $K'_0 = \{O'; \vec{e}_1; \vec{e}_2\}$ mit dem „neuen“ Ursprung und den Basisvektoren der Standardbasis ermitteln. Für die Koordinaten beliebiger Punkte gilt $\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 3 \\ -2 \end{pmatrix} + \begin{pmatrix} x' \\ y' \end{pmatrix}$; durch Einsetzen von $x = x' + 3$ und $y = y' - 2$ erhält man als Gleichung der Parabel

$$y' - 2 = 2(x' + 3 - 3)^2 - 2 \quad \text{bzw.} \quad y' = 2x'^2.$$

Einer der Vektoren der Standardbasis kann offensichtlich für das „neue“ Koordinatensystem K verwendet werden, z. B. $\vec{b}_1 = \vec{e}_1$. Für die zugehörigen Koordinaten gilt dann $\lambda_1 = x'$. Die Parabelgleichung hat damit die Gestalt $\frac{1}{2}y' = \lambda_1^2$.

Es ist noch ein geeigneter Basisvektor $\vec{b}_2$ zu finden, bezüglich dessen sich $\frac{1}{2}y'$ durch λ_2 ersetzen lässt. Damit für beliebige Punkte durch das „alte“ und das „neue“ Koordinatensystem dieselben Positionen beschrieben werden, müssen die Vektoren, die als Produkte des „alten“ und des „neuen“ Basisvektors mit den jeweils zugehörigen Koordinaten entstehen, gleich sein. Somit muss $\lambda_2 \vec{b}_2 = y' \vec{e}_2$ bzw. $\lambda_2 \begin{pmatrix} b_{21} \\ b_{22} \end{pmatrix} = y' \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ gelten. Mit $\lambda_2 = \frac{1}{2}y'$ bedeutet dies $\frac{1}{2} \begin{pmatrix} b_{21} \\ b_{22} \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$, wir erhalten $\vec{b}_2 = \begin{pmatrix} 0 \\ 2 \end{pmatrix}$. Man beachte, dass die *Basisvektoren mit dem Reziproken des Faktors der Koordinaten* zu multiplizieren sind: $\vec{b}_2 = 2\vec{e}_2$ aber $\lambda_2 = \frac{1}{2}y'$.

Bezüglich des Systems $K = \{O'; \vec{b}_1; \vec{b}_2\}$ mit $O' = \begin{pmatrix} 3 \\ -2 \end{pmatrix}$, $\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 0 \\ 2 \end{pmatrix}$ hat die gegebene Parabel also die Gleichung

$$\lambda_2 = \lambda_1^2. \quad \blacksquare$$

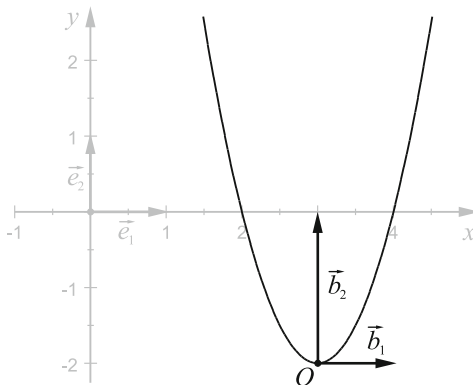


Abb. 5.9: Beschreibung einer Parabel anhand zweier Koordinatensysteme

Beispiel 5.35

In dem Beispiel 2.14 auf S. 79 wurde die Gleichung einer Hyperbel in die Gleichung bezüglich eines anderen Koordinatensystems überführt. Das dort praktizierte Vorgehen lässt sich nun exakt begründen. Gegeben war die Hyperbel mit der Gleichung $y = \frac{1}{x}$ bezüglich des Standardkoordinatensystems (Abb. 5.10). Man benötigt ein Koordinatensystem mit um 45° gedrehten Achsen, um eine Hyperbelgleichung der Form $\frac{\lambda_1^2}{a^2} - \frac{\lambda_2^2}{b^2} = 1$ zu erhalten. Geeignete Basisvektoren sind z. B. $\vec{b}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\vec{b}_2 = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$. Für die Koordinaten λ_1, λ_2 beliebiger Punkte $P = \begin{pmatrix} x \\ y \end{pmatrix}$ bezüglich des Systems $K = \{O; \vec{b}_1; \vec{b}_2\}$ gilt:

$$\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} + \lambda_1 \begin{pmatrix} 1 \\ 1 \end{pmatrix} + \lambda_2 \begin{pmatrix} -1 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{aligned} \lambda_1 - \lambda_2 &= x \\ \lambda_1 + \lambda_2 &= y. \end{aligned}$$

Setzt man dies in die Hyperbelgleichung $y = \frac{1}{x}$ ein, so erhält man

$$\lambda_1 + \lambda_2 = \frac{1}{\lambda_1 - \lambda_2}, \quad (\lambda_1 + \lambda_2)(\lambda_1 - \lambda_2) = 1 \quad \text{bzw.} \quad \lambda_1^2 - \lambda_2^2 = 1.$$

Dabei ist zu beachten, dass die Basisvektoren $\vec{b}_1$ und $\vec{b}_2$ keine Einheitsvektoren sind; sie haben die Beträge $|\vec{b}_1| = \sqrt{1^2 + 1^2} = \sqrt{2}$ und $|\vec{b}_2| = \sqrt{(-1)^2 + 1^2} = \sqrt{2}$. Dieser Aspekt wurde in diesem Abschnitt bisher nicht beachtet (man vergleiche die vorhergehenden Beispiele), da Basisvektoren keine Einheitsvektoren sein müssen. Will man jedoch wie in dem Beispiel 2.14 ein Koordinatensystem mit „gleicher Achsenskalierung“ wie das Standardkoordinatensystem verwenden, so können die Basisvektoren $\vec{b}_1$ und $\vec{b}_2$ zu Einheitsvektoren normiert werden (siehe S. 161). Man erhält damit ein Koordinatensystem $K' = \{O; \vec{b}'_1; \vec{b}'_2\}$ mit $\vec{b}'_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\vec{b}'_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \end{pmatrix}$. Für die Koordinaten λ'_1, λ'_2 beliebiger Punkte $P = \begin{pmatrix} x \\ y \end{pmatrix}$ bezüglich des Systems K' gilt:

$$\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} + \lambda'_1 \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} + \lambda'_2 \frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \end{pmatrix} \quad \text{bzw.} \quad \begin{aligned} \frac{1}{\sqrt{2}} \lambda'_1 - \frac{1}{\sqrt{2}} \lambda'_2 &= x \\ \frac{1}{\sqrt{2}} \lambda'_1 + \frac{1}{\sqrt{2}} \lambda'_2 &= y. \end{aligned}$$

Durch Einsetzen in die Hyperbelgleichung $y = \frac{1}{x}$ erhält man:

$$\begin{aligned} \frac{1}{\sqrt{2}} \lambda'_1 + \frac{1}{\sqrt{2}} \lambda'_2 &= \frac{1}{\frac{1}{\sqrt{2}} \lambda'_1 - \frac{1}{\sqrt{2}} \lambda'_2} \\ \left(\frac{1}{\sqrt{2}} \lambda'_1 + \frac{1}{\sqrt{2}} \lambda'_2 \right) \cdot \left(\frac{1}{\sqrt{2}} \lambda'_1 - \frac{1}{\sqrt{2}} \lambda'_2 \right) &= 1 \\ \frac{\lambda'^2_1}{2} - \frac{\lambda'^2_2}{2} &= 1. \end{aligned}$$

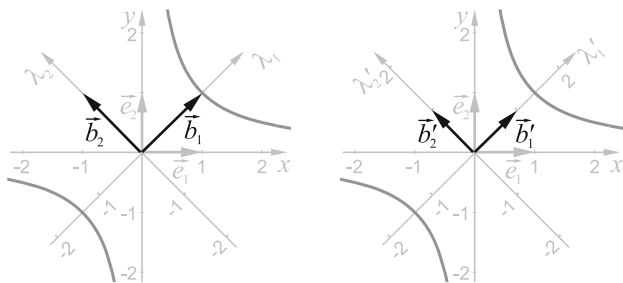


Abb. 5.10:
Beschreibung einer Hyperbel anhand verschiedener Koordinatensysteme

5.5.3 Affine Unterräume und ihre Lagebeziehungen

Definition 5.12

Ist A ein affiner Punktraum über einem Vektorraum V , $P_0 \in A$ ein Punkt und $U \subseteq V$ ein linearer Unterraum von V , so heißt die Menge aller Punkte P mit $P = P_0 + \vec{u}$, $\vec{u} \in U$ *affiner Unterraum* N von A , d. h.:

$$N = \{P \mid P = P_0 + \vec{u}; \vec{u} \in U\}.$$

◆

Beispiel 5.36

Geraden in der Ebene oder im Raum sind affine Unterräume. Dies folgt daraus, dass sich jede Gerade durch eine Parameterdarstellung $\vec{p} = \vec{p}_0 + t\vec{a}$ ($\vec{a} \neq \vec{0}$) beschreiben lässt (siehe Abschnitt 4.1.1), die in Punktschreibweise die Gestalt

$$P = P_0 + t\vec{a} \quad (t \in \mathbb{R})$$

hat. Dabei ist P_0 ein fester Punkt der Ebene oder des Raumes (bzw. des affinen Raumes $\mathbb{R}^2$ oder $\mathbb{R}^3$). Die Menge aller Vektoren $\vec{u} = t\vec{a}$ ($t \in \mathbb{R}$), also die Menge aller zu einem Vektor $\vec{a}$ kollinearen Vektoren, ist ein linearer Unterraum von $\mathbb{R}^2$ bzw. $\mathbb{R}^3$, siehe Beispiel 5.7 auf S. 173. ■

Beispiel 5.37

Ebenso wie Geraden sind auch *Ebenen* affine Unterräume von $\mathbb{R}^3$. Ebenen lassen sich durch Parameterdarstellungen $\vec{p} = \vec{p}_0 + s\vec{a} + t\vec{b}$ mit zwei nicht kollinearen Vektoren $\vec{a}$ und $\vec{b}$ beschreiben (siehe Abschnitt 4.2.1). Dabei ist $\vec{p}_0$ der Ortsvektor eines festen Punktes (Aufpunktes) und $\vec{p}$ sind Ortsvektoren beliebiger Punkte P der jeweils beschriebenen Ebene. Die Parameterdarstellung einer Ebene lässt sich somit auch in der Form

$$P = P_0 + s\vec{a} + t\vec{b} \quad (s, t \in \mathbb{R})$$

schreiben. Die Menge aller Vektoren $\vec{u} = s\vec{a} + t\vec{b}$ ($s, t \in \mathbb{R}$) ist die lineare Hülle $\langle \vec{a}, \vec{b} \rangle$ der Vektoren $\vec{a}$ und $\vec{b}$, von der nach Satz 5.4 (S. 178) bekannt ist, dass es sich um einen linearen Unterraum handelt. Somit lässt sich eine Ebene als Punktmenge auch in der folgenden Form angeben, anhand derer der Bezug zu der Definition 5.12 besonders deutlich wird:

$$\varepsilon = \{P \mid P = P_0 + \vec{u}; \vec{u} \in \langle \vec{a}, \vec{b} \rangle\}.$$

■

Satz 5.24

Ist $N = \{P \mid P = P_0 + \vec{u}; \vec{u} \in U\}$ ein affiner Unterraum und $Q \in N$, so wird durch $M = \{P \mid P = Q + \vec{v}; \vec{v} \in U\}$ derselbe affine Unterraum beschrieben, d. h. $M = N$.

Satz 5.25

Ist A ein affiner Punktraum und N ein affiner Unterraum von A , so ist N selbst ein affiner Punktraum.

Die Beweise der Sätze 5.25 und 5.24 seien der Leserin bzw. dem Leser überlassen.

Wie anhand der Beispiele 5.36 und 5.37 deutlich wird, stellen affine Unterräume Verallgemeinerungen von Geraden und Ebenen dar, die bislang in zwei- oder dreidimensionalen Räumen behandelt wurden. Es lassen sich nun m -dimensionale affine Unterräume beliebiger (n -dimensionaler) affiner Räume betrachten (wobei $m \leq n$ sein muss).

Definition 5.13

Als *Dimension eines affinen Punktraumes* bzw. *eines affinen Unterraumes* bezeichnet man die Dimension des zugehörigen Vektorraumes bzw. linearen Unterraumes. $\blacklozenge$

Nach dieser Definition sind *Punkte* affine Unterräume der Dimension Null, denn der zugehörige Vektorraum enthält nur den Nullvektor (vgl. Definition 5.9 auf S. 197). *Geraden* sind ein- und *Ebenen* zweidimensionale affine Unterräume.

Satz 5.26

Es sei A ein affiner Punktraum über einem Vektorraum V , und es seien N_1 und N_2 affine Unterräume von A mit zugehörigen linearen Unterräumen U_1 und U_2 von V . Ist die Schnittmenge $N_1 \cap N_2$ nicht leer, so ist $N_1 \cap N_2$ ein affiner Unterraum von A mit dem zugehörigen linearen Unterraum $U_1 \cap U_2$.

Bemerkung: Die Bedingung $N_1 \cap N_2 \neq \{\}$ muss unbedingt gestellt werden, denn sie ist bei affinen Unterräumen nicht selbstverständlich. So ist die Schnittmenge zweier paralleler oder windschiefer Geraden leer und somit kein affiner Unterraum. Hingegen besitzen beliebige lineare Unterräume stets eine nicht leere Schnittmenge, denn der Nullvektor ist Element jedes linearen Unterraumes.

Beweis: Es seien $N_1 = \{P \mid P = P_1 + \vec{u}; \vec{u} \in U_1\}$, $N_2 = \{P \mid P = P_2 + \vec{v}; \vec{v} \in U_2\}$ affine Unterräume. Nach Voraussetzung ist $N_1 \cap N_2 \neq \{\}$, es existiert also ein Punkt S mit $S \in N_1$ und $S \in N_2$. Wegen des Satzes 5.24 lassen sich N_1 und N_2 auch folgendermaßen beschreiben:

$$N_1 = \{P \mid P = S + \vec{u}; \vec{u} \in U_1\}, \quad N_2 = \{P \mid P = S + \vec{v}; \vec{v} \in U_2\}.$$

Die Schnittmenge $N_1 \cap N_2$ besteht somit genau aus denjenigen Punkten P , für die Vektoren $\vec{u} \in U_1$, $\vec{v} \in U_2$ mit $P = S + \vec{u} = S + \vec{v}$ existieren. Nach Definition 5.10 existiert zu zwei Punkten S und P eines affinen Raumes *genau* ein Vektor $\vec{x} \in V$ mit $P = S + \vec{x}$. Somit muss $\vec{u} = \vec{v}$ sein, und aus den Bedingungen $\vec{u} \in U_1$ und $\vec{v} \in U_2$ resultiert $\vec{u} \in U_1 \cap U_2$. Somit ist

$$N_1 \cap N_2 = \{P \mid P = S + \vec{u}; \vec{u} \in U_1 \cap U_2\},$$

d. h. ein affiner Unterraum mit dem zugehörigen linearen Unterraum $U_1 \cap U_2$. $\square$

Folgerung: Sind N_1 und N_2 affine Unterräume mit zugehörigen linearen Unterräumen U_1 , U_2 und ist $N_1 \cap N_2 \neq \{\}$, so gilt $\dim(N_1 \cap N_2) = \dim(U_1 \cap U_2)$.

Beispiel 5.38

Die Notwendigkeit der Voraussetzung $N_1 \cap N_2 \neq \{\}$ in dem Satz 5.26 und in der Folgerung zeigt sich am Beispiel zweier Ebenen $\varepsilon_1: P = P_0 + s\vec{a} + t\vec{b}$ und $\varepsilon_2: P = Q_0 + \lambda\vec{c} + \mu\vec{d}$ ($s, t, \lambda, \mu \in \mathbb{R}$) in $\mathbb{R}^3$. Die zugehörigen linearen Unterräume sind die linearen Hüllen von $\vec{a}, \vec{b}$ bzw. von $\vec{c}, \vec{d}$, also $U_1 = \langle \vec{a}, \vec{b} \rangle$ und $U_2 = \langle \vec{c}, \vec{d} \rangle$. Es bestehen drei Möglichkeiten der gegenseitigen Lage von ε_1 und ε_2 .

- Die beiden Ebenen sind identisch; dann ist $\varepsilon_1 \cap \varepsilon_2 = \varepsilon_1 = \varepsilon_2$, und für die Dimension der Schnittmenge gilt $\dim(\varepsilon_1 \cap \varepsilon_2) = 2$. Auch die zugehörigen linearen Unterräume sind identisch, somit gilt $\dim(U_1 \cap U_2) = 2$.

- ε_1 und ε_2 schneiden sich in einer Geraden, d. h. $\dim(\varepsilon_1 \cap \varepsilon_2) = 1$. Die Schnittmenge der linearen Unterräume U_1 und U_2 bilden alle Vektoren, die Vielfache eines Richtungsvektors der Schnittgeraden sind, also $\dim(U_1 \cap U_2) = 1$.
- Falls ε_1 und ε_2 parallel sind, so ist ihre Schnittmenge leer, besitzt also keine Dimension. Die zugehörigen linearen Unterräume sind hingegen identisch, also ist $\dim(U_1 \cap U_2) = 2$. ■

Eine effektive und systematische Untersuchung von Lagebeziehungen beliebiger affiner Unterräume wird mithilfe einer Dimensionsformel für affine Unterräume möglich sein. Die in dem Abschnitt 5.4.6 bewiesene Dimensionsformel für lineare Unterräume ist nicht ohne Einschränkungen auf affine Unterräume übertragbar, wie anhand des Beispiels 5.38 sichtbar wurde. Zudem benötigen wir ein Analogon zu der Summe zweier linearer Unterräume (vgl. Def. 5.3 auf S. 176).

Definition 5.14

Es seien A ein affiner Punktraum über einem Vektorraum V sowie N_1 und N_2 affine Unterräume von A . Der kleinste affine Unterraum, der N_1 und N_2 enthält, wird als *affine Hülle* (oder auch *Verbindungsraum*) $A(N_1 \cup N_2)$ bezeichnet. ◆

Beispiel 5.39

Die Geraden $g: P = P_0 + t\vec{a} = \begin{pmatrix} 4 \\ 2 \\ 1 \end{pmatrix} + t \begin{pmatrix} -2 \\ -1 \\ 0,5 \end{pmatrix}$, $h: P = Q_0 + s\vec{b} = \begin{pmatrix} -2 \\ 4 \\ 1 \end{pmatrix} + s \begin{pmatrix} -1 \\ -3 \\ 1 \end{pmatrix}$ schneiden sich in $S = \begin{pmatrix} -4 \\ -2 \\ 3 \end{pmatrix}$, siehe Beispiel 4.8 (S. 145). Die affine Hülle von g und h ist die Ebene $\varepsilon: P = S + t\vec{a} + s\vec{b}$. Der zugehörige lineare Unterraum ist die lineare Hülle $\langle \vec{a}, \vec{b} \rangle$ der beiden Richtungsvektoren der Ebene ε . Die zu g und h gehörenden linearen Unterräume U_1 und U_2 sind die linearen Hüllen der Richtungsvektoren der Geraden, d. h. $U_1 = \langle \vec{a} \rangle$, $U_2 = \langle \vec{b} \rangle$; ihre Schnittmenge besteht nur aus dem Nullvektor: $U_1 \cap U_2 = \{\vec{0}\}$. Es gilt $\dim g = \dim h = 1$, $\dim(U_1 \cap U_2) = 0$, $\dim A(g \cup h) = \dim(\langle \vec{a}, \vec{b} \rangle) = \dim(U_1 + U_2) = 2$. ■

Beispiel 5.40

Die Geraden $g: P = P_0 + t\vec{a} = \begin{pmatrix} -1 \\ -2 \\ 2 \end{pmatrix} + t \begin{pmatrix} -4 \\ 6 \\ -2 \end{pmatrix}$, $h: P = Q_0 + s\vec{b} = \begin{pmatrix} 1 \\ 3 \\ 2 \end{pmatrix} + s \begin{pmatrix} -3 \\ 3 \\ -1 \end{pmatrix}$ sind windschief, siehe Beispiel 4.9 auf S. 145. Die affine Hülle von g und h ist in diesem Falle der dreidimensionale Raum; sie umfasst alle Linearkombinationen von $\vec{a}$, $\vec{b}$ und $\overrightarrow{P_0Q_0}$ (diese drei Vektoren sind nicht komplanar). Die Schnittmenge der zu g und h gehörenden linearen Unterräume enthält auch hier nur den Nullvektor: $U_1 \cap U_2 = \{\vec{0}\}$. Es gilt also $\dim g = \dim h = 1$, $\dim(U_1 \cap U_2) = 0$, $\dim(\langle \vec{a}, \vec{b} \rangle) = \dim(U_1 + U_2) = 2$, aber $\dim A(g \cup h) = 3$. ■

Beispiel 5.41

Es seien $\varepsilon_1: P = P_0 + s\vec{a} + t\vec{b}$ und $\varepsilon_2: P = Q_0 + \lambda\vec{c} + \mu\vec{d}$ Ebenen in $\mathbb{R}^3$ sowie $U_1 = \langle \vec{a}, \vec{b} \rangle$, $U_2 = \langle \vec{c}, \vec{d} \rangle$ die zugehörigen linearen Unterräume.

- Sind ε_1 und ε_2 identisch, so ist $A(\varepsilon_1 \cup \varepsilon_2) = \varepsilon_1 = \varepsilon_2$, also $\dim A(\varepsilon_1 \cup \varepsilon_2) = 2$. Außerdem gilt, da wegen $\varepsilon_1 = \varepsilon_2$ auch $U_1 = U_2$ sein muss, $\dim(U_1 + U_2) = 2$.

- Falls sich ε_1 und ε_2 in einer Geraden schneiden, so gehört mindestens einer der Richtungsvektoren von ε_2 nicht zu der linearen Hülle $\langle \vec{a}, \vec{b} \rangle$, somit ist $\dim(U_1 + U_2) = 3$. Der kleinste affine Raum, der beide Ebenen enthält, ist der gesamte Raum $\mathbb{R}^3$, d. h. $\dim A(\varepsilon_1 \cup \varepsilon_2) = 3$.
- Ist $\varepsilon_1 \cap \varepsilon_2 = \{\}$, d. h. sind ε_1 und ε_2 parallel, so sind die zugehörigen linearen Unterräume identisch, und somit gilt $\dim(U_1 + U_2) = 2$. Der kleinste affine Raum, der ε_1 und ε_2 enthält, ist allerdings keine Ebene, sondern der gesamte Raum $\mathbb{R}^3$, d. h. $\dim A(\varepsilon_1 \cup \varepsilon_2) = 3$. ■

Analoge Überlegungen lassen sich für eine Gerade und eine Ebene in $\mathbb{R}^3$ (bei Unterscheidung der dabei möglichen Lagebeziehungen, vgl. Abschnitt 4.2.2) anstellen – der Leserin bzw. dem Leser wird empfohlen, vor dem weiteren Studium die diesbezügliche Aufgabe 9 (S. 218) zu bearbeiten. Ohne Beweis vermerken wir, dass – in Verallgemeinerung aller betrachteten Beispiele – der folgende Satz gilt.

Satz 5.27

Es seien A ein affiner Raum über einem Vektorraum V sowie N_1 und N_2 affine Unterräume von A mit zugehörigen linearen Unterräumen U_1 und U_2 von V .

- (i) *Ist $N_1 \cap N_2 \neq \{\}$, so gilt $\dim A(N_1 \cup N_2) = \dim(U_1 + U_2)$.*
- (ii) *Ist $N_1 \cap N_2 = \{\}$, so gilt $\dim A(N_1 \cup N_2) = \dim(U_1 + U_2) + 1$*

Aus dem Satz 5.27 und der Dimensionsformel für lineare Unterräume (Satz 5.21 auf S. 199) ergeben sich die folgenden *Dimensionsformeln für affine Unterräume*.

Satz 5.28

Es seien A ein affiner Raum über einem Vektorraum V sowie N_1 und N_2 affine Unterräume von A mit zugehörigen linearen Unterräumen U_1 und U_2 von V .

- (i) *Ist $N_1 \cap N_2 \neq \{\}$, so gilt*

$$\dim N_1 + \dim N_2 = \dim A(N_1 \cup N_2) + \dim(U_1 \cap U_2).$$
- (ii) *Ist $N_1 \cap N_2 = \{\}$, so gilt*

$$\dim N_1 + \dim N_2 = \dim A(N_1 \cup N_2) + \dim(U_1 \cap U_2) - 1.$$

Mithilfe dieser Dimensionsformeln lassen sich systematisch Lagebeziehungen affiner Unterräume von affinen Punkträumen beliebiger Dimension untersuchen. Für die Unterräume eines affinen Raumes der Dimension 3 ist dies mit den Beispielen 5.38-5.41 sowie den Aufgaben 8 und 9 (siehe S. 218) bereits erfolgt – der Leserin bzw. dem Leser sei empfohlen, die auftretenden Fälle in einer Übersicht wie der folgenden zusammenzufassen.

Lagebeziehungen zweier „Ebenen“ eines vierdimensionalen affinen Raumes

Es werden nun alle möglichen Lagebeziehungen zweier zweidimensionaler Unterräume ε_1 und ε_2 eines vierdimensionalen affinen Raumes A untersucht. Es gilt $2 \leq \dim A(\varepsilon_1 \cup \varepsilon_2) \leq 4$, da beide Ebenen innerhalb von A liegen und ihre affine Hülle daher maximal A selbst sein kann, zugleich aber mindestens die Ebenen selbst enthält. Wir untersuchen für $\varepsilon_1 \cap \varepsilon_2 = \{\}$ und $\varepsilon_1 \cap \varepsilon_2 \neq \{\}$ systematisch die möglichen Fälle $\dim A(\varepsilon_1 \cup \varepsilon_2) = 2, 3, 4$.

- ε_1 und ε_2 besitzen keinen Schnittpunkt, d. h. $\varepsilon_1 \cap \varepsilon_2 = \{\}$.

In diesem Falle gilt wegen $\dim \varepsilon_1 = \dim \varepsilon_2 = 2$ nach Satz 5.28 (ii):

$$4 = \dim A(\varepsilon_1 \cup \varepsilon_2) + \dim (U_1 \cap U_2) - 1 \Rightarrow \dim A(\varepsilon_1 \cup \varepsilon_2) + \dim (U_1 \cap U_2) = 5.$$

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 2$. In diesem Falle müsste $\dim (U_1 \cap U_2) = 5 - 2 = 3$ sein, was wegen $\dim U_1 = \dim U_2 = 2$ *nicht möglich* ist.

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 3$. Es gilt $\dim (U_1 \cap U_2) = 5 - 3 = 2$; damit ist $U_1 = U_2$; die Ebenen ε_1 und ε_2 *sind parallel*.

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 4$. Es gilt $\dim (U_1 \cap U_2) = 1$. U_1 und U_2 haben *eine gemeinsame Richtung*, d. h. es gibt eine Gerade g_1 in der Ebene ε_1 , die zu einer Geraden g_2 in der Ebene ε_1 parallel ist.

- ε_1 und ε_2 besitzen mindestens einen gemeinsamen Punkt, d. h. $\varepsilon_1 \cap \varepsilon_2 \neq \{\}$.

In diesem Falle gilt nach Satz 5.28 (i): $\dim A(\varepsilon_1 \cup \varepsilon_2) + \dim (U_1 \cap U_2) = 4$.

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 2$. In diesem Falle ist $\dim (U_1 \cap U_2) = 2$, d. h. $U_1 = U_2$. Die Ebenen ε_1 und ε_2 müssen daher parallel oder identisch sein. Wegen $\varepsilon_1 \cap \varepsilon_2 \neq \{\}$ sind ε_1 und ε_2 somit *identisch*.

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 3$. Es gilt $\dim (U_1 \cap U_2) = 1$; da die Schnittmenge $\varepsilon_1 \cap \varepsilon_2$ nicht leer ist, haben die Ebenen eine *Schnittgerade*.

$\dim A(\varepsilon_1 \cup \varepsilon_2) = 4$. In diesem Falle ist $\dim (U_1 \cap U_2) = 4 - 4 = 0$. $U_1 \cap U_2$ besteht nur aus dem Nullvektor. Die Ebenen ε_1 und ε_2 haben *genau einen gemeinsamen Punkt*.

Damit sind alle möglichen Fälle der gegenseitigen Lage zweier Ebenen in einem vierdimensionalen Raum untersucht. Insbesondere der Fall, dass zwei Ebenen genau einen Schnittpunkt haben, ist anschaulich schwer vorstellbar – ein vierdimensionaler Raum entzieht sich der unmittelbaren Anschauung. Für einen dreidimensionalen Raum lässt sich anhand der Dimensionsformeln erklären, dass sich zwei Ebenen nicht in genau einem Punkt schneiden können. Dazu müsste nämlich $\dim (U_1 \cap U_2) = 0$ sein. Da aber, falls überhaupt gemeinsame Punkte existieren, $\dim \varepsilon_1 + \dim \varepsilon_2 = \dim A(\varepsilon_1 \cup \varepsilon_2) + \dim (U_1 \cap U_2)$ gilt, würde daraus $\dim A(\varepsilon_1 \cup \varepsilon_2) = 4$ folgen, was im Dreidimensionalen nicht möglich ist.

5.5.4 Anwendung auf die Theorie der linearen Gleichungssysteme

Zusammenhänge zwischen Lösungsmengen linearer Gleichungssysteme und affinen Unterräumen waren in diesem Buch bereits mehrfach von Bedeutung:

- Geraden (d. h. eindimensionale affine Unterräume) im Raum (bzw. in $\mathbb{R}^3$) sind Lösungsmengen linearer Gleichungssysteme mit drei Variablen und zwei Gleichungen, siehe Abschnitt 4.1.2. Umgekehrt hat jedes lösbare LGS mit drei Variablen vom Rang 2 eine einparametrische Lösungsmenge (siehe S. 35)

$$\text{der Form } L = \left\{ \begin{pmatrix} x \\ y \\ z \end{pmatrix} \left| \begin{pmatrix} x \\ y \\ z \end{pmatrix} = t \cdot \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} + \begin{pmatrix} x_0 \\ y_0 \\ z_0 \end{pmatrix}; t \in \mathbb{R} \right. \right\}.$$

- Zweidimensionale affine Unterräume (Ebenen) von $\mathbb{R}^3$ lassen sich durch Gleichungen der Form $ax + by + cz = d$ beschreiben (siehe Abschnitt 4.2.1). Umgekehrt hat jede lineare Gleichung mit drei Variablen eine zweiparametrische Lösungsmenge der Form $L = \{\vec{x} \mid \vec{x} = \vec{x}_0 + s\vec{m} + t\vec{n}; s, t \in \mathbb{R}\}$, beschreibt also eine Ebene (siehe Abschnitt 1.2.1, speziell S. 18).
- Lösungsmengen homogener linearer Gleichungssysteme in n Variablen sind lineare Unterräume von $\mathbb{R}^n$, siehe Beispiel 5.9 auf S. 173.

Eine Verallgemeinerung dieser Tatsachen stellt der folgende Satz dar.

Satz 5.29

Die Lösungsmenge eines lösbaren linearen Gleichungssystems in n Variablen vom Rang r ist ein $n-r$ -dimensionaler affiner Unterraum des affinen Raumes $\mathbb{R}^n$. Der zugehörige lineare Unterraum ist die Lösungsmenge des dem gegebenen LGS zugeordneten homogenen linearen Gleichungssystems.

Umgekehrt ist jeder affine Unterraum der Dimension m von $\mathbb{R}^n$ als Lösungsmenge eines LGS mit n Variablen und dem Rang $n-m$ darstellbar.

Beweis: Es sei ein beliebiges lösbares LGS mit n Variablen gegeben:

$$\begin{aligned}
 & a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = b_1 \\
 (*) \quad & a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = b_2 \\
 & \vdots \qquad \qquad \qquad \vdots \qquad \qquad \qquad \vdots \qquad \qquad \qquad \vdots \\
 & a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n = b_m
 \end{aligned}$$

bzw. in Kurzschreibweise

$$(*) \quad \sum_{j=1}^n a_{ij}x_j = b_i \quad (i = 1 \dots m).$$

Da dieses LGS lösbar ist, existiert eine Lösung $L^{(0)} = \begin{pmatrix} x_1^{(0)} \\ \vdots \\ x_n^{(0)} \end{pmatrix}$.

Es sei U die Lösungsmenge des zu (*) gehörenden homogenen LGS

$$(**) \quad \sum_{j=1}^n a_{ij}x_j = 0 \quad (i = 1 \dots m).$$

Aus dem Beispiel 5.9 (S. 173) ist bekannt, dass U ein linearer Unterraum von $\mathbb{R}^n$ ist. Wegen des Satzes 1.3 (S. 38) ist jede Lösung L des Gleichungssystems (*) als Summe der Lösung $L^{(0)}$ und einer Lösung $\vec{u} \in U$ des homogenen LGS (**) darstellbar, und umgekehrt ist für jede Lösung $\vec{u} \in U$ von (**) die Summe $L^{(0)} + \vec{u}$ eine Lösung des gegebenen (inhomogenen) LGS (*). Die Lösungsmenge des gegebenen LGS ist also genau die Menge

$$\mathcal{L} = \left\{ L \mid L = L^{(0)} + \vec{u}; \vec{u} \in U \right\}.$$

Nach Definition 5.12 handelt es sich dabei um einen affinen Unterraum von $\mathbb{R}^n$.

Bereits in dem Abschnitt 1.3.3 wurde festgestellt, dass ein lösbares LGS in n Variablen mit dem Rang r eine $n-r$ -parametrische Lösungsmenge besitzt (siehe S. 35). Die Anzahl der Parameter ist gleich der Anzahl der Basisvektoren

des zu dem Lösungsraum $\mathcal{L}$ gehörenden linearen Unterraumes U . Somit ist die Lösungsmenge des gegebenen LGS ein $n-r$ -dimensionaler affiner Unterraum.

Wir verzichten auf den Beweis der umgekehrten Richtung des Satzes; die grundsätzliche Beweisidee wird anhand des folgenden Beispiels deutlich. $\square$

Beispiel 5.42

Gegeben ist der affine Unterraum $N = \{P \mid P = P_0 + \vec{u}; \vec{u} \in U\}$ von $\mathbb{R}^4$ mit

$$P_0 = \begin{pmatrix} 5 \\ 3 \\ -1 \\ 0 \end{pmatrix} \text{ und } U = \left\langle \begin{pmatrix} 2 \\ -1 \\ 4 \\ -5 \end{pmatrix}; \begin{pmatrix} -3 \\ 1 \\ -2 \\ 1 \end{pmatrix}; \begin{pmatrix} 4 \\ 2 \\ 0 \\ -3 \end{pmatrix} \right\rangle.$$

Gesucht ist ein Gleichungssystem, als dessen Lösungsmenge sich N ergibt.

Für jeden Punkt $P = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \in N$ muss die Vektorgleichung

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = \begin{pmatrix} 5 \\ 3 \\ -1 \\ 0 \end{pmatrix} + \lambda_1 \begin{pmatrix} 2 \\ -1 \\ 4 \\ -5 \end{pmatrix} + \lambda_2 \begin{pmatrix} -3 \\ 1 \\ -2 \\ 1 \end{pmatrix} + \lambda_3 \begin{pmatrix} 4 \\ 2 \\ 0 \\ -3 \end{pmatrix}$$

bzw. das LGS

$$\begin{aligned} \text{(I)} \quad & 2\lambda_1 - 3\lambda_2 + 4\lambda_3 = x_1 - 5 \\ \text{(II)} \quad & -\lambda_1 + \lambda_2 + 2\lambda_3 = x_2 - 3 \\ \text{(III)} \quad & 4\lambda_1 - 2\lambda_2 = x_3 + 1 \\ \text{(IV)} \quad & -5\lambda_1 + \lambda_2 - 3\lambda_3 = x_4 \end{aligned}$$

erfüllt sein. Dieses LGS wird nun in ein Gleichungssystem umgeformt, in dem nur noch x_1, x_2, x_3, x_4 auftreten. Um die zu eliminierenden Parameter $\lambda_1, \lambda_2, \lambda_3$ zu berechnen (d. h. durch $x_1, \dots, x_4$ auszudrücken), benötigt man drei Gleichungen, die ein LGS mit dem Rang 3 bilden. Wir lösen das aus den Gleichungen (I)-(III) bestehende LGS (welches diese Bedingung erfüllt) und erhalten:

$$\begin{aligned} \lambda_1 &= \frac{5}{12}x_3 + \frac{1}{3}x_2 - \frac{1}{6}x_1 + \frac{1}{4}, \\ \lambda_2 &= \frac{1}{3}x_3 + \frac{2}{3}x_2 - \frac{1}{3}x_1, \\ \lambda_3 &= \frac{1}{24}x_3 + \frac{1}{3}x_2 + \frac{1}{12}x_1 - \frac{11}{8}. \end{aligned}$$

Die so gewonnenen Ausdrücke setzen wir nun für $\lambda_1, \lambda_2, \lambda_3$ in die verbliebene (noch nicht verwendete) Gleichung (IV) ein:

$$\frac{1}{4}x_1 - 2x_2 - \frac{15}{8}x_3 - x_4 = -\frac{23}{8}.$$

Der affine Unterraum N ist die Lösungsmenge dieses aus einer Gleichung bestehenden LGS; es ist $r=1$ und $\dim N=3$, dies entspricht Satz 5.29. $\blacksquare$

Beispiel 5.43

In den Beispielen 5.12 auf S. 174 und 5.29 auf S. 197 wurde die Menge der magischen Quadrate der Kantenlänge 3 betrachtet. Mit variabler Zeilen-, Spalten- und Diagonalensumme ist diese ein dreidimensionaler linearer Unterraum von $\mathbb{R}^{10}$. Wir betrachten nun *magische Quadrate mit fester Zeilen-, Spalten- und Diagonalensumme* s . Damit enthält das LGS auf S. 175 nur noch 9 Variablen, da s nun eine Konstante ist, und wird zu einem inhomogenen LGS:

$$\begin{array}{rclcl}
a_{11} + a_{12} + a_{13} & & & & = s \\
& a_{21} + a_{22} + a_{23} & & & = s \\
& & a_{31} + a_{32} + a_{33} & & = s \\
a_{11} & + a_{21} & + a_{31} & & = s \\
& a_{12} & + a_{22} & + a_{32} & = s \\
& & a_{13} & + a_{23} & + a_{33} = s \\
a_{11} & & + a_{22} & & + a_{33} = s \\
& a_{13} & + a_{22} & + a_{31} & = s
\end{array}$$

Dieses LGS hat die folgende Lösungsmenge:

$$L = \left\{ \begin{pmatrix} a_{11} \\ a_{12} \\ a_{13} \\ a_{21} \\ a_{22} \\ a_{23} \\ a_{31} \\ a_{32} \\ a_{33} \end{pmatrix} \mid \begin{pmatrix} a_{11} \\ a_{12} \\ a_{13} \\ a_{21} \\ a_{22} \\ a_{23} \\ a_{31} \\ a_{32} \\ a_{33} \end{pmatrix} = \begin{pmatrix} \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \\ \frac{1}{3}s \end{pmatrix} + \lambda_1 \begin{pmatrix} -1 \\ 1 \\ 0 \\ 1 \\ 0 \\ -1 \\ 0 \\ -1 \\ 1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 1 \\ -1 \\ -1 \\ 0 \\ 1 \\ 1 \\ -1 \\ 0 \end{pmatrix}; \lambda_1, \lambda_2 \in \mathbb{R} \right\}.$$

Diese Lösungsmenge lässt sich auch anhand der Tatsache ermitteln, dass sich jede Lösung eines inhomogenen LGS als Summe einer festen Lösung dieses LGS und einer Lösung des zugehörigen homogenen LGS ergibt (Satz 1.3 auf S. 38). Eine spezielle Lösung des inhomogenen LGS lässt sich aus dem Gleichungssystem leicht entnehmen, nämlich das 9-Tupel, dessen sämtliche Komponenten $\frac{1}{3}s$ sind. Diese entspricht einem Punkt P_0 des affinen Unterraumes, der durch das inhomogene LGS beschrieben wird. Das zugehörige homogene LGS (welches entsteht, wenn auf den rechten Seiten aller Gleichungen s durch 0 ersetzt wird) lässt sich etwas einfacher als das inhomogene Gleichungssystem lösen. Man erhält als Lösungsmenge einen zweidimensionalen linearen Unterraum von $\mathbb{R}^9$ mit den beiden in der Lösungsmenge oben angegebenen Basisvektoren. In der Matrizenschreibweise ist die Menge aller magischen Quadrate mit der festen Zeilen-, Spalten- und Diagonalsumme s also der affine Teilraum $N = \{P \mid P = P_0 + \vec{u}; \vec{u} \in U\}$ von $\mathbb{R}^9$ mit

$$P_0 = \begin{pmatrix} \frac{1}{3}s & \frac{1}{3}s & \frac{1}{3}s \\ \frac{1}{3}s & \frac{1}{3}s & \frac{1}{3}s \\ \frac{1}{3}s & \frac{1}{3}s & \frac{1}{3}s \end{pmatrix} \quad \text{und} \quad U = \left\langle \begin{pmatrix} -1 & 1 & 0 \\ 1 & 0 & -1 \\ 0 & -1 & 1 \end{pmatrix}, \begin{pmatrix} 0 & 1 & -1 \\ -1 & 0 & 1 \\ 1 & -1 & 0 \end{pmatrix} \right\rangle.$$

Kommen wir nun zurück zu dem Ausgangspunkt unserer Überlegungen zu magischen Quadraten, zu dem berühmtesten magischen Quadrat, siehe S. 174. Dieses hat die Zeilen-, Spalten- und Diagonalsumme $s = 15$; wir erhalten es mit $\lambda_1 = 1$ und $\lambda_2 = 3$ als „Punkt“ des ermittelten affinen Unterraumes:

$$\begin{pmatrix} 5 & 5 & 5 \\ 5 & 5 & 5 \\ 5 & 5 & 5 \end{pmatrix} + 1 \cdot \begin{pmatrix} -1 & 1 & 0 \\ 1 & 0 & -1 \\ 0 & -1 & 1 \end{pmatrix} + 3 \cdot \begin{pmatrix} 0 & 1 & -1 \\ -1 & 0 & 1 \\ 1 & -1 & 0 \end{pmatrix} = \begin{pmatrix} 4 & 9 & 2 \\ 3 & 5 & 7 \\ 8 & 1 & 6 \end{pmatrix}. \quad \blacksquare$$

5.5.5 Aufgaben zu Abschnitt 5.5

1. Beweisen Sie die Behauptungen 2, 3, 5 und 6 des Satzes 5.22 auf S. 204.
2. Bestimmen Sie die Koordinaten des Punktes $P = \begin{pmatrix} -6 \\ -11 \end{pmatrix}$ bezüglich des Koordinatensystems $K = \{O_K; \vec{b}_1; \vec{b}_2\}$ mit $O_K = \begin{pmatrix} 7 \\ 1 \end{pmatrix}$, $b_1 = \begin{pmatrix} -1 \\ 4 \end{pmatrix}$, $b_2 = \begin{pmatrix} 3 \\ 8 \end{pmatrix}$.
3. Bestimmen Sie die Koordinaten der Punkte $P = \begin{pmatrix} -4 \\ 8 \\ 37 \end{pmatrix}$ und $Q = \begin{pmatrix} 19 \\ -17 \\ -52 \end{pmatrix}$ bezüglich des folgenden Koordinatensystems von $\mathbb{R}^3$: $K = \{O_K; \vec{b}_1; \vec{b}_2; \vec{b}_3\}$ mit $O_K = \begin{pmatrix} 5 \\ 0 \\ 5 \end{pmatrix}$, $\vec{b}_1 = \begin{pmatrix} 3 \\ 2 \\ 12 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} -8 \\ 7 \\ 13 \end{pmatrix}$ und $\vec{b}_3 = \begin{pmatrix} 2 \\ 5 \\ -4 \end{pmatrix}$.
4. Gegeben ist eine Parabel durch die Gleichung $y = \frac{1}{2}x^2 + 4x + 11$ bezüglich des Standardkoordinatensystems von $\mathbb{R}^2$. Ermitteln Sie ein Koordinatensystem $K = \{O_K; \vec{b}_1; \vec{b}_2\}$, bezüglich dessen die Parabel durch die Gleichung $\lambda_2 = \lambda_1^2$ beschrieben wird.
5. Eine Parabel hat die Gleichung $\lambda_2 = \lambda_1^2$ bezüglich des Koordinatensystems $K = \{O_K; \vec{b}_1; \vec{b}_2\}$ mit $O_K = \begin{pmatrix} -2 \\ 5 \end{pmatrix}$, $\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 0 \\ 4 \end{pmatrix}$. Geben Sie die Gleichung dieser Parabel bezüglich des Standardkoordinatensystems von $\mathbb{R}^2$ an.
6. Beweisen Sie den Satz 5.24 auf S. 210.
7. Beweisen Sie den Satz 5.25: Ist A ein affiner Punktraum und N ein affiner Unterraum von A , so ist N selbst ein affiner Punktraum.
8. Gegeben sind zwei Geraden g_1 und g_2 in $\mathbb{R}^3$ durch die Parameterdarstellungen $g_1: P = P_0 + s\vec{a}$ und $g_2: P = Q_0 + t\vec{b}$ ($s, t \in \mathbb{R}$). Fassen Sie g_1 und g_2 als affine Unterräume auf und geben Sie die zugehörigen linearen Unterräume U_1 und U_2 an. Vergleichen Sie für die möglichen Lagebeziehungen von g_1 und g_2 die Dimensionen $\dim(g_1 \cap g_2)$ und $\dim(U_1 \cap U_2)$.
9. Es seien $g: P = P_0 + t\vec{a}$ eine Gerade und $\varepsilon: P = Q_0 + r\vec{b} + s\vec{c}$ eine Ebene in $\mathbb{R}^3$ sowie $U_g = \langle \vec{a} \rangle$ und $U_\varepsilon = \langle \vec{b}, \vec{c} \rangle$ die zugehörigen linearen Unterräume. Überlegen Sie sich für die möglichen Lagebeziehungen von g und ε (siehe Abschnitt 4.2.2), welche Dimensionen die Schnittmenge $U_\varepsilon \cap U_g$, die Summe $U_\varepsilon + U_g$ sowie die affine Hülle $A(g \cup \varepsilon)$ haben. Verifizieren Sie, dass die Sätze 5.27 und 5.28 in allen Fällen zutreffen.
10. Untersuchen Sie alle möglichen Lagebeziehungen einer Geraden g und einer Ebene ε in einem vierdimensionalen affinen Raum.
11. Gegeben ist der affine Unterraum $N = \{P \mid P = P_0 + \vec{u}; \vec{u} \in U\}$ von $\mathbb{R}^4$ mit $P_0 = \begin{pmatrix} 2 \\ 1 \\ 3 \\ 4 \end{pmatrix}$ und $U = \left\langle \begin{pmatrix} 3 \\ 2 \\ -1 \\ 0 \end{pmatrix}; \begin{pmatrix} 4 \\ 5 \\ 7 \\ -3 \end{pmatrix} \right\rangle$. Geben Sie ein lineares Gleichungssystem an, als dessen Lösungsmenge sich N ergibt.

5.6 Euklidische Vektor- und Punkträume

In diesem Abschnitt werden Vektor- und affine Punkträume um *metrische Aspekte*, d. h. um Möglichkeiten des „Messens“ von Längen und Winkeln ergänzt. Dazu wird, wie bereits in dem Abschnitt 3.5, ein Skalarprodukt eingeführt, hier nun jedoch auf verallgemeinerter Grundlage für beliebige Vektorräume.

5.6.1 Positiv definite symmetrische Bilinearformen – verallgemeinerte Skalarprodukte

Definition 5.15

Eine *positiv definite symmetrische Bilinearform* (bzw. ein *verallgemeinertes Skalarprodukt*) auf einem Vektorraum V ist eine Abbildung $B: V \times V \rightarrow \mathbb{R}$ (die also jedem Paar von Vektoren eine reelle Zahl zuordnet) mit folgenden Eigenschaften (für beliebige $\vec{u}, \vec{v}, \vec{w} \in V$, $\lambda \in \mathbb{R}$):

$$\text{B1. a) } B(\vec{u} + \vec{v}, \vec{w}) = B(\vec{u}, \vec{w}) + B(\vec{v}, \vec{w}) \quad (\text{Additivität})$$

$$\text{b) } B(\lambda \vec{u}, \vec{v}) = \lambda B(\vec{u}, \vec{v}) \quad (\text{Homogenität})$$

$$\text{B2. } B(\vec{u}, \vec{v}) = B(\vec{v}, \vec{u}) \quad (\text{Symmetrie})$$

$$\text{B3. } B(\vec{u}, \vec{u}) \geq 0 \text{ sowie } B(\vec{u}, \vec{u}) = 0 \Leftrightarrow \vec{u} = \vec{0} \quad (B \text{ ist positiv definit})$$

Ein Vektorraum mit einer positiv definiten symmetrischen Bilinearform heißt *euklidischer Vektorraum*. Ein affiner Punktraum mit einem zugehörigen euklidischen Vektorraum heißt *euklidischer Punktraum*. $\blacklozenge$

Bemerkungen:

- Additivität und Homogenität werden zusammengefasst als *Linearität* bezeichnet. Aus der Symmetrie folgt, dass B auch im zweiten Argument additiv und homogen ist, d. h. es gilt für beliebige $\vec{u}, \vec{v}, \vec{w} \in V$, $\lambda \in \mathbb{R}$ auch

$$B(\vec{u}, \vec{v} + \vec{w}) = B(\vec{u}, \vec{v}) + B(\vec{u}, \vec{w}) \quad \text{und} \quad B(\vec{u}, \lambda \vec{v}) = \lambda B(\vec{u}, \vec{v}).$$

Damit ist B linear in beiden Argumenten, also *bilinear*.

- Die in der Definition 5.15 enthaltenen Eigenschaften wurden bereits in dem Satz 3.13 (S. 123) als Eigenschaften eines innerhalb eines konkreten Vektorraumes (nämlich $\mathbb{R}^n$) definierten Skalarproduktes formuliert und bewiesen. Durch die Definition 5.15 werden nun für beliebige Vektorräume Skalarprodukte anhand geforderter Eigenschaften axiomatisch definiert.

Beispiel 5.44

Die Vektorräume $\mathbb{R}^2$, $\mathbb{R}^3$ bzw. allgemeiner $\mathbb{R}^n$ sind mit dem in dem Abschnitt 3.5 definierten Skalarprodukt (welches wegen des Satzes 3.13 auf S. 123 eine positiv definite symmetrische Bilinearform ist) euklidische Vektorräume. Jedoch können auch andere positiv definite symmetrische Bilinearformen verwendet werden. So ist z. B. $\mathbb{R}^2$ mit

$$B: \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R}, \quad B\left(\begin{pmatrix} x_u \\ y_u \end{pmatrix}, \begin{pmatrix} x_v \\ y_v \end{pmatrix}\right) = 4x_u x_v - 2x_u y_v - 2y_u x_v + 3y_u y_v$$

ein euklidischer Vektorraum. Es gilt für beliebige $\vec{u} = \begin{pmatrix} x_u \\ y_u \end{pmatrix}$, $\vec{v} = \begin{pmatrix} x_v \\ y_v \end{pmatrix}$, $\vec{w} = \begin{pmatrix} x_w \\ y_w \end{pmatrix}$:

$$\begin{aligned}
B1. \quad B(\vec{u} + \vec{v}, \vec{w}) &= 4(x_u + x_v)x_w - 2(x_u + x_v)y_w - 2(y_u + y_v)x_w + 3(y_u + y_v)y_w \\
&= 4x_u x_w - 2x_u y_w - 2y_u x_w + 3y_u y_w \\
&\quad + 4x_v x_w - 2x_v y_w - 2y_v x_w + 3y_v y_w \\
&= B(\vec{u}, \vec{w}) + B(\vec{v}, \vec{w})
\end{aligned}$$

$$\begin{aligned}
B(\lambda \vec{u}, \vec{v}) &= 4\lambda x_u x_v - 2\lambda x_u y_v - 2\lambda y_u x_v + 3\lambda y_u y_v \\
&= \lambda(4x_u x_v - 2x_u y_v - 2y_u x_v + 3y_u y_v) = \lambda B(\vec{u}, \vec{v})
\end{aligned}$$

$$\begin{aligned}
B2. \quad B(\vec{u}, \vec{v}) &= 4x_u x_v - 2x_u y_v - 2y_u x_v + 3y_u y_v \\
&= 4x_v x_u - 2y_v x_u - 2x_v y_u + 3y_v y_u = B(\vec{v}, \vec{u})
\end{aligned}$$

$$\begin{aligned}
B3. \quad B(\vec{u}, \vec{u}) &= 4x_u x_u - 2x_u y_u - 2y_u x_u + 3y_u y_u = 4x_u^2 - 4x_u y_u + 3y_u^2 \\
&= (2x_u - y_u)^2 + 2y_u^2
\end{aligned}$$

Da weder $(2x_u - y_u)^2$ noch $2y_u^2$ negativ sein kann, gilt $B(\vec{u}, \vec{u}) \geq 0$ für beliebige $\vec{u} \in \mathbb{R}^2$. Damit $B(\vec{u}, \vec{u}) = 0$ ist, müssen $(2x_u - y_u)$ sowie y_u und damit auch x_u Null sein. Somit gilt $B(\vec{u}, \vec{u}) = 0$ nur, wenn $\vec{u}$ der Nullvektor ist.

Durch B ist somit eine positiv definite symmetrische Bilinearform gegeben. ■

Beispiel 5.45

Der Vektorraum $P_2 = \{p \mid p(x) = a_2 x^2 + a_1 x + a_0; a_0, a_1, a_2 \in \mathbb{R}\}$ der Polynome höchstens zweiten Grades (siehe Beispiel 5.5 auf S. 170) ist mit

$$B: P_2 \times P_2 \rightarrow \mathbb{R}, \quad B(p, q) = a_2 b_2 + a_1 b_1 + a_0 b_0$$

für $p(x) = a_2 x^2 + a_1 x + a_0$, $q(x) = b_2 x^2 + b_1 x + b_0$ ein euklidischer Vektorraum. Es lässt sich leicht zeigen, dass die Bedingungen der Definition 5.15 hierfür erfüllt sind. (Man beachte auch die Analogie dieser Bilinearform B zu dem Standardskalarprodukt von $\mathbb{R}^3$). ■

Satz 5.30

Ist V ein euklidischer Vektorraum mit einer positiv definiten symmetrischen Bilinearform B , so gilt für beliebige $\vec{u}, \vec{v}, \vec{w} \in V$, $\lambda, \mu \in \mathbb{R}$:

$$B(\vec{u}, \lambda \vec{v} + \mu \vec{w}) = \lambda B(\vec{u}, \vec{v}) + \mu B(\vec{u}, \vec{w}).$$

Der Beweis dieses Satzes sei der Leserin bzw. dem Leser überlassen.

Für positiv definite symmetrische Bilinearformen gelten somit die für Skalarprodukte bekannten Rechenregeln. Wir bezeichnen positiv definite symmetrische Bilinearformen daher im Folgenden kurz als *Skalarprodukte* und verwenden statt $B(\vec{u}, \vec{v})$ die Symbolik $\vec{u} \cdot \vec{v}$. Dabei sind die unterschiedlichen Bedeutungen des Zeichens „ $\cdot$ “ zu beachten, siehe die Bemerkung zu der Definition 3.9, S. 122.

In Verallgemeinerung des Satzes 5.30 gilt für das Skalarprodukt eines Vektors $\vec{u}$ mit einer Linearkombination $\sum_{i=1}^k \lambda_i \vec{v}_i$ von k Vektoren:

$$B\left(\vec{u}, \sum_{i=1}^k \lambda_i \vec{v}_i\right) = \sum_{i=1}^k \lambda_i B(\vec{u}, \vec{v}_i) \quad \text{bzw.} \quad \vec{u} \cdot \sum_{i=1}^k \lambda_i \vec{v}_i = \sum_{i=1}^k \lambda_i \vec{u} \cdot \vec{v}_i.$$

Sind zwei Vektoren durch $\vec{u} = \sum_{i=1}^k \lambda_i \vec{u}_i$ und $\vec{v} = \sum_{j=1}^l \mu_j \vec{v}_j$ gegeben, so gilt

$$\vec{u} \cdot \vec{v} = \sum_{i=1}^k \lambda_i \vec{u}_i \cdot \sum_{j=1}^l \mu_j \vec{v}_j = \sum_{i=1}^k \lambda_i \cdot \sum_{j=1}^l \mu_j \vec{u}_i \cdot \vec{v}_j = \sum_{i=1}^k \sum_{j=1}^l \lambda_i \mu_j \vec{u}_i \cdot \vec{v}_j.$$

Besonders interessant wird diese Rechenregel, wenn Darstellungen von Vektoren bezüglich einer Basis gegeben sind.

Es sei V ein euklidischer Vektorraum mit einem Skalarprodukt „ $\cdot$ “, und es sei $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Basis von V . Wenn alle Skalarprodukte $\vec{b}_i \cdot \vec{b}_j$ ($i = 1 \dots n, j = 1 \dots n$) der Basisvektoren untereinander gegeben sind, so sind die Skalarprodukte aller Vektoren von V bestimmt.

Sind $\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i$ und $\vec{v} = \sum_{j=1}^n \mu_j \vec{b}_j$ beliebige Vektoren von V , so gilt

$$\vec{u} \cdot \vec{v} = \sum_{i=1}^n \sum_{j=1}^n \lambda_i \mu_j \vec{b}_i \cdot \vec{b}_j. \quad (5.1)$$

Die bereits in dem Abschnitt 3.5 behandelte Komponentendarstellung

$$\vec{u} \cdot \vec{v} = u_1 v_1 + u_2 v_2 + \dots + u_n v_n = \sum_{i=1}^n u_i v_i.$$

des Skalarproduktes zweier Vektoren in $\mathbb{R}^n$ ist ein Spezialfall von (5.1). Die Komponenten von n -Tupeln sind nämlich gleichzeitig ihre Koordinaten bezüglich der Standardbasis $B_0 = \{\vec{e}_1; \vec{e}_2; \dots; \vec{e}_n\}$ von $\mathbb{R}^n$, siehe Beispiel 5.23 auf S. 187. Für die Basisvektoren der Standardbasis gilt $\vec{e}_i \cdot \vec{e}_i = 1$ ($i = 1 \dots n$) sowie $\vec{e}_i \cdot \vec{e}_j = 0$, falls $i \neq j$ ($i = 1 \dots n, j = 1 \dots n$). Die Gleichung (5.1) vereinfacht sich damit (da alle Summanden mit $i \neq j$ verschwinden) zu $\vec{u} \cdot \vec{v} = \sum_{i=1}^n u_i v_i \vec{e}_i \cdot \vec{e}_i = \sum_{i=1}^n u_i v_i$.

Beispiel 5.46

Wir betrachten die Basis $B = \{\vec{b}_1; \vec{b}_2; \vec{b}_3\}$ mit $\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$, $\vec{b}_3 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$ von $\mathbb{R}^3$ (Beispiel 5.23, S. 187). Die Skalarprodukte $\vec{b}_i \cdot \vec{b}_j$ der Basisvektoren sind:

$$\begin{aligned} \vec{b}_1 \cdot \vec{b}_1 &= 1, & \vec{b}_2 \cdot \vec{b}_2 &= 2, & \vec{b}_3 \cdot \vec{b}_3 &= 3, \\ \vec{b}_1 \cdot \vec{b}_2 &= \vec{b}_2 \cdot \vec{b}_1 = 1, & \vec{b}_1 \cdot \vec{b}_3 &= \vec{b}_3 \cdot \vec{b}_1 = 1, & \vec{b}_2 \cdot \vec{b}_3 &= \vec{b}_3 \cdot \vec{b}_2 = 2. \end{aligned}$$

Wir berechnen nun mithilfe der Gleichung (5.1) das Skalarprodukt der Vektoren $\vec{u} = \begin{pmatrix} 4 \\ -5 \\ 6 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 2 \\ 5 \\ 2 \end{pmatrix}$. Der Vektor $\vec{u}$ hat bezüglich der Basis B die Koordinaten $\lambda_1 = 9, \lambda_2 = -11, \lambda_3 = 6$ (Beispiel 5.24, S. 188), die Koordinaten von $\vec{v}$ bezüglich B sind $\mu_1 = -3, \mu_2 = 3, \mu_3 = 2$. Durch Einsetzen in (5.1) ergibt sich:

$$\begin{aligned} \vec{u} \cdot \vec{v} &= \sum_{i=1}^3 \sum_{j=1}^3 \lambda_i \mu_j \vec{b}_i \cdot \vec{b}_j = \lambda_1 \mu_1 \vec{b}_1 \cdot \vec{b}_1 + \lambda_1 \mu_2 \vec{b}_1 \cdot \vec{b}_2 + \lambda_1 \mu_3 \vec{b}_1 \cdot \vec{b}_3 + \lambda_2 \mu_1 \vec{b}_2 \cdot \vec{b}_1 \\ &\quad + \lambda_2 \mu_2 \vec{b}_2 \cdot \vec{b}_2 + \lambda_2 \mu_3 \vec{b}_2 \cdot \vec{b}_3 + \lambda_3 \mu_1 \vec{b}_3 \cdot \vec{b}_1 + \lambda_3 \mu_2 \vec{b}_3 \cdot \vec{b}_2 + \lambda_3 \mu_3 \vec{b}_3 \cdot \vec{b}_3 \\ &= 9 \cdot (-3) \cdot 1 + 9 \cdot 3 \cdot 1 + 9 \cdot 2 \cdot 1 - 11 \cdot (-3) \cdot 1 \\ &\quad - 11 \cdot 3 \cdot 2 - 11 \cdot 2 \cdot 2 + 6 \cdot (-3) \cdot 1 + 6 \cdot 3 \cdot 2 + 6 \cdot 2 \cdot 3 = -5. \end{aligned}$$

Natürlich erhält man dieses Ergebnis viel schneller, wenn man das Skalarprodukt (wie in dem Abschnitt 3.5) direkt aus den Komponenten der Vektoren berechnet: $\vec{u} \cdot \vec{v} = x_u x_v + y_u y_v + z_u z_v = 8 - 25 + 12 = -5$. Jedoch ist dies nur in Vektorräumen (wie $\mathbb{R}^n$) mit einer „Standardbasis“ möglich, während das hier praktizierte Verfahren bezüglich beliebiger Basen angewendet werden kann. ■

5.6.2 Beträge von Vektoren, Orthogonalität und Winkel

Definition 5.16

Es sei $\vec{u}$ ein Vektor eines euklidischen Vektorraumes V mit dem Skalarprodukt „ $\cdot$ “. Als *Betrag* von $\vec{u}$ bezeichnet man die Wurzel aus dem Skalarprodukt dieses Vektors mit sich selbst: $|\vec{u}| = \sqrt{\vec{u} \cdot \vec{u}} = \sqrt{\vec{u}^2}$.

Sind P und Q zwei Punkte eines euklidischen Punktraumes, so bezeichnet man den Betrag des Vektors $\overrightarrow{PQ}$ als *Abstand der Punkte P und Q* : $|PQ| = |\overrightarrow{PQ}|$. ♦

Beträge von Vektoren wurden bereits in dem Abschnitt 3.5.2 auf dieselbe Weise definiert. Der Unterschied besteht darin, dass wir dort von einem für spezielle Vektorräume gegebenen Skalarprodukt ausgegangen sind und hier nun eine Verallgemeinerung auf beliebige euklidische Vektorräume vollzogen haben.

In dem Abschnitt 3.5.2 wurde bereits die *Cauchy-Schwarzsche Ungleichung* bewiesen (siehe Satz 3.14 auf S. 124):

Für beliebige Vektoren $\vec{u}$ und $\vec{v}$ gilt $|\vec{u} \cdot \vec{v}| \leq |\vec{u}| \cdot |\vec{v}|$.

Wir vermerken, dass beim Beweis der Cauchy-Schwarzschen Ungleichung auf S. 125 nur Tatsachen verwendet wurden, die in beliebigen euklidischen Vektorräumen gelten, weshalb wir diese Ungleichung im Folgenden verwenden können.

Satz 5.31

Dreiecksungleichung: Für beliebige Vektoren $\vec{u}, \vec{v}$ eines euklidischen Vektorraumes gilt $|\vec{u} + \vec{v}| \leq |\vec{u}| + |\vec{v}|$.

Beweis: Es gilt $|\vec{u} + \vec{v}|^2 = \left(\sqrt{(\vec{u} + \vec{v}) \cdot (\vec{u} + \vec{v})} \right)^2 = (\vec{u} + \vec{v}) \cdot (\vec{u} + \vec{v}) = \vec{u} \cdot \vec{u} + \vec{v} \cdot \vec{v} + 2 \vec{u} \cdot \vec{v}$. Nach der Cauchy-Schwarzschen Ungleichung ist $|\vec{u} \cdot \vec{v}| \leq |\vec{u}| \cdot |\vec{v}|$ und somit erst recht $\vec{u} \cdot \vec{v} \leq |\vec{u}| \cdot |\vec{v}|$. Da außerdem nach Definition 5.16 gilt $\vec{u} \cdot \vec{u} = |\vec{u}|^2$ und $\vec{v} \cdot \vec{v} = |\vec{v}|^2$, folgt daraus $|\vec{u} + \vec{v}|^2 \leq |\vec{u}|^2 + |\vec{v}|^2 + 2 |\vec{u}| \cdot |\vec{v}|$ und daher die Behauptung $|\vec{u} + \vec{v}| \leq |\vec{u}| + |\vec{v}|$. □

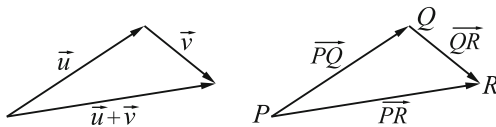


Abb. 5.11:

Dreiecksungleichung in euklidischen Vektor- und Punkträumen

Die *Dreiecksungleichung* lässt sich auch für Punkte formulieren: Für beliebige Punkte P, Q, R eines euklidischen Punktraumes gilt $|PR| \leq |PQ| + |QR|$.

Definition 5.17

Zwei vom Nullvektor verschiedene Vektoren $\vec{u}$ und $\vec{v}$ eines euklidischen Vektorraumes heißen *orthogonal* zueinander, wenn $\vec{u} \cdot \vec{v} = 0$ gilt. ♦

Es mag verwundern, dass die Definition 5.17 dieselbe Aussage zu enthalten scheint wie der Satz 3.16 auf S. 126. Jedoch wurde dort von einer anschaulich-geometrischen Orthogonalität ausgegangen, während der Begriff hier auf der Grundlage der linearen Algebra für beliebige euklidische Vektorräume definiert wird.

Anhand der Definitionen von Abstand und Orthogonalität lassen sich bereits einige Sätze der Elementargeometrie beweisen. Beispielsweise wurde der Satz des Pythagoras in dem Abschnitt 3.5 noch nicht bewiesen, da Beträge von Vektoren und das Skalarprodukt dort mithilfe dieses Satzes eingeführt bzw. geometrisch interpretiert wurden. Ein Beweis des Satzes des Pythagoras auf dieser Grundlage wäre ein „Zirkelschluss“ gewesen. Nun wird dieser Satz aber auf den Grundlagen der linearen Algebra nachgewiesen.

Beispiel 5.47

Satz des Pythagoras: In einem bei C rechtwinkligen Dreieck $\triangle ABC$ mit den A , B und C gegenüberliegenden Seiten a , b bzw. c gilt $a^2 + b^2 = c^2$.

Zum Beweis des Satzes des Pythagoras betrachten wir die Vektoren $\vec{a} = \overrightarrow{CB}$, $\vec{b} = \overrightarrow{AC}$ und $\vec{c} = \overrightarrow{AB}$, siehe Abb. 5.12. Nach Voraussetzung sind $\vec{a}$ und $\vec{b}$ orthogonal, es gilt also $\vec{a} \cdot \vec{b} = 0$. Wegen $\vec{c} = \vec{a} + \vec{b}$ folgt daraus

$$\vec{c}^2 = \vec{c} \cdot \vec{c} = (\vec{a} + \vec{b}) \cdot (\vec{a} + \vec{b}) = \vec{a}^2 + \vec{b}^2 + 2\vec{a} \cdot \vec{b} = \vec{a}^2 + \vec{b}^2.$$

Wegen $c = |AB| = |\overrightarrow{AB}| = |\vec{c}|$ (und analog $a = |\vec{a}|$, $b = |\vec{b}|$) sowie $a^2 = |\vec{a}|^2 = \vec{a} \cdot \vec{a} = \vec{a}^2$ (analog für $\vec{b}^2$ und $\vec{c}^2$) ergibt sich daraus die Behauptung $a^2 + b^2 = c^2$. ■

Definition 5.18

Es seien $\vec{u}$ und $\vec{v}$ zwei von $\vec{o}$ verschiedene Vektoren eines euklidischen Vektorraumes. Als *Winkelmaß der Vektoren $\vec{u}$ und $\vec{v}$* bezeichnet man die Zahl φ mit

$$\cos \varphi = \frac{\vec{u} \cdot \vec{v}}{|\vec{u}| \cdot |\vec{v}|}.$$

Mithilfe dieser Definition des Winkelmaßes – die ohne Rückgriff auf Tatsachen, die aus der Geometrie bekannt sind, angegeben wurde – lässt sich eine Vielzahl bekannter geometrischer Sätze auf der Grundlage der linearen Algebra beweisen.

Beispiel 5.48

Kosinussatz: In einem beliebigen Dreieck (mit Bezeichnungen wie in Abb. 5.13) gilt: $a^2 = b^2 + c^2 + 2bc \cos \alpha$, $c^2 = a^2 + b^2 + 2ab \cos \gamma$, $b^2 = a^2 + c^2 + 2ac \cos \beta$.

Beweis: Es ist $\vec{a} = -\vec{c} + \vec{b}$ und daher

$$\vec{a}^2 = \vec{a} \cdot \vec{a} = \vec{b}^2 + \vec{c}^2 - 2\vec{b} \cdot \vec{c} = \vec{b}^2 + \vec{c}^2 - 2|\vec{b}||\vec{c}| \frac{\vec{b} \cdot \vec{c}}{|\vec{b}||\vec{c}|} = |\vec{b}|^2 + |\vec{c}|^2 - 2|\vec{b}||\vec{c}| \cos \alpha.$$

Die anderen beiden Behauptungen lassen sich völlig analog nachweisen. ■

Es sei vermerkt, dass die in den Abschnitten 3.5.3 und 3.5.4 bewiesenen Sätze der Geometrie auf dieselbe Weise mit den nun zur Verfügung stehenden Tatsachen für beliebige euklidische Punkt- und Vektorräume bewiesen werden könnten. Darüber hinaus lässt sich auf dieser verallgemeinerten Grundlage metrische Geometrie von Geraden und Ebenen in zwei- bzw. dreidimensionalen euklidischen Räumen ebenso betreiben wie in dem Abschnitt 4.3.

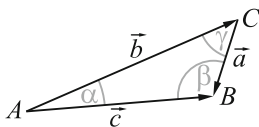
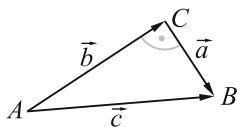


Abb. 5.12: Satz des Pythagoras

Abb. 5.13: Kosinussatz

5.6.3 Orthonormalbasen und kartesische Koordinatensysteme

Definition 5.19

Es sei $(V, \cdot)$ ein beliebiger euklidischer Vektorraum. Eine Basis $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ von V heißt *Orthonormalbasis*, falls die Basisvektoren $\vec{b}_1, \dots, \vec{b}_n$

- Einheitsvektoren sind, d. h. $\vec{b}_1^2 = \vec{b}_2^2 = \dots = \vec{b}_n^2 = 1$,
- paarweise orthogonal zueinander sind, d. h. $\vec{b}_i \cdot \vec{b}_j = 0$ für $i \neq j$, $i, j = 1 \dots n$.

Ein Koordinatensystem $K = \{O; \vec{b}_1; \dots; \vec{b}_n\}$ eines euklidischen Punktraumes A heißt *kartesisches Koordinatensystem*, falls die enthaltene Basis $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Orthonormalbasis des zu A gehörenden euklidischen Vektorraumes ist. ♦

Beispiel 5.49

Das in dem Beispiel 5.35 (S. 209) genutzte Koordinatensystem $K' = \{O; \vec{b}'_1; \vec{b}'_2\}$ mit $\vec{b}'_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$, $\vec{b}'_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ ist ein kartesisches Koordinatensystem von $\mathbb{R}^2$. Um dies zu überprüfen, werden die Beträge (bzw. ihre Quadrate) sowie die Skalarprodukte der Basisvektoren $\vec{b}'_1$ und $\vec{b}'_2$ berechnet:

$$\begin{aligned}\vec{b}'_1{}^2 &= \left(\frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) \cdot \left(\frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix} \right) = \frac{1}{\sqrt{2}} \cdot \frac{1}{\sqrt{2}} \cdot \begin{pmatrix} 1 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \frac{1}{2} (1+1) = 1 \\ \vec{b}'_2{}^2 &= \left(\frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right) \cdot \left(\frac{1}{\sqrt{2}} \begin{pmatrix} -1 \\ 1 \end{pmatrix} \right) = \frac{1}{\sqrt{2}} \cdot \frac{1}{\sqrt{2}} \cdot \begin{pmatrix} -1 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 1 \end{pmatrix} = \frac{1}{2} (1+1) = 1 \\ \vec{b}'_1 \cdot \vec{b}'_2 &= \vec{b}'_2 \cdot \vec{b}'_1 = \frac{1}{\sqrt{2}} \cdot \frac{1}{\sqrt{2}} \cdot \begin{pmatrix} 1 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 1 \end{pmatrix} = \frac{1}{2} (-1+1) = 0.\end{aligned}$$

$B' = \{\vec{b}'_1; \vec{b}'_2\}$ ist somit eine Orthonormalbasis von $\mathbb{R}^2$. ■

Beispiel 5.50

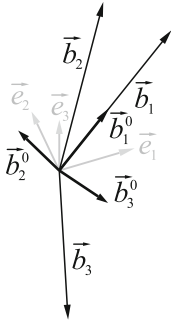
Es ist die Basis $B = \{\vec{b}_1; \vec{b}_2; \vec{b}_3\}$ von $\mathbb{R}^3$ mit $\vec{b}_1 = \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix}$, $\vec{b}_3 = \begin{pmatrix} -1 \\ -3 \\ 1 \end{pmatrix}$

gegeben. Gesucht ist eine Orthonormalbasis $B^0 = \{\vec{b}_1^0; \vec{b}_2^0; \vec{b}_3^0\}$, deren Vektoren so weit wie möglich dieselben Richtungen beschreiben wie $\vec{b}_1, \vec{b}_2, \vec{b}_3$, d. h. $\vec{b}_1^0$ und $\vec{b}_1$ sollen kollinear sein, und $\vec{b}_1^0, \vec{b}_2^0$ sollen dieselbe lineare Hülle haben wie $\vec{b}_1, \vec{b}_2$.

1. Den Basisvektor $\vec{b}_1^0$ erhält man durch Normierung: $\vec{b}_1^0 = \frac{1}{|\vec{b}_1|} \vec{b}_1 = \frac{1}{\sqrt{6}} \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix}$.
2. Der Basisvektor $\vec{b}_2^0$ soll in der linearen Hülle $\langle \vec{b}_1, \vec{b}_2 \rangle$ liegen und zu $\vec{b}_1^0$ bzw. zu $\vec{b}_1$ orthogonal sein, d. h. $\vec{b}_2^0 = \lambda \vec{b}_1 + \mu \vec{b}_2$ ($\lambda, \mu \in \mathbb{R}$) und $\vec{b}_2^0 \cdot \vec{b}_1 = 0$. Zu bestimmen sind Koeffizienten λ und μ , welche diese Bedingungen erfüllen. Dazu setzt man die erste in die zweite Bedingung ein: $(\lambda \vec{b}_1 + \mu \vec{b}_2) \cdot \vec{b}_1 = 0$. Nach Berechnung der Skalarprodukte $\vec{b}_1 \cdot \vec{b}_1 = 6$ und $\vec{b}_2 \cdot \vec{b}_1 = 5$ erhält man daraus die Gleichung $\lambda \vec{b}_1 \cdot \vec{b}_1 + \mu \vec{b}_2 \cdot \vec{b}_1 = 6\lambda + 5\mu = 0$, die beispielsweise für $\lambda_0 = -5$, $\mu_0 = 6$ erfüllt ist. Somit erhält man als Vektor, der die Orthogonalitätsbedingung zu $\vec{b}_1^0$ erfüllt und in der linearen Hülle $\langle \vec{b}_1, \vec{b}_2 \rangle$ liegt:

$$\vec{b}_2' = \lambda_0 \vec{b}_1 + \mu_0 \vec{b}_2 = -5 \vec{b}_1 + 6 \vec{b}_2 = -5 \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} + 6 \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} = \begin{pmatrix} -4 \\ 1 \\ 7 \end{pmatrix}.$$

Dieser Vektor wird normiert, um den Basisvektor $\vec{b}_2^0$ zu erhalten: $\vec{b}_2^0 = \frac{1}{|\vec{b}_2'|} \vec{b}_2' = \frac{1}{\sqrt{66}} \begin{pmatrix} -4 \\ 1 \\ 7 \end{pmatrix}$.

**Abb. 5.14:**

Konstruktion einer Orthonormalbasis aus einer gegebenen Basis von $\mathbb{R}^3$

Die grau dargestellten Vektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3$ sind die Vektoren der Standardbasis von $\mathbb{R}^3$.

3. Der fehlende Basisvektor $\vec{b}_3^0$ lässt sich als Vektorprodukt der Basisvektoren $\vec{b}_1^0, \vec{b}_2^0$ ermitteln; nach Definition 3.11 (S. 133) erfüllt $\vec{b}_3^0 = \vec{b}_1^0 \times \vec{b}_2^0$ die Bedingungen $|\vec{b}_3^0| = 1$, $\vec{b}_3^0 \cdot \vec{b}_1^0 = 0$ und $\vec{b}_3^0 \cdot \vec{b}_2^0 = 0$. Wir setzen daher:

$$\vec{b}_3^0 = \vec{b}_1^0 \times \vec{b}_2^0 = \left(\frac{1}{\sqrt{6}} \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} \right) \times \left(\frac{1}{\sqrt{66}} \begin{pmatrix} -4 \\ 1 \\ 7 \end{pmatrix} \right) = \frac{1}{\sqrt{396}} \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} \times \begin{pmatrix} -4 \\ 1 \\ 7 \end{pmatrix} = \frac{1}{\sqrt{396}} \begin{pmatrix} 6 \\ -18 \\ 6 \end{pmatrix}.$$

Durch eine kleine Vereinfachung erhält man $\vec{b}_3^0 = \frac{1}{\sqrt{11}} \begin{pmatrix} 1 \\ -3 \\ 1 \end{pmatrix}$. ■

Die Schritte 1. und 2. lassen sich immer anwenden, wenn eine Basis eines euklidischen Vektorraumes gegeben und eine Orthonormalbasis gesucht ist. Hingegen ist der Schritt 3. auf dreidimensionale Räume beschränkt, da nur dort ein Vektorprodukt definiert ist, siehe Abschnitt 3.6. Wir bestimmen daher $\vec{b}_3^0$ ein weiteres Mal – nun jedoch nach einem Verfahren, das dem zur Bestimmung von $\vec{b}_2^0$ in Schritt 2. ähnelt und das universell anwendbar ist.

- 3.' Der Basisvektor $\vec{b}_3^0$ soll zu der linearen Hülle $\langle \vec{b}_1^0, \vec{b}_2^0, \vec{b}_3 \rangle$ gehören² und zu den Vektoren $\vec{b}_1^0$ und $\vec{b}_2^0$ orthogonal sein, d. h. $\vec{b}_3^0 = \lambda \vec{b}_1^0 + \mu \vec{b}_2^0 + \nu \vec{b}_3$ ($\lambda, \mu, \nu \in \mathbb{R}$) sowie $\vec{b}_3^0 \cdot \vec{b}_1^0 = 0$ und $\vec{b}_3^0 \cdot \vec{b}_2^0 = 0$. Zu bestimmen sind Koeffizienten λ, μ, ν , welche diese Bedingungen erfüllen. Dazu setzt man $\vec{b}_3^0 = \lambda \vec{b}_1^0 + \mu \vec{b}_2^0 + \nu \vec{b}_3$ in die beiden Orthogonalitätsbedingungen ein:

$$0 = (\lambda \vec{b}_1^0 + \mu \vec{b}_2^0 + \nu \vec{b}_3) \cdot \vec{b}_1^0 = \lambda \vec{b}_1^0 \cdot \vec{b}_1^0 + \mu \vec{b}_2^0 \cdot \vec{b}_1^0 + \nu \vec{b}_3 \cdot \vec{b}_1^0 = \lambda - \frac{4}{\sqrt{6}} \nu$$

$$0 = (\lambda \vec{b}_1^0 + \mu \vec{b}_2^0 + \nu \vec{b}_3) \cdot \vec{b}_2^0 = \lambda \vec{b}_1^0 \cdot \vec{b}_2^0 + \mu \vec{b}_2^0 \cdot \vec{b}_2^0 + \nu \vec{b}_3 \cdot \vec{b}_2^0 = \mu + \frac{8}{\sqrt{66}} \nu.$$

Damit haben wir ein LGS, das wir nach λ, μ, ν lösen. Wir erhalten eine einparametrische Lösungsmenge: $\lambda = t$, $\mu = -\frac{2}{\sqrt{11}} t$, $\nu = \frac{\sqrt{6}}{4} t$. Setzen wir $t = 1$, also $\lambda = 1$, $\mu = -\frac{2}{\sqrt{11}}$, $\nu = \frac{\sqrt{6}}{4}$, so erhalten wir einen Vektor $\vec{b}_3'$, der die Orthogonalitätsbedingungen erfüllt und in der linearen Hülle $\langle \vec{b}_1^0, \vec{b}_2^0, \vec{b}_3 \rangle$ liegt:

$$\vec{b}_3' = \frac{1}{\sqrt{6}} \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} - \frac{2}{\sqrt{11}\sqrt{66}} \begin{pmatrix} -4 \\ 1 \\ 7 \end{pmatrix} + \frac{\sqrt{6}}{4} \begin{pmatrix} -1 \\ -3 \\ 1 \end{pmatrix} = \frac{27}{22\sqrt{6}} \begin{pmatrix} 1 \\ -3 \\ 1 \end{pmatrix}.$$

Durch Normieren des Vektors $\vec{b}_3'$ erhalten wir $\vec{b}_3^0 = \frac{1}{|\vec{b}_3'|} \vec{b}_3' = \frac{1}{\sqrt{11}} \begin{pmatrix} 1 \\ -3 \\ 1 \end{pmatrix}$.

²Die Forderung $\vec{b}_3^0 \in \langle \vec{b}_1^0, \vec{b}_2^0, \vec{b}_3 \rangle$ ist für das betrachtete Beispiel trivial, da diese lineare Hülle der gesamte Vektorraum $\mathbb{R}^3$ ist. Der daraus resultierende Ansatz $\vec{b}_3^0 = \lambda \vec{b}_1^0 + \mu \vec{b}_2^0 + \nu \vec{b}_3$ ist trotzdem sinnvoll und insbesondere auf beliebige, auch höherdimensionale, euklidische Vektorräume übertragbar.

Das in dem Beispiel 5.50 mit den Schritten 1., 2., 3.' praktizierte Verfahren der Konstruktion einer Orthonormalbasis aus einer gegebenen Basis lässt sich für beliebige Basen in euklidischen Vektorräumen praktizieren; in n -dimensionalen Vektorräumen sind dazu n Schritte erforderlich. Das Verfahren ist auch für Basen linearer Unterräume euklidischer Vektorräume anwendbar, denn lineare Unterräume sind selbst Vektorräume. Somit kann zum Beispiel eine Orthonormalbasis des Vektorraumes der magischen 3×3 -Quadrate (der ein Unterraum von $\mathbb{R}^9$ ist, siehe Beispiel 5.12 auf S. 174) konstruiert werden, siehe Aufgabe 6.

Das praktizierte Verfahren zur Konstruktion einer Orthonormalbasis heißt *Gram-Schmidtsches Orthonormierungsverfahren* (nach Jorgen Pedersen Gram und Erhard Schmidt). Da es in beliebigen euklidischen Vektorräumen anwendbar ist, gilt der folgende Satz.

Satz 5.32

Jeder euklidische Vektorraum endlicher Dimension besitzt eine Orthonormalbasis, jeder zugehörige euklidische Punktraum ein kartesisches Koordinatensystem.

Die Aussagen des Satzes 5.32 sind Existenzaussagen. Wie wir in den behandelten Beispielen gesehen haben, besitzen euklidische Vektorräume mehrere Orthonormalbasen und Punkträume mehrere (i. Allg. sogar unendlich viele) verschiedene kartesische Koordinatensysteme.

5.6.4 Aufgaben zu Abschnitt 5.6

1. Weisen Sie nach, dass $\mathbb{R}^2$ mit

$$B: \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R}, \quad B\left(\begin{pmatrix} x_u \\ y_u \end{pmatrix}, \begin{pmatrix} x_v \\ y_v \end{pmatrix}\right) = x_u x_v + x_u y_v + y_u x_v + 2 y_u y_v$$

ein euklidischer Vektorraum ist.

2. Beweisen Sie den Satz 5.30 auf S. 220.

3. Gegeben sind die Vektoren $\vec{u}$ mit den Koordinaten $\lambda_1 = 1, \lambda_2 = \frac{4}{3}, \lambda_3 = -\frac{2}{3}$ und $\vec{v}$ mit den Koordinaten $\mu_1 = -2, \mu_2 = -\frac{7}{3}, \mu_3 = \frac{2}{3}$ bezüglich der Basis

$$B = \{\vec{b}_1; \vec{b}_2; \vec{b}_3\} \text{ mit } \vec{b}_1 = \begin{pmatrix} 3 \\ 2 \\ 12 \end{pmatrix}, \vec{b}_2 = \begin{pmatrix} -8 \\ 7 \\ 13 \end{pmatrix}, \vec{b}_3 = \begin{pmatrix} 2 \\ 5 \\ -4 \end{pmatrix}.$$

Berechnen Sie die Skalarprodukte der Basisvektoren untereinander und darauf basierend das Skalarprodukt der Vektoren $\vec{u}$ und $\vec{v}$ (vgl. Beispiel 5.46 auf S. 221).

4. Beweisen Sie die folgende Ergänzung der Cauchy-Schwarzschen Ungleichung: Sind zwei Vektoren $\vec{u}$ und $\vec{v}$ eines euklidischen Vektorraumes linear abhängig, so gilt $|\vec{u} \cdot \vec{v}| = |\vec{u}| \cdot |\vec{v}|$.
5. Beweisen Sie, dass für beliebige Vektoren $\vec{u}$ und $\vec{v}$ eines euklidischen Vektorraumes und beliebige $\lambda, \mu \in \mathbb{R}$ gilt: $|\lambda \vec{u} + \mu \vec{v}| \leq |\lambda| |\vec{u}| + |\mu| |\vec{v}|$.
6. Konstruieren Sie eine Orthonormalbasis des Vektorraumes der magischen Quadrate der Kantenlänge 3. Gehen Sie dabei von der in dem Beispiel 5.29 auf S. 197 ermittelten Basis aus. Orientieren Sie sich an den Schritten 1., 2., 3.' des Beispiels 5.50.

6 Matrizen

Übersicht

6.1	Begriffsbestimmung, Ränge von Matrizen	228
6.2	Matrizenmultiplikation und -inversion	235
6.3	Ein Ausblick auf Determinanten	251
6.4	Matrizenrechnung mithilfe des Computers	256

Matrizen traten bereits in dem Kapitel 1 und in einigen Beispielen des Kapitels 5 auf. Sie werden nun systematisch behandelt, d. h. es wird *Matrizenrechnung* betrieben, und es werden *strukturelle Betrachtungen* anhand von Matrizen angestellt. Dabei wird sich zeigen, dass der Matrizenkalkül elegante Zusammenfassungen zentraler Sachverhalte der linearen Algebra (wie der linearen Abhängigkeit/Unabhängigkeit von Vektormengen) ermöglicht. Eine besonders hohe Bedeutung haben Matrizen für die Beschreibung der in dem Kapitel 7 behandelten linearen Abbildungen; auch für diese sind Matrizen sowohl Rechenhilfsmittel als auch „Fingerabdrücke“, welche wichtige Eigenschaften beschreiben.

Das vorliegende Kapitel verbindet somit in besonders enger Weise rechnerische Aspekte und inhaltlich-strukturelle Überlegungen. Ergänzend werden einige Anwendungsbeispiele zur Materialverflechtung sowie zur Beschreibung von Populationen behandelt. Dabei wird nicht der Anspruch verfolgt, tatsächlich „realistische“ Anwendungen zu behandeln; diese sind i. Allg. um ein Vielfaches komplexer, als dies hier dargestellt werden kann. Es soll aber mit den behandelten Beispielen zumindest ein kleiner Einblick in nichtgeometrische Anwendungen der linearen Algebra gewährt werden, und vor allem sollen die Beispiele das Verständnis zentraler mathematischer Inhalte erleichtern sowie Rechenregeln motivieren und verständlicher machen.

Bei der Matrizenrechnung ist oftmals viel Zeit für aufwändige Berechnungen erforderlich. Dem Einsatz des Computers kommt daher besondere Bedeutung zu. Eine Übersicht über die Nutzung des Computeralgebrasystems Maxima für alle in diesem Kapitel vorgenommenen Berechnungen enthält der Abschnitt 6.4. Der Leserin bzw. dem Leser wird empfohlen, schon beim Studium der vorherigen Abschnitte und dem Lösen der dort gestellten Übungsaufgaben einen Blick in diesen Abschnitt zu werfen und den Computer zu nutzen.

6.1 Begriffsbestimmung, Ränge von Matrizen

6.1.1 Zeilen- und Spaltenvektoren, transponierte Matrizen

Matrizen traten bereits an mehreren Stellen dieses Buches auf:

- In dem Abschnitt 1.3 (S. 29ff.) wurden einfache und erweiterte Koeffizientenmatrizen linearer Gleichungssysteme betrachtet und Lösbarkeitskriterien sowie Aussagen über Lösungsmengen von LGS anhand der Ränge ihrer Koeffizientenmatrizen formuliert.
- In dem Beispiel 5.6 auf S. 171 stellte sich heraus, dass die Menge aller Matrizen reeller Zahlen mit n Spalten und m Zeilen sowie einer dort definierten Matrizenaddition und eine Multiplikation von Matrizen mit reellen Zahlen ein Vektorraum ist. Die Dimension dieses Vektorraumes ist, wie anhand späterer Beispiele deutlich wurde, $n \cdot m$.
- Als spezielle 3×3 -Matrizen wurden an mehreren Stellen magische Quadrate betrachtet (siehe u. a. Beispiel 5.12 auf S. 174).

Wir werden im Folgenden den Begriff der Matrix fundieren, insbesondere Ränge von Matrizen genauer behandeln und dabei Bezüge zur linearen Unabhängigkeit von Vektoren sowie zu Basen und Dimensionen linearer Unterräume herstellen.

Definition 6.1

Ein Schema der Form

$$\mathbf{A} = (a_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$$

mit $a_{ij} \in \mathbb{R}$, $1 \leq i \leq m$, $1 \leq j \leq n$ heißt *Matrix reeller Zahlen* mit m Zeilen und n Spalten bzw. kurz $m \times n$ -Matrix; i ist der *Zeilenindex* und j der *Spaltenindex*. Die Zahlen a_{ij} werden als *Koeffizienten* der Matrix $\mathbf{A}$ bezeichnet.

Die Vektoren $\vec{a}_{S1} = \begin{pmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{pmatrix}$, $\vec{a}_{S2} = \begin{pmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{m2} \end{pmatrix}$, $\dots$, $\vec{a}_{Sn} = \begin{pmatrix} a_{1n} \\ a_{2n} \\ \vdots \\ a_{mn} \end{pmatrix}$ heißen *Spaltenvektoren*, die Vektoren $\vec{a}_{Z1} = (a_{11}; a_{12}; \dots; a_{1n})$, $\vec{a}_{Z2} = (a_{21}; a_{22}; \dots; a_{2n})$, $\dots$, $\vec{a}_{Zm} = (a_{m1}; a_{m2}; \dots; a_{mn})$ heißen *Zeilenvektoren* der Matrix $\mathbf{A}$.

Die Menge aller reellen $m \times n$ -Matrizen wird mit $\mathbb{R}^{m \times n}$ bezeichnet. $\blacklozenge$

Bemerkung: $\mathbb{R}^{m \times n}$ ist mit den in dem Beispiel 5.6 auf S. 171 definierten Verknüpfungen „+“ und „ $\cdot$ “ ein Vektorraum.

Definition 6.2

Es sei $\mathbf{A}$ eine Matrix (mit den Bezeichnungen der Koeffizienten wie in der Definition 6.1). Die Matrix

$$\mathbf{A}^T = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$$

heißt *transponierte Matrix* zu $\mathbf{A}$ bzw. *Transponierte* der Matrix $\mathbf{A}$. $\blacklozenge$

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1j} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2j} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mj} & \dots & a_{mn} \end{pmatrix} \longrightarrow \mathbf{A}^T = \begin{pmatrix} a_{11} & a_{21} & \dots & a_{m1} \\ a_{12} & a_{22} & \dots & a_{m2} \\ \vdots & \vdots & & \vdots \\ a_{1j} & a_{2j} & \dots & a_{mj} \\ \vdots & \vdots & & \vdots \\ a_{1n} & a_{2n} & \dots & a_{mn} \end{pmatrix}$$

Abb. 6.1: Transponieren einer Matrix

Die Transponierte einer Matrix entsteht dadurch, dass aus den Spalten der Ausgangsmatrix die Zeilen der transponierten Matrix werden (und umgekehrt), siehe Abb. 6.1. Für beliebige Koeffizienten a_{ij}^T der transponierten Matrix und a_{ji} der Ausgangsmatrix gilt $a_{ij}^T = a_{ji}$.

Es lassen sich auch Spalten- und Zeilenvektoren transponieren, z. B.

$$\vec{a}_{Sj}^T = \begin{pmatrix} a_{1j} \\ a_{2j} \\ \vdots \\ a_{mj} \end{pmatrix}^T = (a_{1j}; a_{2j}; \dots; a_{mj}) \quad \text{und} \quad \vec{a}_{Zi}^T = (a_{i1}; a_{i2}; \dots; a_{in})^T = \begin{pmatrix} a_{i1} \\ a_{i2} \\ \vdots \\ a_{in} \end{pmatrix}.$$

Wir werden diese Schreibweise mitunter nutzen, um Spaltenvektoren als Zeilen zu schreiben bzw. umgekehrt.

Definition 6.3

Matrizen mit gleicher Zeilen- und Spaltenzahl ($n=m$) bezeichnet man als *quadratische Matrizen*. Ist die Transponierte einer quadratischen Matrix $\mathbf{A} \in \mathbb{R}^{n \times n}$ gleich der Matrix selbst (gilt also $a_{ji} = a_{ij}$ für alle $i, j = 1, \dots, n$), so heißt $\mathbf{A}$ *symmetrische Matrix*; die *Hauptdiagonale* einer symmetrischen Matrix (welche aus den Koeffizienten $a_{11}, a_{22}, \dots, a_{nn}$ besteht) ist ihre Symmetrieachse:

$$\mathbf{A} = \begin{pmatrix} \mathbf{a}_{11} & a_{12} & \dots & a_{1n} \\ a_{12} & \mathbf{a}_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{1n} & a_{2n} & \dots & \mathbf{a}_{nn} \end{pmatrix}.$$

Sind in einer quadratischen Matrix $\mathbf{A}$ alle Koeffizienten a_{ij} mit $i \neq j$ (d. h. alle Koeffizienten, die nicht der Hauptdiagonalen angehören) gleich Null, so heißt $\mathbf{A}$ *Diagonalmatrix*:

$$\mathbf{A} = \begin{pmatrix} \mathbf{a}_{11} & 0 & \dots & 0 \\ 0 & \mathbf{a}_{22} & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \mathbf{a}_{nn} \end{pmatrix}. \quad \blacklozenge$$

6.1.2 Zeilenrang und Spaltenrang einer Matrix

In dem Abschnitt 1.3.2 wurde der Rang einer Matrix als Anzahl der nach Überführung in die Zeilenstufenform verbliebenen Zeilen (die von reinen Nullzeilen verschieden sind) bzw. – gleichbedeutend damit – als Anzahl der „nicht überflüssigen“ Gleichungen eines LGS eingeführt. Diese Bedeutung des Ranges einer Matrix wird nun mit anderen Bedeutungen (Anzahlen von linear unabhängigen Zeilen- und Spaltenvektoren, Dimensionen der dadurch erzeugten linearen Unterräume) verknüpft, womit „Rang“ dann zu einem der zentralen Begriffe der Linearen Algebra wird.

Definition 6.4

Als Spaltenrang einer Matrix $\mathbf{A}$ bezeichnet man die maximale Anzahl linear unabhängiger Spaltenvektoren, als Zeilenrang die maximale Anzahl linear unabhängiger Zeilenvektoren von $\mathbf{A}$. ♦

Folgerung: Der Spaltenrang einer Matrix $\mathbf{A} = (a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}}$ ist nach der Definition 6.4 gleich der Dimension des von den Spaltenvektoren $\vec{a}_{S1}, \vec{a}_{S2}, \dots, \vec{a}_{Sn}$ erzeugten linearen Unterraumes $\langle \vec{a}_{S1}, \vec{a}_{S2}, \dots, \vec{a}_{Sn} \rangle$ von $\mathbb{R}^m$; ihr Zeilenrang ist gleich der Dimension des von den Zeilenvektoren $\vec{a}_{Z1}, \vec{a}_{Z2}, \dots, \vec{a}_{Zm}$ erzeugten Unterraumes $\langle \vec{a}_{Z1}, \vec{a}_{Z2}, \dots, \vec{a}_{Zm} \rangle$ von $\mathbb{R}^n$, siehe die Abschnitte 5.3 und 5.4.

Beispiel 6.1

Es werden der Zeilenrang und der Spaltenrang der folgenden Matrix bestimmt:

$$\mathbf{A} = \begin{pmatrix} 1 & 3 & 5 & 2 & -4 \\ 0 & 2 & 3 & 1 & -3 \\ 1 & -3 & -4 & -1 & 5 \\ 4 & 6 & 11 & 5 & 2 \end{pmatrix}.$$

Um den *Spaltenrang* der Matrix $\mathbf{A}$ zu bestimmen, muss eine maximale Menge linear unabhängiger Vektoren unter den Spaltenvektoren

$$\vec{a}_{S1} = \begin{pmatrix} 1 \\ 0 \\ 1 \\ 4 \end{pmatrix}, \vec{a}_{S2} = \begin{pmatrix} 3 \\ 2 \\ -3 \\ 6 \end{pmatrix}, \vec{a}_{S3} = \begin{pmatrix} 5 \\ 3 \\ -4 \\ 11 \end{pmatrix}, \vec{a}_{S4} = \begin{pmatrix} 2 \\ 1 \\ -1 \\ 5 \end{pmatrix}, \vec{a}_{S5} = \begin{pmatrix} -4 \\ -3 \\ 5 \\ 2 \end{pmatrix}$$

ermittelt werden. Da es sich um Vektoren in $\mathbb{R}^4$ handelt, kann es höchstens vier linear unabhängige Vektoren unter $\vec{a}_{S1}, \dots, \vec{a}_{S5}$ geben. Um festzustellen, ob darunter vier linear unabhängige Vektoren sind, kann man zunächst prüfen, ob $\vec{a}_{S1}, \dots, \vec{a}_{S4}$ linear unabhängig sind, indem man (analog zu dem Beispiel 5.19, S. 182) die Vektorgleichung $\lambda_1 \vec{a}_{S1} + \lambda_2 \vec{a}_{S2} + \lambda_3 \vec{a}_{S3} + \lambda_4 \vec{a}_{S4} = \vec{0}$ löst. Es ergibt sich eine zweiparametrische Lösungsmenge: $\lambda_1 = -\frac{s+t}{2}$, $\lambda_2 = -\frac{3s+t}{2}$, $\lambda_3 = s$, $\lambda_4 = t$. Somit gilt für beliebige $r, s \in \mathbb{R}$:

$$-\frac{s+t}{2} \vec{a}_{S1} - \frac{3s+t}{2} \vec{a}_{S2} + s \vec{a}_{S3} + t \vec{a}_{S4} = \vec{0}.$$

Die Vektoren $\vec{a}_{S1}, \dots, \vec{a}_{S4}$ sind somit linear abhängig; mehr noch: es existieren nicht einmal drei linear unabhängige Vektoren unter diesen vier Vektoren, denn entfernt man $\vec{a}_{S3}$ oder $\vec{a}_{S4}$, so hat die verbleibende Gleichung immer noch nicht triviale Lösungen. Somit kann es unter $\vec{a}_{S1}, \dots, \vec{a}_{S5}$ höchstens drei linear unabhängige Vektoren geben, unter denen dann $\vec{a}_{S5}$ sein müsste. Wir prüfen zunächst, ob $\vec{a}_{S1}, \vec{a}_{S2}, \vec{a}_{S5}$ linear unabhängig sind. Lösen der Vektorgleichung $\lambda_1 \vec{a}_{S1} + \lambda_2 \vec{a}_{S2} + \lambda_5 \vec{a}_{S5} = \vec{0}$ ergibt die eindeutige Lösung $\lambda_1 = \lambda_2 = \lambda_5 = 0$. Somit sind die Vektoren $\vec{a}_{S1}, \vec{a}_{S2}, \vec{a}_{S5}$ linear unabhängig, und eine größere Zahl linear unabhängiger Vektoren existiert unter den Spaltenvektoren $\vec{a}_{S1}, \dots, \vec{a}_{S5}$ der Matrix $\mathbf{A}$ nicht. Der *Spaltenrang* von $\mathbf{A}$ ist somit 3; eine *Basis des von den Spaltenvektoren erzeugten linearen Unterraumes* (von $\mathbb{R}^4$) ist $\{\vec{a}_{S1}; \vec{a}_{S2}; \vec{a}_{S5}\}$.

Es wird nun der *Zeilenrang* der Matrix $\mathbf{A}$ bestimmt. Dazu muss eine maximale Menge linear unabhängiger Vektoren unter den Zeilenvektoren ermittelt werden. Wir stellen diese im Sinne einer besseren Übersichtlichkeit in transponierter Schreibweise als Spaltenvektoren dar:

$$\vec{a}_{Z1}^T = \begin{pmatrix} 1 \\ 3 \\ 5 \\ 2 \\ -4 \end{pmatrix}, \quad \vec{a}_{Z2}^T = \begin{pmatrix} 0 \\ 2 \\ 3 \\ 1 \\ -3 \end{pmatrix}, \quad \vec{a}_{Z3}^T = \begin{pmatrix} 1 \\ -3 \\ -4 \\ -1 \\ 5 \end{pmatrix}, \quad \vec{a}_{Z4}^T = \begin{pmatrix} 4 \\ 6 \\ 11 \\ 5 \\ 2 \end{pmatrix}.$$

Durch Lösen der Vektorgleichung $\lambda_1 \vec{a}_{Z1}^T + \lambda_2 \vec{a}_{Z2}^T + \lambda_3 \vec{a}_{Z3}^T + \lambda_4 \vec{a}_{Z4}^T = \vec{0}$ erhält man eine einparametrische Lösungsmenge: $\lambda_1 = -t$, $\lambda_2 = 3t$, $\lambda_3 = t$, $\lambda_4 = 0$ ($t \in \mathbb{R}$). Somit besitzt die obige Vektorgleichung nicht triviale Lösungen; die Vektoren $\vec{a}_{Z1}^T, \dots, \vec{a}_{Z4}^T$ sind linear abhängig. Dies gilt natürlich auch für die „nicht transponierten“ Zeilenvektoren $\vec{a}_{Z1}, \dots, \vec{a}_{Z4}$, denn die Operationen, die zur Überprüfung der linearen Abhängigkeit angewendet werden, sind hierfür dieselben; es wird lediglich eine andere Schreibweise genutzt.

Da sich bei der Lösung der Gleichung $\lambda_1 \vec{a}_{Z1}^T + \lambda_2 \vec{a}_{Z2}^T + \lambda_3 \vec{a}_{Z3}^T + \lambda_4 \vec{a}_{Z4}^T = \vec{0}$ zwingend $\lambda_4 = 0$ ergibt, ist der Vektor $\vec{a}_{Z4}^T$ nicht als Linearkombination der restlichen Vektoren darstellbar. Da zugleich offensichtlich $\vec{a}_{Z1}^T$ und $\vec{a}_{Z2}^T$ linear unabhängig sind, sollte die Vektormenge $\{\vec{a}_{Z1}^T; \vec{a}_{Z2}^T; \vec{a}_{Z4}^T\}$ linear unabhängig sein. Dies bestätigt man durch Lösen der Vektorgleichung $\lambda_1 \vec{a}_{Z1}^T + \lambda_2 \vec{a}_{Z2}^T + \lambda_4 \vec{a}_{Z4}^T = \vec{0}$; hierfür erhält man tatsächlich nur die triviale Lösung $\lambda_1 = \lambda_2 = \lambda_4 = 0$. Die Zeilenvektoren $\vec{a}_{Z1}, \vec{a}_{Z2}, \vec{a}_{Z4}$ sind also linear unabhängig und bilden, da $\vec{a}_{Z1}, \dots, \vec{a}_{Z4}$ linear abhängig sind, eine maximale Menge linear unabhängiger Zeilenvektoren der Matrix **A**. Der Zeilenrang von **A** ist somit 3; eine *Basis des von den Zeilenvektoren erzeugten linearen Unterraumes* (von $\mathbb{R}^5$) ist $\{\vec{a}_{Z1}; \vec{a}_{Z2}; \vec{a}_{Z4}\}$. ■

Anmerkungen und weiterführende Überlegungen zu dem Beispiel 6.1

- Der *Spaltenrang* der betrachteten Matrix ist *gleich* ihrem *Zeilenrang*, obwohl es sich um die Ränge völlig unterschiedlicher Vektormengen handelt. Allerdings könnte dies nur zufälligerweise bei dem behandelten Beispiel der Fall sein. Wir werden diesem Zusammenhang weiter nachgehen und dabei feststellen, dass er in jedem Falle zutrifft.
- Die Bestimmung des Zeilen- und des Spaltenranges nach dem in dem Beispiel 6.1 praktizierten Verfahren kann recht langwierig sein, da jeweils unter einer mitunter größeren Zahl von Zeilen- und Spaltenvektoren eine maximale Anzahl linear unabhängiger Vektoren „gesucht“ werden muss. Die Nutzung eines Computeralgebrasystems erweist sich dabei als hilfreich (siehe dazu den Abschnitt 6.4 sowie eine Datei für das CAS Maxima, die auf der Internetseite zu diesem Buch zur Verfügung steht), der Aufwand bleibt dennoch hoch.
- Wir bestimmen im Folgenden den Rang der in dem Beispiel 6.1 betrachteten Matrix **A** nach der in dem Abschnitt 1.3.2 gegebenen Definition (siehe S. 33) als Anzahl der nach Überführung in die Zeilenstufenform verbleibenden Zeilen, die von reinen Nullzeilen verschieden sind. Dazu formen wir **A** mithilfe des Gauß-Jordan-Algorithmus (vgl. Abschnitt 1.2.2) um.

$$\left(\begin{array}{ccccc|c} 1 & 3 & 5 & 2 & -4 & \\ 0 & 2 & 3 & 1 & -3 & \\ 1 & -3 & -4 & -1 & 5 & \\ 4 & 6 & 11 & 5 & 2 & \end{array} \right) \begin{array}{l} | \cdot (-1) \\ | \cdot (-4) \\ | \cdot 1 \\ | \cdot 1 \end{array} \longrightarrow \left(\begin{array}{ccccc|c} 1 & 3 & 5 & 2 & -4 & \\ 0 & 2 & 3 & 1 & -3 & \\ 0 & -6 & -9 & -3 & 9 & \\ 0 & -6 & -9 & -3 & 18 & \end{array} \right) \begin{array}{l} \\ | \cdot 3 \\ | \cdot 1 \\ | \cdot 1 \end{array}$$

$$\begin{aligned}
&\rightarrow \begin{pmatrix} 1 & 3 & 5 & 2 & -4 \\ 0 & 2 & 3 & 1 & -3 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 9 \end{pmatrix} \xrightarrow{\text{Zeilen vertauschen}} \begin{pmatrix} 1 & 3 & 5 & 2 & -4 \\ 0 & 2 & 3 & 1 & -3 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \begin{array}{l} | \cdot 9 \\ | \cdot 3 \\ | \cdot 1 \\ | \cdot 4 \end{array} \\
&\rightarrow \begin{pmatrix} 9 & 27 & 45 & 18 & 0 \\ 0 & 6 & 9 & 3 & 0 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \begin{array}{l} | \cdot 1 \\ | \cdot (-6) \end{array} \rightarrow \begin{pmatrix} 9 & -9 & -9 & 0 & 0 \\ 0 & 6 & 9 & 3 & 0 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}
\end{aligned}$$

Offensichtlich ist der Rang dieser Matrix 3, also gleich ihrem Zeilen- und Spaltenrang. Wir nehmen weitere Vereinfachungen vor, aus denen sich Überlegungen zur Gleichheit aller bestimmten „Ränge“ ergeben werden.

Alle bisher vorgenommenen Matrixumformungen (Vertauschen von Zeilen, Multiplikation von Zeilen mit Zahlen, Addition von Zeilen) waren *Zeilenumformungen*, die zunächst für das Lösen linearer Gleichungssysteme eingeführt wurden. Bei der Umformung einer Matrix sind jedoch *Spaltenumformungen* ebenso sinnvoll (diese sind zugleich Zeilenumformungen der transponierten Matrix). Mit Spaltenumformungen vereinfachen wir die Matrix weiter.

$$\begin{pmatrix} 9 & -9 & -9 & 0 & 0 \\ 0 & 6 & 9 & 3 & 0 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \xrightarrow{\begin{array}{l} \cdot 1 \cdot (-3) \\ \cdot 1 \quad \cdot (-2) \end{array}} \begin{pmatrix} 9 & -9 & -9 & 0 & 0 \\ 0 & 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix} \xrightarrow{\begin{array}{l} \cdot 1 \quad \cdot 1 \\ \cdot 1 \quad \cdot 1 \end{array}} \begin{pmatrix} 9 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 9 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Durch Vertauschen von Spalten erhalten wir schließlich die Matrix

$$\begin{pmatrix} 9 & 0 & 0 & 0 & 0 \\ 0 & 3 & 0 & 0 & 0 \\ 0 & 0 & 9 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{pmatrix}.$$

Im Ergebnis aller Umformungen ist eine Matrix entstanden, die aus einer Diagonalmatrix und ansonsten nur aus Nullzeilen und -spalten besteht. Der Zeilenrang einer dergestalt umgeformten Matrix ist offensichtlich gleich dem Spaltenrang und gleich dem in dem Abschnitt 1.3.2 eingeführten Rang. Es lässt sich sogar jede Matrix durch Zeilen- und Spaltenumformungen in eine derartige Form bringen. Um allgemein konstatieren zu können, dass der Spaltenrang einer beliebigen Matrix gleich ihrem Zeilenrang ist, ist noch zu zeigen, dass beide Ränge durch die vorgenommenen Zeilen- und Spaltenumformungen nicht verändert werden.

Satz 6.1

Der Zeilenrang und der Spaltenrang einer Matrix werden durch die folgenden Zeilen- und Spaltenumformungen nicht verändert:

- Z1. Vertauschen zweier Zeilen,
- Z2. Multiplikation einer Zeile mit einer reellen Zahl (außer Null),
- Z3. Addition einer Zeile zu einer anderen Zeile,
- S1. Vertauschen zweier Spalten,
- S2. Multiplikation einer Spalte mit einer reellen Zahl (außer Null),
- S3. Addition einer Spalte zu einer anderen Spalte.

Beweis: Der Satz enthält insgesamt 12 Aussagen. Durch Übertragung auf transponierte Matrizen ergibt sich aber, dass sich Spaltenumformungen auf den Spaltenrang ebenso auswirken wie Zeilenumformungen auf den Zeilenrang sowie dass Zeilenumformungen auf den Spaltenrang dieselbe Wirkung haben wie entsprechende Spaltenumformungen auf den Zeilenrang. Es genügt somit, zu zeigen, dass die Umformungen Z1-S3 *Spaltenränge* beliebiger Matrizen nicht verändern.

S1. Durch Spaltenvertauschungen ändert sich der Spaltenrang einer Matrix nicht, da bei der Bestimmung der linearen Unabhängigkeit einer Menge von Vektoren, deren Reihenfolge irrelevant ist, vgl. die Definition 5.6 auf S. 182.

Z1. Der Spaltenrang einer Matrix bleibt bei Zeilenvertauschungen unverändert, wenn sich dabei keine Veränderungen hinsichtlich der linearen Abhängigkeit bzw. Unabhängigkeit an Teilmengen von Spaltenvektoren der betrachteten Matrix ergeben. Wir betrachten dazu eine beliebige Teilmenge von k Spaltenvektoren einer $m \times n$ -Matrix. Diese können wir so umbenennen bzw. mit anderen Spaltenvektoren vertauschen, dass sie die ersten k Spaltenvektoren der gegebenen Matrix bilden und daher mit $\vec{a}_{S1}, \vec{a}_{S2}, \dots, \vec{a}_{Sk}$ bezeichnen.

Eine Menge $\{\vec{a}_{S1}; \vec{a}_{S2}; \dots; \vec{a}_{Sk}\}$ von Spaltenvektoren ist genau dann linear unabhängig, wenn aus $\lambda_1 \vec{a}_{S1} + \lambda_2 \vec{a}_{S2} + \dots + \lambda_k \vec{a}_{Sk} = \vec{0}$ folgt $\lambda_1 = \dots = \lambda_k = 0$. Gleichbedeutend damit ist, dass das LGS

$$(*) \quad \begin{array}{cccccc} \lambda_1 a_{11} & + & \lambda_2 a_{12} & + & \dots & + & \lambda_k a_{1k} & = & 0 \\ \lambda_1 a_{21} & + & \lambda_2 a_{22} & + & \dots & + & \lambda_k a_{2k} & = & 0 \\ \vdots & & \vdots & & \vdots & & \vdots & & \vdots \\ \lambda_1 a_{m1} & + & \lambda_2 a_{m2} & + & \dots & + & \lambda_k a_{mk} & = & 0 \end{array}$$

nur die triviale Lösung besitzt. Dies ist unabhängig von der Reihenfolge der Gleichungen und somit von der Reihenfolge der Zeilen der gegebenen Matrix. Zeilenvertauschungen verändern also nicht die lineare Abhängigkeit bzw. Unabhängigkeit beliebiger Mengen von Spaltenvektoren einer Matrix und somit auch nicht deren Spaltenrang.

Z2. Die Multiplikation einer Zeile einer Matrix mit einer von Null verschiedenen reellen Zahl bewirkt eine Multiplikation einer der Gleichungen des LGS (*) mit dieser Zahl. Dabei handelt es sich um eine Äquivalenzumformung, welche die Lösungsmenge des LGS nicht verändert (vgl. Abschnitt 1.2.2). Somit bleibt die lineare Abhängigkeit bzw. Unabhängigkeit einer beliebigen Menge von Spaltenvektoren einer Matrix und damit auch ihr Spaltenrang erhalten.

Z3. Auch die Addition einer Zeile zu einer anderen Zeile einer Matrix macht sich bei der Untersuchung der linearen Abhängigkeit bzw. Unabhängigkeit einer beliebigen Menge von Spaltenvektoren dadurch bemerkbar, dass in dem LGS (*) eine Gleichung zu einer anderen Gleichung addiert wird. Die Lösungsmenge dieses LGS ändert sich auch hierbei nicht; somit bleibt auch bei dieser Umformung der Spaltenrang der gegebenen Matrix erhalten. $\square$

Zum Nachweis der Tatsache, dass auch Spaltenumformungen der Typen S2 und S3 Spaltenränge von Matrizen nicht ändern, siehe die Aufgabe 3 auf S. 234.

Satz 6.2

Jede Matrix lässt sich durch Zeilen- und Spaltenumformungen der Typen Z1-S3 (siehe Satz 6.1) in eine Matrix umformen, die aus einer Diagonalmatrix und ansonsten nur aus Nullzeilen und -spalten besteht:

$$\begin{pmatrix} \mathbf{a}_{11} & 0 & \dots & 0 & 0 & \dots & 0 \\ 0 & \mathbf{a}_{22} & \ddots & & & & \\ \vdots & \ddots & \ddots & & 0 & & \\ 0 & \dots & 0 & \mathbf{a}_{rr} & 0 & \dots & 0 \\ 0 & \dots & & & 0 & 0 & \dots & 0 \\ \vdots & & & & \vdots & \vdots & \vdots & \\ 0 & \dots & & & 0 & 0 & \dots & 0 \end{pmatrix}.$$

Wir verzichten auf den Beweis des Satzes 6.2; die prinzipielle Vorgehensweise bestünde in der Umformung einer allgemein vorgegebenen Matrix mit den auf S. 231f. an einer speziellen Matrix vorgenommenen Umformungsschritten.

Aus den Sätzen 6.1 und 6.2 ergibt sich unmittelbar der folgende Satz.

Satz 6.3

Der Zeilenrang einer beliebigen Matrix ist gleich ihrem Spaltenrang.

Nachdem nun gesichert ist, dass alle bislang betrachteten Ränge einer Matrix (Zeilenrang, Spaltenrang, Anzahl der nach Überführung in die Zeilenstufenform verbleibenden Zeilen, die keine Nullzeilen sind) gleich sind, sprechen wir nur noch von dem *Rang* einer Matrix und verwenden dafür die Schreibweise $\text{rg } \mathbf{A}$.

Von besonderem Interesse sind oft quadratische Matrizen mit „vollem Rang“.

Definition 6.5

Eine $n \times n$ -Matrix $\mathbf{A}$, deren Rang gleich der Zeilen- und Spaltenanzahl von $\mathbf{A}$ ist (d. h. $\text{rg } \mathbf{A} = n$), heißt *reguläre Matrix*. ♦

6.1.3 Aufgaben zu Abschnitt 6.1

- Bestimmen Sie die Ränge der folgenden Matrizen mithilfe des Gauß-Algorithmus (vgl. Abschnitt 1.3):

$$\mathbf{A} = \begin{pmatrix} 2 & -3 & 2 & -1 \\ 3 & -2 & 1 & 1 \\ 3 & 4 & 0 & 7 \\ 4 & 5 & -1 & 9 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 2 & 6 & -10 \\ -3 & -9 & 15 \\ 5 & 2 & 1 \\ 1 & -3 & 7 \end{pmatrix}, \quad \mathbf{C} = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 6 & -4 & -1 & 3 & 6 \\ 4 & 5 & 0 & 9 & 3 \end{pmatrix}.$$

- Bestimmen Sie die Spalten- und die Zeilenränge der folgenden Matrizen und geben Sie Basen der von den Zeilen- und den Spaltenvektoren dieser Matrizen erzeugten linearen Unterräume an (vgl. Beispiel 6.1 auf S. 230):

$$\mathbf{A} = \begin{pmatrix} 1 & 5 & 9 & 1 \\ 2 & 6 & 10 & 1 \\ 3 & 7 & 11 & -1 \\ 4 & 8 & 12 & -1 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 7 & 5 & 9 \\ 2 & 11 & 0 \\ -3 & 11 & 3 \\ -5 & -6 & 1 \end{pmatrix}, \quad \mathbf{C} = \begin{pmatrix} 11 & 5 & 1 & 5 \\ 0 & 6 & 0 & 4 \\ -1 & -3 & 12 & 3 \\ -5 & -4 & -5 & 2 \end{pmatrix}.$$

- Beweisen Sie, dass der Spaltenrang einer beliebigen Matrix bei Spaltenumformungen der Typen S2 und S3 (vgl. Satz 6.1, S. 232) unverändert bleibt.
- Welche der in den Aufgaben 1 und 2 gegebenen Matrizen sind regulär?

6.2 Matrizenmultiplikation und -inversion

6.2.1 Einführungsbeispiel: Materialverflechtung

Beispiel 6.2

Aus drei verschiedenen Rohstoffen R_1 , R_2 und R_3 werden in einem Produktionsablauf zwei Zwischenprodukte Z_1 , Z_2 hergestellt, welche dann zu vier Endprodukten E_1 , E_2 , E_3 und E_4 weiterverarbeitet werden. Wie viele Mengeneinheiten der verschiedenen Rohstoffe zur Herstellung der jeweiligen Zwischenprodukte und wie viele Zwischenprodukte zur Herstellung der verschiedenen Endprodukte benötigt werden, geht aus der Abb. 6.2 hervor.

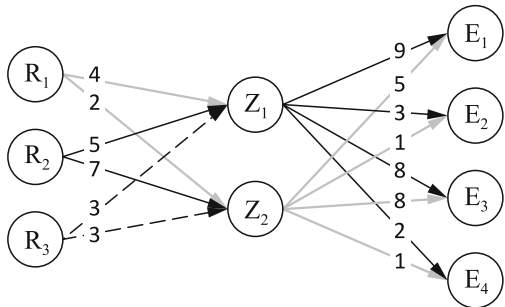


Abb. 6.2: Materialverflechtung

In den folgenden Tabellen sind die der Abb. 6.2 entnommenen Mengenangaben wiedergegeben. Dabei geben die Spalten den Bedarf an Rohstoffen bzw. Zwischenprodukten für die jeweiligen Zwischen- bzw. Endprodukte an.

	Z ₁	Z ₂		E ₁	E ₂	E ₃	E ₄
R ₁	4	2	Z ₁	9	3	8	2
R ₂	5	7	Z ₂	5	1	8	1
R ₃	3	3					

Es soll berechnet werden, wie viele Mengeneinheiten (ME) der verschiedenen Rohstoffe für die Produktion von 60 (ME) des Endproduktes E_1 , 150 (ME) E_2 , 40 (ME) E_3 sowie 200 (ME) E_4 erforderlich sind. Dies lässt sich natürlich einzeln anhand der Tabellen ausrechnen, z. B. für die gewünschte Menge an E_1 :

$$\begin{aligned} 60 E_1 &= 60 \cdot 9 Z_1 + 60 \cdot 5 Z_2 \\ &= 60 \cdot 9 (4 R_1 + 5 R_2 + 3 R_3) + 60 \cdot 5 (2 R_1 + 7 R_2 + 3 R_3) \\ &= 2760 R_1 + 4800 R_2 + 2520 R_3, \end{aligned}$$

wobei die Mengeneinheit (ME) weggelassen wurde. Wir verzichten darauf, für die weiteren Endprodukte auf diese Weise den Rohstoffbedarf zu ermitteln, sondern entwickeln ein Verfahren, dies auf übersichtlichere Weise mithilfe von Matrizen zu tun. Dazu stellen wir zunächst die gewünschten Mengen e_1, e_2, e_3, e_4 der Endprodukte E_1, E_2, E_3, E_4 als Spaltenvektor dar:

$$\text{Outputvektor: } \vec{e} = \begin{pmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \end{pmatrix} = \begin{pmatrix} 60 \\ 150 \\ 40 \\ 200 \end{pmatrix}.$$

Zunächst ermitteln wir den Bedarf an Zwischenprodukten, indem wir $\vec{e}$ mit einer *Verflechtungsmatrix* verknüpfen, welche die rechte Bedarfstabelle enthält:

$$\mathbf{A}_{Z \rightarrow E} = \begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{pmatrix} = \begin{pmatrix} 9 & 3 & 8 & 2 \\ 5 & 1 & 8 & 1 \end{pmatrix}.$$

Das Ergebnis einer geeigneten Verknüpfung $\mathbf{A}_{Z \rightarrow E} \circ \vec{e}$ muss ein Vektor $\vec{z} = \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}$ sein, der den Bedarf an Zwischenprodukten angibt. Die Komponenten dieses Vektors (d. h. die benötigten Mengen der jeweiligen Zwischenprodukte) ergeben sich, indem für jede Zeile der Verflechtungsmatrix deren Koeffizienten mit den zugehörigen Komponenten des Outputvektors multipliziert und die Ergebnisse addiert werden:

$$z_1 = a_{11} e_1 + a_{12} e_2 + a_{13} e_3 + a_{14} e_4 = 60 \cdot 9 + 150 \cdot 3 + 40 \cdot 8 + 200 \cdot 2 = 1710$$

$$z_2 = a_{21} e_1 + a_{22} e_2 + a_{23} e_3 + a_{24} e_4 = 60 \cdot 5 + 150 \cdot 1 + 40 \cdot 8 + 200 \cdot 1 = 970.$$

Die Ermittlung der Komponenten des Vektors $\vec{z}$ erinnert an die Einführung des Skalarproduktes anhand des Beispiels 3.11 und der Definition 3.9 auf S. 122. Ebenso wie das Skalarprodukt zweier Vektoren berechnet wird, erfolgt die „Multiplikation“ einer Zeile der Verflechtungsmatrix mit dem Outputvektor als Summe der Produkte von Komponenten des Outputvektors mit den entsprechenden Koeffizienten der Verflechtungsmatrix. Da jedoch Matrizen meist mehr als eine Zeile haben, ergibt sich als Ergebnis nicht eine einzige Zahl, sondern ein Vektor, dessen Komponenten den jeweiligen Zeilen der Verflechtungsmatrix entsprechen:

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{pmatrix} \circ \begin{pmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \end{pmatrix} = \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}, \quad \begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{pmatrix} \circ \begin{pmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \end{pmatrix} = \begin{pmatrix} z_1 \\ z_2 \end{pmatrix}.$$

Der so ermittelte „Zwischenvektor“ $\vec{z} = \begin{pmatrix} 1710 \\ 970 \end{pmatrix}$ wird nun auf dieselbe Weise mit der Verflechtungsmatrix

$$\mathbf{B}_{R \rightarrow Z} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{pmatrix} = \begin{pmatrix} 4 & 2 \\ 5 & 7 \\ 3 & 3 \end{pmatrix}$$

verknüpft, die den Rohstoffbedarf für die Herstellung von Zwischenprodukten angibt. Als Ergebnis erhält man einen „Inputvektor“ $\vec{r}$, dessen Komponenten die benötigten Rohstoffmengen sind:

$$\begin{aligned} \vec{r} = \mathbf{B}_{R \rightarrow Z} \circ \vec{z} &= \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{pmatrix} \circ \begin{pmatrix} z_1 \\ z_2 \end{pmatrix} = \begin{pmatrix} b_{11} z_1 + b_{12} z_2 \\ b_{21} z_1 + b_{22} z_2 \\ b_{31} z_1 + b_{32} z_2 \end{pmatrix} \\ &= \begin{pmatrix} 4 & 2 \\ 5 & 7 \\ 3 & 3 \end{pmatrix} \circ \begin{pmatrix} 1710 \\ 970 \end{pmatrix} = \begin{pmatrix} 4 \cdot 1710 + 2 \cdot 970 \\ 5 \cdot 1710 + 7 \cdot 970 \\ 3 \cdot 1710 + 3 \cdot 970 \end{pmatrix} = \begin{pmatrix} 8780 \\ 15340 \\ 8040 \end{pmatrix}. \end{aligned}$$

Für die Herstellung der gewünschten Mengen an Endprodukten werden also 8780 (ME) des Rohstoffs R_1 , 15340 (ME) R_2 und 8040 (ME) R_3 benötigt.

Um den Rohstoffbedarf für andere Outputvektoren zu berechnen, müsste das vollzogene zweischrittige Rechenverfahren erneut durchgeführt werden. Daher ist es sinnvoll, die Verflechtungsmatrizen $\mathbf{B}_{R \rightarrow Z}$ und $\mathbf{A}_{Z \rightarrow E}$ zu einer einzigen Matrix $\mathbf{C}_{R \rightarrow E}$ zu verknüpfen, mithilfe derer sich der Inputvektor durch $\vec{r} = \mathbf{C}_{R \rightarrow E} \circ \vec{e}$ direkt aus dem Outputvektor ermitteln lässt. Dazu muss sich die erste Komponente des Inputvektors als Verknüpfung der ersten Zeile von $\mathbf{C}_{R \rightarrow E}$ mit dem Outputvektor ergeben (analog für die weiteren Komponenten

und Zeilen). Um also die erste Zeile der Matrix $\mathbf{C}_{R \rightarrow E}$ zu ermitteln, vollziehen wir zurück, wie die erste Komponente des Inputvektors berechnet wurde:

$$\begin{aligned} r_1 &= b_{11} z_1 + b_{12} z_2 = b_{11} (a_{11} e_1 + a_{12} e_2 + a_{13} e_3 + a_{14} e_4) \\ &\quad + b_{12} (a_{21} e_1 + a_{22} e_2 + a_{23} e_3 + a_{24} e_4) \\ &= (b_{11} a_{11} + b_{12} a_{21}) e_1 + (b_{11} a_{12} + b_{12} a_{22}) e_2 \\ &\quad + (b_{11} a_{13} + b_{12} a_{23}) e_3 + (b_{11} a_{14} + b_{12} a_{24}) e_4 \\ &= (4 \cdot 9 + 2 \cdot 5) \cdot 60 + (4 \cdot 3 + 2 \cdot 1) \cdot 150 + (4 \cdot 8 + 2 \cdot 8) \cdot 40 + (4 \cdot 2 + 2 \cdot 1) \cdot 200 = 8780. \end{aligned}$$

Damit sich also durch Verknüpfung der ersten Zeile der gesuchten „Gesamtverflechtungsmatrix“ $\mathbf{C}_{R \rightarrow E}$ mit dem Outputvektor die richtige Komponente r_1 des Inputvektors ergibt, muss der erste Zeilenvektor der Matrix $\mathbf{C}_{R \rightarrow E}$ die Gestalt

$$\vec{a}_{Z1} = (b_{11} a_{11} + b_{12} a_{21}; b_{11} a_{12} + b_{12} a_{22}; b_{11} a_{13} + b_{12} a_{23}; b_{11} a_{14} + b_{12} a_{24})$$

haben. Völlig analog dazu überlegt man sich, wie die anderen Zeilen beschaffen sein müssen und erhält

$$\mathbf{C}_{R \rightarrow E} = \begin{pmatrix} b_{11} a_{11} + b_{12} a_{21} & b_{11} a_{12} + b_{12} a_{22} & b_{11} a_{13} + b_{12} a_{23} & b_{11} a_{14} + b_{12} a_{24} \\ b_{21} a_{11} + b_{22} a_{21} & b_{21} a_{12} + b_{22} a_{22} & b_{21} a_{13} + b_{22} a_{23} & b_{21} a_{14} + b_{22} a_{24} \\ b_{31} a_{11} + b_{32} a_{21} & b_{31} a_{12} + b_{32} a_{22} & b_{31} a_{13} + b_{32} a_{23} & b_{31} a_{14} + b_{32} a_{24} \end{pmatrix}.$$

Diese Matrix, die durch „Verknüpfung“ der Matrizen $\mathbf{B}_{R \rightarrow Z}$ und $\mathbf{A}_{Z \rightarrow E}$ entstanden ist, heißt *Produkt der Matrizen* $\mathbf{B}_{R \rightarrow Z}$ und $\mathbf{A}_{Z \rightarrow E}$:

$$\mathbf{C}_{R \rightarrow E} = \mathbf{B}_{R \rightarrow Z} \circ \mathbf{A}_{Z \rightarrow E} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{pmatrix} \circ \begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \end{pmatrix}.$$

Damit das Produkt $\mathbf{C} = \mathbf{B} \circ \mathbf{A}$ zweier Matrizen gebildet werden kann, muss die Anzahl der Zeilen der rechten Matrix $\mathbf{A}$ mit der Spaltenanzahl der linken Matrix $\mathbf{B}$ übereinstimmen. Die Zeilenanzahl der Produktmatrix ist gleich der Zeilenanzahl von $\mathbf{B}$, ihre Spaltenanzahl gleich der Anzahl der Spalten von $\mathbf{A}$.

Wir bilden nun das Produkt mit den konkreten Zahlen der in diesem Beispiel betrachteten Verflechtungsmatrizen:

$$\mathbf{C}_{R \rightarrow E} = \mathbf{B}_{R \rightarrow Z} \circ \mathbf{A}_{Z \rightarrow E} = \begin{pmatrix} 4 & 2 \\ 5 & 7 \\ 3 & 3 \end{pmatrix} \circ \begin{pmatrix} 9 & 3 & 8 & 2 \\ 5 & 1 & 8 & 1 \end{pmatrix} = \begin{pmatrix} 46 & 14 & 48 & 10 \\ 80 & 22 & 96 & 17 \\ 42 & 12 & 48 & 9 \end{pmatrix}.$$

Bilden wir nun schließlich (in der zuvor für $\mathbf{A}_{Z \rightarrow E} \circ \vec{e}$ beschriebenen Weise) das Produkt dieser Matrix mit dem Outputvektor $\vec{e}$, so erhalten wir

$$\mathbf{C}_{R \rightarrow E} \circ \vec{e} = \begin{pmatrix} 46 & 14 & 48 & 10 \\ 80 & 22 & 96 & 17 \\ 42 & 12 & 48 & 9 \end{pmatrix} \circ \begin{pmatrix} 60 \\ 150 \\ 40 \\ 200 \end{pmatrix} = \begin{pmatrix} 8780 \\ 15340 \\ 8040 \end{pmatrix},$$

also dasselbe Ergebnis wie bereits zuvor. Wir fassen noch einmal zusammen, auf welchen Wegen wir den Inputvektor berechnet haben:

$$\begin{aligned} \vec{e} &\longrightarrow \vec{z} = \mathbf{A}_{Z \rightarrow E} \circ \vec{e} \longrightarrow \vec{r} = \mathbf{B}_{R \rightarrow Z} \circ \vec{z} = \mathbf{B}_{R \rightarrow Z} \circ (\mathbf{A}_{Z \rightarrow E} \circ \vec{e}) \\ \vec{e} &\longrightarrow \vec{r} = \mathbf{C}_{R \rightarrow E} \circ \vec{e} = (\mathbf{B}_{R \rightarrow Z} \circ \mathbf{A}_{Z \rightarrow E}) \circ \vec{e} \end{aligned}$$

Es lässt sich also (zumindest für das hier behandelte Beispiel) ein Assoziativgesetz für die Multiplikation von Matrizen und Spaltenvektoren (die sich auch als einspaltige Matrizen auffassen lassen) erkennen:

$$\mathbf{B}_{R \rightarrow Z} \circ (\mathbf{A}_{Z \rightarrow E} \circ \vec{e}) = (\mathbf{B}_{R \rightarrow Z} \circ \mathbf{A}_{Z \rightarrow E}) \circ \vec{e}. \quad \blacksquare$$

6.2.2 Matrizenmultiplikation – Definition und Rechenregeln

Wir verallgemeinern die in dem vorangegangenen Abschnitt anhand konkreter Beispiele eingeführte Matrizenmultiplikation zu einer allgemeinen Definition.

Definition 6.6

Es seien $\mathbf{A} \in \mathbb{R}^{l \times m}$ eine $l \times m$ -Matrix und $\mathbf{B} \in \mathbb{R}^{m \times n}$ eine $m \times n$ -Matrix:

$$\mathbf{A} = (a_{ij})_{\substack{i=1 \dots l \\ j=1 \dots m}} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1m} \\ a_{21} & a_{22} & \dots & a_{2m} \\ \vdots & \vdots & & \vdots \\ a_{l1} & a_{l2} & \dots & a_{lm} \end{pmatrix}, \quad \mathbf{B} = (b_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} = \begin{pmatrix} b_{11} & b_{12} & \dots & b_{1n} \\ b_{21} & b_{22} & \dots & b_{2n} \\ \vdots & \vdots & & \vdots \\ b_{m1} & b_{m2} & \dots & b_{mn} \end{pmatrix}.$$

Als Produkt der Matrizen $\mathbf{A}$ und $\mathbf{B}$ wird die Matrix

$$\mathbf{A} \circ \mathbf{B} = \begin{pmatrix} a_{11}b_{11} + a_{12}b_{21} & a_{11}b_{12} + a_{12}b_{22} & \dots & a_{11}b_{1n} + a_{12}b_{2n} \\ + \dots + a_{1m}b_{m1} & + \dots + a_{1m}b_{m2} & & + \dots + a_{1m}b_{mn} \\ a_{21}b_{11} + a_{22}b_{21} & a_{21}b_{12} + a_{22}b_{22} & \dots & a_{21}b_{1n} + a_{22}b_{2n} \\ + \dots + a_{2m}b_{m1} & + \dots + a_{2m}b_{m2} & & + \dots + a_{2m}b_{mn} \\ \vdots & \vdots & & \vdots \\ a_{l1}b_{11} + a_{l2}b_{21} & a_{l1}b_{12} + a_{l2}b_{22} & \dots & a_{l1}b_{1n} + a_{l2}b_{2n} \\ + \dots + a_{lm}b_{m1} & + \dots + a_{lm}b_{m2} & & + \dots + a_{lm}b_{mn} \end{pmatrix}$$

bezeichnet. In Kurzschreibweise lässt sich diese folgendermaßen angeben:

$$\begin{aligned} \mathbf{A} \circ \mathbf{B} &= \begin{pmatrix} \sum_{k=1}^m a_{1k}b_{k1} & \sum_{k=1}^m a_{1k}b_{k2} & \dots & \sum_{k=1}^m a_{1k}b_{kn} \\ \sum_{k=1}^m a_{2k}b_{k1} & \sum_{k=1}^m a_{2k}b_{k2} & \dots & \sum_{k=1}^m a_{2k}b_{kn} \\ \vdots & \vdots & & \vdots \\ \sum_{k=1}^m a_{lk}b_{k1} & \sum_{k=1}^m a_{lk}b_{k2} & \dots & \sum_{k=1}^m a_{lk}b_{kn} \end{pmatrix} \\ &= \left(\sum_{k=1}^m a_{ik}b_{kj} \right)_{\substack{i=1 \dots l \\ j=1 \dots n}}. \end{aligned}$$

Bemerkungen:

- Wie bereits in dem Beispiel 6.2 erwähnt wurde (siehe S. 237, dort jedoch mit umgekehrter Reihenfolge der Matrizen $\mathbf{A}$ und $\mathbf{B}$), muss die Spaltenanzahl m der (linken) Matrix $\mathbf{A}$ mit der Anzahl der Zeilen der (rechten) Matrix $\mathbf{B}$ übereinstimmen, damit das Produkt $\mathbf{A} \circ \mathbf{B}$ gebildet werden kann.
- Die Zeilenanzahl der Produktmatrix $\mathbf{C}$ ist gleich der Zeilenanzahl von $\mathbf{A}$, ihre Spaltenanzahl gleich der Anzahl der Spalten von $\mathbf{B}$, d. h. wenn $\mathbf{A} \in \mathbb{R}^{l \times m}$ und $\mathbf{B} \in \mathbb{R}^{m \times n}$, so ist $\mathbf{A} \circ \mathbf{B} \in \mathbb{R}^{l \times n}$.
- Ein Element c_{ij} der Produktmatrix $\mathbf{C} = \mathbf{A} \circ \mathbf{B}$ ist nach der Definition 6.6 die Summe der Produkte der Komponenten der i -ten Zeile von $\mathbf{A}$ mit den entsprechenden Komponenten der j -ten Spalte von $\mathbf{B}$ bzw., gleichbedeutend damit, das Skalarprodukt des Zeilenvektors $\vec{a}_{zi}$ mit dem Spaltenvektor $\vec{b}_{sj}$.

$$\begin{pmatrix} a_{11} & a_{12} & \dots & a_{1m} \\ \vdots & \vdots & & \vdots \\ \boxed{a_{i1} \quad a_{i2} \quad \dots \quad a_{im}} \\ \vdots & \vdots & & \vdots \\ a_{l1} & a_{l2} & \dots & a_{lm} \end{pmatrix} \circ \begin{pmatrix} b_{11} & \dots & \boxed{b_{1j}} & \dots & b_{1n} \\ b_{21} & \dots & \boxed{b_{2j}} & \dots & b_{2n} \\ \vdots & & \vdots & & \vdots \\ b_{m1} & \dots & \boxed{b_{mj}} & \dots & b_{mn} \end{pmatrix} = \begin{pmatrix} c_{11} & \dots & c_{1j} & \dots & c_{1n} \\ \vdots & & \vdots & & \vdots \\ c_{i1} & \dots & \boxed{c_{ij}} & \dots & c_{in} \\ \vdots & & \vdots & & \vdots \\ c_{l1} & \dots & c_{lj} & \dots & c_{ln} \end{pmatrix}$$

- Mitunter wird das Produkt $\mathbf{A} \circ \mathbf{B}$ zweier Matrizen in der Literatur auch mit $\mathbf{A} \cdot \mathbf{B}$ oder einfach mit $\mathbf{AB}$ bezeichnet.

Satz 6.4

Die Multiplikation von Matrizen ist assoziativ, d. h. für beliebige Matrizen

$\mathbf{A} \in \mathbb{R}^{l \times m}$, $\mathbf{B} \in \mathbb{R}^{m \times n}$ und $\mathbf{C} \in \mathbb{R}^{n \times p}$ gilt $(\mathbf{A} \circ \mathbf{B}) \circ \mathbf{C} = \mathbf{A} \circ (\mathbf{B} \circ \mathbf{C})$.

Wir verzichten auf einen Beweis des Satzes 6.4 und betrachten statt dessen (nachdem sich bereits in dem Beispiel 6.2 mit $p=1$, $n=4$, $m=2$, $l=3$ Assoziativität gezeigt hatte, siehe S. 237) ein weiteres Beispiel.

Beispiel 6.3

Gegeben sind die Matrizen $\mathbf{A} = \begin{pmatrix} 0 & 3 & -2 \\ 2 & 4 & 7 \\ 1 & -8 & 0 \\ 6 & 5 & 8 \end{pmatrix}$, $\mathbf{B} = \begin{pmatrix} 9 & 4 & 3 \\ 12 & 4 & 2 \\ 0 & 5 & 1 \end{pmatrix}$, $\mathbf{C} = \begin{pmatrix} -1 & -9 \\ 0 & 5 \\ 3 & -2 \end{pmatrix}$.

Wir berechnen zunächst $\mathbf{A} \circ \mathbf{B} = \begin{pmatrix} 36 & 2 & 4 \\ 66 & 59 & 21 \\ -87 & -28 & -13 \\ 114 & 84 & 36 \end{pmatrix}$ und $\mathbf{B} \circ \mathbf{C} = \begin{pmatrix} 0 & -67 \\ -6 & -92 \\ 3 & 23 \end{pmatrix}$.

Damit ergibt sich $(\mathbf{A} \circ \mathbf{B}) \circ \mathbf{C} = \begin{pmatrix} -24 & -322 \\ -3 & -341 \\ 48 & 669 \\ -6 & -678 \end{pmatrix}$ und $\mathbf{A} \circ (\mathbf{B} \circ \mathbf{C}) = \begin{pmatrix} -24 & -322 \\ -3 & -341 \\ 48 & 669 \\ -6 & -678 \end{pmatrix}$,

also $(\mathbf{A} \circ \mathbf{B}) \circ \mathbf{C} = \mathbf{A} \circ (\mathbf{B} \circ \mathbf{C})$. ■

Im Gegensatz zur Assoziativität gilt die *Kommutativität* der Matrizenmultiplikation i. Allg. *nicht*, wie die folgenden Gegenbeispiele zeigen.

Beispiel 6.4

Für Matrizen mit unterschiedlichen Zeilen- und Spaltenanzahlen ist es unmittelbar einsichtig, dass die Kommutativität der Matrizenmultiplikation nicht gegeben sein kann; z. B. gilt für $\mathbf{A} = \begin{pmatrix} 2 & 3 & -5 \\ 6 & 9 & 7 \end{pmatrix}$ und $\mathbf{B} = \begin{pmatrix} 1 & 2 \\ -3 & 0 \\ 11 & -6 \end{pmatrix}$:

$$\mathbf{A} \circ \mathbf{B} = \begin{pmatrix} -62 & 34 \\ 56 & -30 \end{pmatrix}, \quad \mathbf{B} \circ \mathbf{A} = \begin{pmatrix} 14 & 21 & 9 \\ -6 & -9 & 15 \\ -14 & -21 & -97 \end{pmatrix}.$$

Aber auch für quadratische Matrizen gilt die Kommutativität der Matrizenmultiplikation i. Allg. nicht. So ist für $\mathbf{A} = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$, $\mathbf{B} = \begin{pmatrix} 6 & 5 \\ 7 & 1 \end{pmatrix}$:

$$\mathbf{A} \circ \mathbf{B} = \begin{pmatrix} 20 & 7 \\ 46 & 19 \end{pmatrix}, \quad \mathbf{B} \circ \mathbf{A} = \begin{pmatrix} 21 & 32 \\ 10 & 18 \end{pmatrix},$$

also $\mathbf{A} \circ \mathbf{B} \neq \mathbf{B} \circ \mathbf{A}$. ■

In speziellen Fällen können zwei Matrizen $\mathbf{A}$ und $\mathbf{B}$ „vertauschbar“ sein (d. h. für diese Matrizen gilt $\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A}$), siehe die Aufgaben 3 und 4 auf S. 250.

Das folgende Beispiel zeigt eine weitere schwere Abweichung der Matrizenmultiplikation von den für Zahlen gültigen Rechenregeln. In beliebigen Zahlbereichen folgt aus $a \cdot b = 0$ stets, dass $a = 0$ oder $b = 0$ gelten muss; diese Eigenschaft heißt Nullteilerfreiheit. Hingegen ist die *Matrizenmultiplikation nicht nullteilerfrei*. Eine Nullmatrix (d. h. eine Matrix, die nur aus Nullen besteht) kann auch als Produkt zweier Matrizen entstehen, von denen keine eine Nullmatrix ist.

Beispiel 6.5

Für die Matrizen $\mathbf{A} = \begin{pmatrix} 2 & 2 \\ 3 & 3 \end{pmatrix}$ und $\mathbf{B} = \begin{pmatrix} 1 & 1 \\ -1 & -1 \end{pmatrix}$ ist $\mathbf{A} \circ \mathbf{B} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$. ■

Definition 6.7

Die quadratische Matrix $\mathbf{E}_n \in \mathbb{R}^{n \times n}$, deren Hauptdiagonale nur Einsen enthält und die ansonsten nur aus Nullen besteht, heißt *Einheitsmatrix*:

$$\mathbf{E}_n = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{pmatrix}.$$

◆

Satz 6.5

Für die Multiplikation beliebiger quadratischer Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ mit der Einheitsmatrix $\mathbf{E}_n$ gilt $\mathbf{A} \circ \mathbf{E}_n = \mathbf{E}_n \circ \mathbf{A} = \mathbf{A}$.

Die Einheitsmatrix $\mathbf{E}_n$ hat also unter den $n \times n$ -Matrizen die Funktion des *Einselements* bzw. des *neutralen Elements* bezüglich der Matrizenmultiplikation. Die Gültigkeit des Satzes 6.5 folgt unmittelbar aus der Definition 6.6, denn es ist $e_{kj} = 0$ für $k \neq j$ und $e_{kj} = 1$ für $k = j$. Somit gilt $\sum_{k=1}^n a_{ik} e_{kj} = a_{ij}$ für beliebige Koeffizienten einer Matrix $\mathbf{A}$, und nach der Definition 6.6 (mit $l = m = n$ und $\mathbf{B} = \mathbf{E}_n$) folgt daraus $\mathbf{A} \circ \mathbf{E}_n = \mathbf{A}$ (analog begründet man $\mathbf{E}_n \circ \mathbf{A} = \mathbf{A}$).

Satz 6.6

Falls das Produkt $\mathbf{C} = \mathbf{A} \circ \mathbf{B}$ existiert, so ist $\mathbf{C}^T = \mathbf{B}^T \circ \mathbf{A}^T$.

(Dabei sind $\mathbf{B}^T$, $\mathbf{A}^T$ und $\mathbf{C}^T$ die zu $\mathbf{A}$, $\mathbf{B}$ bzw. $\mathbf{C}$ transponierten Matrizen, siehe die Definition 6.2 auf S. 228.)

Wir verzichten darauf, den Satz 6.6 zu beweisen. Es wird empfohlen, zur inhaltlichen Erschließung dieses Satzes die Aufgabe 5 auf S. 250 zu lösen.

Matrizenmultiplikation sowie -addition und skalare Multiplikation

In dem Beispiel 5.6 (siehe S. 171) wurden die Addition „+“ von Matrizen (mit gleicher Zeilen- und Spaltenanzahl) sowie die Multiplikation „ $\cdot$ “ von Matrizen mit Skalaren eingeführt, in Kurzschreibweise:

$$\begin{aligned} \mathbf{A} + \mathbf{B} &= (a_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} + (b_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} = (a_{ij} + b_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} \\ \lambda \cdot \mathbf{A} &= \lambda \cdot (a_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}} = (\lambda \cdot a_{ij})_{\substack{i=1 \dots m \\ j=1 \dots n}}. \end{aligned}$$

Es wurde festgestellt, dass die Menge aller $m \times n$ -Matrizen mit diesen Verknüpfungen ein Vektorraum ist. Im Folgenden werden Rechenregeln des Matrizenproduktes im Zusammenhang mit den Verknüpfungen „+“ und „ $\cdot$ “ formuliert.

Satz 6.7

Es seien $\mathbf{A}$, $\mathbf{A}^{(1)}$, $\mathbf{A}^{(2)}$ $l \times m$ -Matrizen und $\mathbf{B}$, $\mathbf{B}^{(1)}$, $\mathbf{B}^{(2)}$ $m \times n$ -Matrizen sowie $\lambda \in \mathbb{R}$. Dann gilt

- a) $(\mathbf{A}^{(1)} + \mathbf{A}^{(2)}) \circ \mathbf{B} = \mathbf{A}^{(1)} \circ \mathbf{B} + \mathbf{A}^{(2)} \circ \mathbf{B}$, $\mathbf{A} \circ (\mathbf{B}^{(1)} + \mathbf{B}^{(2)}) = \mathbf{A} \circ \mathbf{B}^{(1)} + \mathbf{A} \circ \mathbf{B}^{(2)}$,
- b) $\lambda \cdot (\mathbf{A} \circ \mathbf{B}) = (\lambda \cdot \mathbf{A}) \circ \mathbf{B} = \mathbf{A} \circ (\lambda \cdot \mathbf{B})$.

Bemerkung: Da die Matrizenmultiplikation i. Allg. nicht kommutativ ist, mussten unter a) zwei Distributivgesetze formuliert werden (während bei kommutativen Verknüpfungen ein Distributivgesetz ausreicht).

Beweis: Es seien $\mathbf{A}^{(1)} = (a_{ij}^{(1)})_{\substack{i=1\dots l \\ j=1\dots m}}$, $\mathbf{A}^{(2)} = (a_{ij}^{(2)})_{\substack{i=1\dots l \\ j=1\dots m}}$, $\mathbf{B} = (b_{ij})_{\substack{i=1\dots m \\ j=1\dots n}}$. Es gilt

$$\mathbf{A}^{(1)} + \mathbf{A}^{(2)} = (a_{ij}^{(1)} + a_{ij}^{(2)})_{\substack{i=1\dots l \\ j=1\dots m}}, \quad (\mathbf{A}^{(1)} + \mathbf{A}^{(2)}) \circ \mathbf{B} = \left(\sum_{k=1}^m (a_{ik}^{(1)} + a_{ik}^{(2)}) b_{kj} \right)_{\substack{i=1\dots l \\ j=1\dots n}},$$

$$\mathbf{A}^{(1)} \circ \mathbf{B} = \left(\sum_{k=1}^m a_{ik}^{(1)} b_{kj} \right)_{\substack{i=1\dots l \\ j=1\dots n}}, \quad \mathbf{A}^{(2)} \circ \mathbf{B} = \left(\sum_{k=1}^m a_{ik}^{(2)} b_{kj} \right)_{\substack{i=1\dots l \\ j=1\dots n}} \quad \text{und somit}$$

$$\mathbf{A}^{(1)} \circ \mathbf{B} + \mathbf{A}^{(2)} \circ \mathbf{B} = \left(\sum_{k=1}^m (a_{ik}^{(1)} + a_{ik}^{(2)}) b_{kj} \right)_{\substack{i=1\dots l \\ j=1\dots n}}.$$

Damit gilt die Behauptung $(\mathbf{A}^{(1)} + \mathbf{A}^{(2)}) \circ \mathbf{B} = \mathbf{A}^{(1)} \circ \mathbf{B} + \mathbf{A}^{(2)} \circ \mathbf{B}$. Der zweite Teil des Distributivgesetzes lässt sich analog dazu nachweisen; zum Nachweis des Satzes 6.7 b siehe die Aufgabe 6 auf S. 250. $\square$

6.2.3 Anwendungsbeispiele zur Matrizenmultiplikation

Nachdem in dem Abschnitt 6.2.1 das Beispiel der Materialverflechtung zur Einführung der Matrizenmultiplikation genutzt wurde, erfolgt nun die Behandlung einer inner- sowie einer weiteren außermathematischen Anwendung.

Beispiel 6.6

In der Aufgabe 10 auf S. 201 wurden *Fibonacci-Zahlen* beschrieben. Für die ursprüngliche Fibonacci-Folge (f_n) gilt

$$f_1 = 1, f_2 = 1 \quad \text{sowie} \quad f_{n+2} = f_n + f_{n+1} \quad (\text{für alle } n \in \mathbb{N}).$$

Daraus lassen sich nacheinander beliebig viele Glieder der Fibonacci-Folge berechnen: $f_1 = 1, f_2 = 1, f_3 = 2, f_4 = 3, f_5 = 5, f_6 = 8, f_7 = 13, f_8 = 21, f_9 = 34, \dots$. Mithilfe von Matrizen lässt sich nun eine recht einfache explizite Vorschrift für die Berechnung von Fibonacci-Zahlen angeben. Es gilt

$$\begin{pmatrix} f_{n+1} & f_n \\ f_n & f_{n-1} \end{pmatrix} = \mathbf{A}^n \quad \text{mit} \quad \mathbf{A} = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}.$$

Dabei versteht man unter $\mathbf{A}^n$ in gewohnter Weise $\mathbf{A} \circ \mathbf{A} \circ \dots \circ \mathbf{A}$ (n -mal). Es ist

$$\mathbf{A}^2 = \begin{pmatrix} 2 & 1 \\ 1 & 1 \end{pmatrix}, \quad \mathbf{A}^3 = \begin{pmatrix} 3 & 2 \\ 2 & 1 \end{pmatrix}, \quad \mathbf{A}^4 = \begin{pmatrix} 5 & 3 \\ 3 & 2 \end{pmatrix}, \quad \mathbf{A}^5 = \begin{pmatrix} 8 & 5 \\ 5 & 3 \end{pmatrix}, \quad \mathbf{A}^6 = \begin{pmatrix} 13 & 8 \\ 8 & 5 \end{pmatrix}, \dots$$

Aus der matriziellen Darstellung lassen sich einige Gesetzmäßigkeiten für Fibonacci-Zahlen ableiten. So gilt beispielsweise

$$\begin{pmatrix} f_{m+n+1} & f_{m+n} \\ f_{m+n} & f_{m+n-1} \end{pmatrix} = \mathbf{A}^{m+n} = \mathbf{A}^m \circ \mathbf{A}^n = \begin{pmatrix} f_{m+1} & f_m \\ f_m & f_{m-1} \end{pmatrix} \circ \begin{pmatrix} f_{n+1} & f_n \\ f_n & f_{n-1} \end{pmatrix}.$$

Daraus lässt sich durch Bildung des Produkts der beiden letzten Matrizen und Vergleich mit den Koeffizienten der vorderen Matrix ableiten, dass

$$f_{m+n} = f_{m+1}f_n + f_m f_{n-1} \quad \text{und} \quad f_{m+n} = f_m f_{n+1} + f_{m-1}f_n$$

ist. Somit gilt für den Spezialfall $m = n+1$:

$$f_{2n+1} = f_{n+1}^2 + f_n^2.$$

Mit dieser Gleichung lässt sich nun z. B. das 17. Glied der Fibonacci-Folge anhand der oben angegebenen Glieder f_8 und f_9 berechnen, ohne die dazwischen liegenden Glieder berechnen zu müssen:

$$f_{17} = f_9^2 + f_8^2 = 34^2 + 21^2 = 1597. \quad \blacksquare$$

Beispiel 6.7**Populationsmatrizen**

Matrizen lassen sich gut zur Beschreibung diskreter „verflochtener“ Wachstumsprozesse (wie der Entwicklung von Populationen) nutzen, siehe hierzu ausführlicher Lehmann (1983). Ein einfaches Beispiel ist das folgende:

Aus den Eiern eines Käfers schlüpfen nach einem Monat Larven. Nach einem weiteren Monat werden diese zu Käfern, die nach einem Monat jeweils 8 Eier legen und dann sofort sterben. Nur aus einem Viertel der Eier werden Larven, die anderen Eier werden gefressen oder verenden. Von den Larven wird die Hälfte zu Käfern, die andere Hälfte stirbt.

Die folgende Tabelle fasst die Umwandlungen zusammen:

Zeitpunkt t	Zeitpunkt $t+1$ Monat
Käfer	8 Eier
Ei	$\frac{1}{4}$ Larve
Larve	$\frac{1}{2}$ Käfer

Eine andere Darstellungsweise des Vorgangs (die mit der Darstellung von Materialverflechtungen in Tabellen vergleichbar ist, siehe S. 235) ist die folgende:

Ei wird zu	Larve wird zu	Käfer wird zu	
		8	Eier
$\frac{1}{4}$			Larven
	$\frac{1}{2}$		Käfer

Diese Tabelle schreiben wir nun als *Populationsmatrix*: $\mathbf{P} = \begin{pmatrix} 0 & 0 & 8 \\ \frac{1}{4} & 0 & 0 \\ 0 & \frac{1}{2} & 0 \end{pmatrix}$.

Es seien zu einem Zeitpunkt gleiche Anzahlen von 1000 Eiern, 1000 Larven und 1000 Käfern vorhanden. Diese Angabe fassen wir zu einem *Populationsvektor*

$\vec{p}_0 = \begin{pmatrix} 1000 \\ 1000 \\ 1000 \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \end{matrix}$ zusammen. Durch Multiplikation dieses Vektors mit $\mathbf{P}$

erhält man den Populationsvektor nach einem Monat:

$$\vec{p}_1 = \mathbf{P} \circ \vec{p}_0 = \begin{pmatrix} 0 & 0 & 8 \\ \frac{1}{4} & 0 & 0 \\ 0 & \frac{1}{2} & 0 \end{pmatrix} \circ \begin{pmatrix} 1000 \\ 1000 \\ 1000 \end{pmatrix} = \begin{pmatrix} 8000 \\ 250 \\ 500 \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \end{matrix}.$$

Um den Populationsbestand nach weiteren Monaten zu berechnen, wird der jeweils aktuelle Populationsvektor erneut mit der Populationsmatrix multipliziert:

$$\vec{p}_2 = \mathbf{P} \circ \vec{p}_1 = \begin{pmatrix} 4000 \\ 2000 \\ 125 \end{pmatrix}, \quad \vec{p}_3 = \mathbf{P} \circ \vec{p}_2 = \begin{pmatrix} 1000 \\ 1000 \\ 1000 \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \end{matrix}.$$

Offensichtlich hat sich nach drei Monaten wieder die Ausgangspopulation hergestellt. Dies bedeutet, dass $\mathbf{P}^3 \circ \vec{p}_0 = \vec{p}_0$ ist. Um diese Tatsache näher zu untersuchen, berechnen wir $\mathbf{P}^3$:

$$\mathbf{P}^2 = \mathbf{P} \circ \mathbf{P} = \begin{pmatrix} 0 & 0 & 8 \\ \frac{1}{4} & 0 & 0 \\ 0 & \frac{1}{2} & 0 \end{pmatrix} \circ \begin{pmatrix} 0 & 0 & 8 \\ \frac{1}{4} & 0 & 0 \\ 0 & \frac{1}{2} & 0 \end{pmatrix} = \begin{pmatrix} 0 & 4 & 0 \\ 0 & 0 & 2 \\ \frac{1}{8} & 0 & 0 \end{pmatrix}, \quad \mathbf{P}^3 = \mathbf{P} \circ \mathbf{P}^2 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

$\mathbf{P}^3$ ist also die Einheitsmatrix $\mathbf{E}_3$, somit muss $\mathbf{P}^3 \circ \vec{p}_0 = \vec{p}_0$ gelten; es gilt sogar $\mathbf{P}^3 \circ \vec{x} = \vec{x}$ für beliebige Vektoren $\vec{x} \in \mathbb{R}^3$.

Während bei der Multiplikation mit der Einheitsmatrix *jeder* Vektor unverändert bleibt, kann es durchaus *spezielle* Vektoren geben, die auch bei der Multiplikation mit anderen Matrizen unverändert bleiben. Wir untersuchen nun, ob es einen Vektor $\vec{q}$ gibt, für den $\mathbf{P} \circ \vec{q} = \vec{q}$ ist. Innerhalb der hier betrachteten Sachsituation kommt dies folgender Frage gleich: *Gibt es eine Anfangspopulation von q_E Eiern, q_L Larven und q_K Käfern, deren Anzahlen sich nach einem Monat nicht verändert haben?* Zur Beantwortung dieser Frage setzen wir

$$\vec{q} = \mathbf{P} \circ \vec{q}, \text{ d. h. } \begin{pmatrix} q_E \\ q_L \\ q_K \end{pmatrix} = \begin{pmatrix} 0 & 0 & 8 \\ \frac{1}{4} & 0 & 0 \\ 0 & \frac{1}{2} & 0 \end{pmatrix} \circ \begin{pmatrix} q_E \\ q_L \\ q_K \end{pmatrix} = \begin{pmatrix} 8q_K \\ \frac{1}{4}q_E \\ \frac{1}{2}q_L \end{pmatrix}.$$

Diese Bedingung ist erfüllt für $q_E = 8q_K = 4q_L$. Somit ist z. B. $\vec{q} = \begin{pmatrix} 8 \\ 2 \\ 1 \end{pmatrix}$ ein Vektor mit $\mathbf{P} \circ \vec{q} = \vec{q}$, der eine „stabile Population“ beschreibt. ■

Beispiel 6.8

Für eine andere Käferart wird die in dem Beispiel 6.7 gegebene Situation etwas abgewandelt: Statt *Käfern, die nach einem Monat jeweils 8 Eier legen und dann sofort sterben* betrachten wir *Käfer* (die ebenfalls nach einem Monat jeweils 8 Eier legen), *von denen 75% sterben und 25% noch einen weiteren Monat als „alte Käfer“ leben, dann noch vier Eier legen und nun sterben*. Die folgende Tabelle beschreibt diese Entwicklung.

Ei wird zu	Larve wird zu	Käfer wird zu	alter Käfer wird	
		8	4	Eier
$\frac{1}{4}$				Larven
	$\frac{1}{2}$			Käfer
		$\frac{1}{4}$		alte Käfer

Die zugehörige Populationsmatrix ist $\mathbf{P} = \begin{pmatrix} 0 & 0 & 8 & 4 \\ \frac{1}{4} & 0 & 0 & 0 \\ 0 & \frac{1}{2} & 0 & 0 \\ 0 & 0 & \frac{1}{4} & 0 \end{pmatrix}$. Wir untersuchen die

Populationsentwicklung für einen Ausgangsbestand $\vec{p}_0 = \begin{pmatrix} 1000 \\ 1000 \\ 1000 \\ 1000 \end{pmatrix} \begin{pmatrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \\ \text{(alte Käfer)} \end{pmatrix}$.

Durch mehrfache Multiplikation des Populationsvektors $\vec{p}_0$ mit der Populationsmatrix erhalten wir die Bestände $\vec{p}_1 = \mathbf{P} \circ \vec{p}_0$, $\vec{p}_2 = \mathbf{P} \circ \vec{p}_1$, $\vec{p}_3 = \mathbf{P} \circ \vec{p}_2$ und $\vec{p}_4 = \mathbf{P} \circ \vec{p}_3$ nach 1, 2, 3 bzw. 4 Monaten:

$$\vec{p}_1 = \begin{pmatrix} 12000 \\ 250 \\ 500 \\ 250 \end{pmatrix}, \vec{p}_2 = \begin{pmatrix} 5000 \\ 3000 \\ 125 \\ 125 \end{pmatrix}, \vec{p}_3 = \begin{pmatrix} 1500 \\ 1250 \\ 1500 \\ 31,25 \end{pmatrix}, \vec{p}_4 = \begin{pmatrix} 12125 \\ 375 \\ 625 \\ 375 \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \\ \text{(alte Käfer)}. \end{matrix}$$

Hieran sind noch keine „Regelmäßigkeiten“ zu erkennen. Um einen Eindruck von der langfristigen Entwicklung zu erhalten, lassen sich mithilfe des Computers (siehe dazu den Abschnitt 6.4) z. B. die Populationsvektoren nach 100, 101 und 102 Monaten bestimmen. Man erhält (nach Rundung auf ganzzahlige Werte):

$$\vec{p}_{100} \approx \begin{pmatrix} 250205 \\ 60620 \\ 28985 \\ 6973 \end{pmatrix}, \vec{p}_{101} \approx \begin{pmatrix} 259769 \\ 62551 \\ 30309 \\ 7246 \end{pmatrix}, \vec{p}_{102} \approx \begin{pmatrix} 271462 \\ 64942 \\ 31275 \\ 7577 \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \\ \text{(alte Käfer)}. \end{matrix}$$

Es entsteht der Eindruck eines allmählichen Wachstums aller Populationen. Dieser Eindruck verstärkt sich durch die Betrachtung der Populationen nach 1000, 1001 und 1002 Monaten:

$$\vec{p}_{1000} \approx \begin{pmatrix} 1.61 \cdot 10^{20} \\ 3.87 \cdot 10^{19} \\ 1.86 \cdot 10^{19} \\ 4.48 \cdot 10^{18} \end{pmatrix}, \vec{p}_{1001} \approx \begin{pmatrix} 1.67 \cdot 10^{20} \\ 4.01 \cdot 10^{19} \\ 1.93 \cdot 10^{19} \\ 4.65 \cdot 10^{18} \end{pmatrix}, \vec{p}_{1002} \approx \begin{pmatrix} 1.7 \cdot 10^{20} \\ 4.17 \cdot 10^{19} \\ 2.01 \cdot 10^{19} \\ 4.83 \cdot 10^{18} \end{pmatrix} \begin{matrix} \text{(Eier)} \\ \text{(Larven)} \\ \text{(Käfer)} \\ \text{(alte K.)}. \end{matrix}$$

In dem Beispiel 6.7 war es möglich, eine Anfangspopulation zu ermitteln, die sich nach einem Monat nicht verändert (und daher für alle Zeit konstant bleibt). Jedoch erscheint es aufgrund des in diesem Beispiel nun festzustellenden ständigen Wachstums aller Populationen unwahrscheinlich, dass sich eine solche Anfangspopulation auch hierfür finden lässt. Um dieser Frage auf den Grund zu gehen, untersuchen wir erneut, ob es einen Vektor $\vec{q}$ mit

$$\vec{q} = \mathbf{P} \circ \vec{q}, \text{ d. h. } \begin{pmatrix} q_E \\ q_L \\ q_K \\ q_{aK} \end{pmatrix} = \begin{pmatrix} 0 & 0 & 8 & 4 \\ \frac{1}{4} & 0 & 0 & 0 \\ 0 & \frac{1}{2} & 0 & 0 \\ 0 & 0 & \frac{1}{4} & 0 \end{pmatrix} \circ \begin{pmatrix} q_E \\ q_L \\ q_K \\ q_{aK} \end{pmatrix} = \begin{pmatrix} 8q_K + 4q_{aK} \\ \frac{1}{4}q_E \\ \frac{1}{2}q_L \\ \frac{1}{4}q_K \end{pmatrix}$$

gibt. Dazu ist das folgende lineare Gleichungssystem zu lösen:

$$\begin{array}{ll} q_E = 8q_K + 4q_{aK} & q_E - 8q_K - 4q_{aK} = 0 \\ q_L = \frac{1}{4}q_E & \text{bzw.} \quad -\frac{1}{4}q_E + q_L = 0 \\ q_K = \frac{1}{2}q_L & -\frac{1}{2}q_L + q_K = 0 \\ q_{aK} = \frac{1}{4}q_K & -\frac{1}{4}q_K + q_{aK} = 0. \end{array}$$

Dieses LGS besitzt nur die Lösung $q_E = q_L = q_K = q_{aK} = 0$. Somit existiert, abgesehen von Nullbeständen, keine Anfangspopulation, die konstant bleibt. Damit eine solche existieren würde, müsste das LGS einen niedrigeren Rang als vier haben. Dies wäre der Fall, wenn z. B. die erste Zeile als Linearkombination der restlichen drei Zeilen darstellbar wäre. Man könnte dies u. a. durch Ersetzung des Koeffizienten $-\frac{1}{2}$ vor q_L in der dritten Gleichung durch $-\frac{4}{9}$ erreichen. Die entsprechende Veränderung der Populationsmatrix ließe sich durch die monatliche Vernichtung eines Teils der Larven realisieren, sodass sich die folgende Veränderung der Ausgangssituation ergibt: *Von den Larven werden vier Neuntel zu Käfern, fünf Neuntel sterben* (siehe die Aufgabe 7 auf S. 250). ■

6.2.4 Inverse Matrizen

In der Definition 6.7 auf S. 240 wurde die Einheitsmatrix $\mathbf{E}_n = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & & \vdots \\ \vdots & & \ddots & 0 \\ 0 & \dots & 0 & 1 \end{pmatrix}$ eingeführt und in dem Satz 6.5 festgestellt, dass $\mathbf{E}_n$ in der Menge der $n \times n$ -Matrizen die Funktion des *Einselements* bzw. *neutralen Elements* bezüglich der Multiplikation innehat, also mit der Eins in den reellen Zahlen vergleichbar ist. Im Folgenden untersuchen wir, ob es in der Menge der $n \times n$ -Matrizen bezüglich der Matrizenmultiplikation *inverse Elemente* gibt, d. h. ob sich zu Matrizen $\mathbf{A} \in \mathbb{R}^{n \times n}$ Matrizen $\mathbf{A}^{-1}$ finden lassen, für die $\mathbf{A} \circ \mathbf{A}^{-1} = \mathbf{A}^{-1} \circ \mathbf{A} = \mathbf{E}_n$ gilt.

Beispiel 6.9

Es wird die inverse Matrix $\mathbf{A}^{-1} = \begin{pmatrix} a'_{11} & a'_{12} & a'_{13} \\ a'_{21} & a'_{22} & a'_{23} \\ a'_{31} & a'_{32} & a'_{33} \end{pmatrix}$ zu $\mathbf{A} = \begin{pmatrix} 2 & -5 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix}$ gesucht. Es muss dazu gelten

$$\begin{pmatrix} 2 & -5 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix} \circ \begin{pmatrix} a'_{11} & a'_{12} & a'_{13} \\ a'_{21} & a'_{22} & a'_{23} \\ a'_{31} & a'_{32} & a'_{33} \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Daraus ergibt sich nach der Definition 6.6 der Matrizenmultiplikation, dass die Komponenten $a'_{11} \dots a'_{33}$ von $\mathbf{A}^{-1}$ folgende Gleichungen erfüllen müssen:

$$\begin{array}{rrrr} 2a'_{11} & -5a'_{21} & +3a'_{31} & = 1 \\ -a'_{11} & +2a'_{21} & -2a'_{31} & = 0 \\ 3a'_{11} & & -a'_{31} & = 0 \\ & 2a'_{12} & -5a'_{22} & +3a'_{32} = 0 \\ & -a'_{12} & +2a'_{22} & -2a'_{32} = 1 \\ & 3a'_{12} & & -a'_{32} = 0 \\ & 2a'_{13} & -5a'_{23} & +3a'_{33} = 0 \\ & -a'_{13} & +2a'_{23} & -2a'_{33} = 0 \\ & 3a'_{13} & & -a'_{33} = 1. \end{array}$$

Dieses LGS enthält drei Teilsysteme, die unabhängig voneinander gelöst werden können, da sie verschiedene Variablen enthalten (so treten z. B. a'_{11} , a'_{21} und a'_{31} nur in den ersten drei Gleichungen auf). Man erhält als Lösung $a'_{11} = -\frac{2}{13}$, $a'_{12} = -\frac{5}{13}$, $a'_{13} = \frac{4}{13}$, $a'_{21} = -\frac{7}{13}$, $a'_{22} = -\frac{11}{13}$, $a'_{23} = \frac{1}{13}$, $a'_{31} = -\frac{6}{13}$, $a'_{32} = -\frac{15}{13}$, $a'_{33} = -\frac{1}{13}$. Die Inverse der Matrix $\mathbf{A}$ ist somit

$$\mathbf{A}^{-1} = \begin{pmatrix} -\frac{2}{13} & -\frac{5}{13} & \frac{4}{13} \\ -\frac{7}{13} & -\frac{11}{13} & \frac{1}{13} \\ -\frac{6}{13} & -\frac{15}{13} & -\frac{1}{13} \end{pmatrix}.$$

Streng genommen dürften wir bisher nur sagen, dass $\mathbf{A}^{-1}$ „rechtsinvers“ zu $\mathbf{A}$ ist, denn wir haben $\mathbf{A}^{-1}$ so berechnet, dass $\mathbf{A} \circ \mathbf{A}^{-1} = \mathbf{E}_3$ gilt, und die Matrizenmultiplikation ist i. Allg. nicht kommutativ (vgl. das Beispiel 6.4 auf S. 239). Durch Berechnung des Matrizenprodukts $\mathbf{A}^{-1} \circ \mathbf{A}$ stellt man aber fest, dass sich auch hierbei $\mathbf{E}_3$ ergibt. ■

Ohne Beweis sei angemerkt, dass sogar für beliebige Matrizen $\mathbf{A}$, falls eine Matrix $\mathbf{A}^{-1}$ mit $\mathbf{A} \circ \mathbf{A}^{-1} = \mathbf{E}_n$ existiert, für diese dann auch $\mathbf{A}^{-1} \circ \mathbf{A} = \mathbf{E}_n$ gilt.

In dem Beispiel 6.9 wurde bereits von *der* inversen Matrix $\mathbf{A}^{-1}$ gesprochen. Um den bestimmten Artikel verwenden zu dürfen, muss aber gesichert werden, dass eine Matrix *höchstens* eine Inverse besitzen kann und nicht etwa zwei verschiedene Matrizen $\mathbf{A}^{-1}$ und $(\mathbf{A}^{-1})'$ existieren, die beide invers zu einer Matrix $\mathbf{A}$ sind. Dies geschieht mit dem folgenden Satz.

Satz 6.8

Es seien $\mathbf{A}$, $\mathbf{B}$ und $\mathbf{C}$ $n \times n$ -Matrizen.

Gilt $\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A} = \mathbf{E}_n$ sowie $\mathbf{A} \circ \mathbf{C} = \mathbf{C} \circ \mathbf{A} = \mathbf{E}_n$, so ist $\mathbf{B} = \mathbf{C}$.

Beweis: Nach Voraussetzung und wegen der Assoziativität der Matrizenmultiplikation ist $\mathbf{C} = \mathbf{C} \circ \mathbf{E}_n = \mathbf{C} \circ (\mathbf{A} \circ \mathbf{B}) = (\mathbf{C} \circ \mathbf{A}) \circ \mathbf{B} = \mathbf{E}_n \circ \mathbf{B} = \mathbf{B}$. $\square$

Beispiel 6.10

Es wird die inverse Matrix $\mathbf{A}^{-1} = \begin{pmatrix} a'_{11} & a'_{12} & a'_{13} \\ a'_{21} & a'_{22} & a'_{23} \\ a'_{31} & a'_{32} & a'_{33} \end{pmatrix}$ zu $\mathbf{A} = \begin{pmatrix} 5 & -4 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix}$ gesucht. Wir stellen wie in dem Beispiel 6.9 durch Multiplikation $\mathbf{A} \circ \mathbf{A}^{-1}$ ein LGS auf, schreiben dieses aber jetzt Platz sparend als drei Teilsysteme:

$$\begin{array}{lll} 5a'_{11} - 4a'_{21} + 3a'_{31} = 1 & 5a'_{12} - 4a'_{22} + 3a'_{32} = 0 & 5a'_{13} - 4a'_{23} + 3a'_{33} = 0 \\ -a'_{11} + 2a'_{21} - 2a'_{31} = 0 & -a'_{12} + 2a'_{22} - 2a'_{32} = 1 & -a'_{13} + 2a'_{23} - 2a'_{33} = 0 \\ 3a'_{11} - a'_{31} = 0 & 3a'_{12} - a'_{32} = 0 & 3a'_{13} - a'_{33} = 1. \end{array}$$

Offensichtlich haben alle drei dieser Teil-LGS dieselbe einfache Koeffizientenmatrix und unterscheiden sich nur auf den rechten Seiten. Beim Lösen des ersten LGS ergibt sich bereits, dass dieses nicht lösbar ist:

$$\left(\begin{array}{ccc|c} 5 & -4 & 3 & 1 \\ -1 & 2 & -2 & 0 \\ 3 & 0 & -1 & 0 \end{array} \right) \begin{array}{l} | \cdot 1 \\ | \cdot 5 \\ | \cdot (-5) \end{array} \rightarrow \left(\begin{array}{ccc|c} 5 & -4 & 3 & 1 \\ 0 & 6 & -7 & 1 \\ 0 & -12 & 14 & 3 \end{array} \right) \begin{array}{l} | \cdot 2 \\ | \cdot 1 \end{array} \rightarrow \left(\begin{array}{ccc|c} 5 & -4 & 3 & 1 \\ 0 & 6 & -7 & 1 \\ 0 & 0 & 0 & 5 \end{array} \right).$$

Somit existiert keine zu $\mathbf{A}$ inverse Matrix. $\blacksquare$

Es ergeben sich aus den Beispielen 6.9 und 6.10 zwei *Fragen*:

- Wie lässt sich die Menge derjenigen Matrizen charakterisieren, die Inverse besitzen? Bei der Beantwortung dieser Frage werden sich interessante strukturelle Erkenntnisse über die Menge der invertierbaren Matrizen ergeben.
- Wie lassen sich inverse Matrizen möglichst einfach berechnen? (Mit der in dem Beispiel 6.9 praktizierten Vorgehensweise müssen für die Inversion einer $n \times n$ -Matrix ein LGS mit n^2 Gleichungen und n^2 Variablen bzw. n Gleichungssysteme mit jeweils n Gleichungen und n Variablen gelöst werden, was bei größeren Matrizen einen sehr hohen Rechenaufwand verursacht.)

Zur Beantwortung der ersten Frage betrachten wir erneut das Beispiel 6.10. Die dort betrachtete Matrix $\mathbf{A}$ besitzt keine Inverse, da ein LGS, welches $\mathbf{A}$ als einfache Koeffizientenmatrix besitzt, nicht lösbar ist und daher bei der Durchführung des Gauß-Algorithmus ein Widerspruch auftritt. Dieser Fall kann bei einer Matrix mit „vollem Rang“, d. h. mit linear unabhängigen Zeilen (sowie Spalten, siehe dazu den Abschnitt 6.1) nicht auftreten. So ist der Rang der in dem Beispiel 6.9 gegebenen Matrix 3. Die dabei entstehenden, aus drei Gleichungen mit

drei Variablen bestehenden, LGS sind in jedem Falle lösbar, da die Ränge ihrer einfachen Koeffizientenmatrizen mit den Rängen der erweiterten Koeffizientenmatrizen übereinstimmen, vgl. Abschnitt 1.3.3. Diese Erkenntnis ist verallgemeinerbar, wir fassen sie unter Verwendung des Begriffs der regulären Matrix (siehe die Definition 6.5 auf S. 234) zu dem folgenden Satz zusammen.

Satz 6.9

Eine Matrix $\mathbf{A}$ besitzt genau dann eine Inverse $\mathbf{A}^{-1}$, wenn sie regulär ist.

Aufgrund des Satzes 6.9 können die Begriffe *regulär* und *invertierbar* für Matrizen synonym verwendet werden. Wir verzichten auf einen Beweis dieses Satzes; für Plausibilitätsbetrachtungen dazu sei nochmals auf die Beispiele 6.9 und 6.10 verwiesen. Ebenfalls ohne Beweis wird der folgende Satz angegeben.

Satz 6.10

Sind $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ reguläre Matrizen, so ist auch $\mathbf{A} \circ \mathbf{B}$ regulär.

Aus dem Satz 6.10 ergibt sich die Folgerung, dass das Produkt $\mathbf{C} = \mathbf{A} \circ \mathbf{B}$ zweier invertierbarer Matrizen $\mathbf{A}, \mathbf{B}$ ebenfalls eine inverse Matrix $\mathbf{C}^{-1}$ besitzt.

Satz 6.11

Für beliebige $n \in \mathbb{N}$ ($n \geq 1$) bildet die Menge aller regulären $n \times n$ -Matrizen mit der Matrizenmultiplikation $\circ$ eine Gruppe.

Beweis: Zu dem Begriff der *Gruppe* siehe die Definition 3.4 auf S. 94. Die *Abgeschlossenheit* der Menge aller regulären $n \times n$ -Matrizen bezüglich $\circ$ ist durch den Satz 6.10, ihre *Assoziativität* durch den Satz 6.4 (S. 239) gegeben. Das *neutrale Element* dieser Menge ist die Einheitsmatrix $\mathbf{E}_n$. Schließlich sichert der Satz 6.9 die Existenz eines *inversen Elements* zu jedem Element der Menge aller regulären $n \times n$ -Matrizen. $\square$

Bemerkung: Im Gegensatz zu vielen bekannten Gruppen ist die Gruppe der regulären $n \times n$ -Matrizen mit der Matrizenmultiplikation $\circ$ *nicht kommutativ*.

Satz 6.12

Es seien $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ reguläre Matrizen. Dann gilt:

- a) $(\mathbf{A} \circ \mathbf{B})^{-1} = \mathbf{B}^{-1} \circ \mathbf{A}^{-1}$,
- b) $(\mathbf{A}^{-1})^{-1} = \mathbf{A}$,
- c) $(\mathbf{A}^T)^{-1} = (\mathbf{A}^{-1})^T$.

Beweis: Für den Beweis von a) ist zu zeigen, dass $\mathbf{B}^{-1} \circ \mathbf{A}^{-1}$ invers zu $\mathbf{A} \circ \mathbf{B}$ ist, also $(\mathbf{A} \circ \mathbf{B}) \circ (\mathbf{B}^{-1} \circ \mathbf{A}^{-1}) = \mathbf{E}_n$ gilt. Wegen der Assoziativität von $\circ$ gilt:
 $(\mathbf{A} \circ \mathbf{B}) \circ (\mathbf{B}^{-1} \circ \mathbf{A}^{-1}) = (\mathbf{A} \circ (\mathbf{B} \circ \mathbf{B}^{-1})) \circ \mathbf{A}^{-1} = (\mathbf{A} \circ \mathbf{E}_n) \circ \mathbf{A}^{-1} = \mathbf{A} \circ \mathbf{A}^{-1} = \mathbf{E}_n$
 und somit die Behauptung.

Zu b) überlegt man, dass nach Voraussetzung $\mathbf{A} \circ \mathbf{A}^{-1} = \mathbf{A}^{-1} \circ \mathbf{A} = \mathbf{E}_n$ gilt, $\mathbf{A}$ und $\mathbf{A}^{-1}$ also auch mit „vertauschten Rollen“ betrachtet werden können. Somit ist $\mathbf{A}$ die Inverse von $\mathbf{A}^{-1}$, also $\mathbf{A} = (\mathbf{A}^{-1})^{-1}$.

Der Beweis von c) sei der Leserin bzw. dem Leser überlassen, siehe die Aufgabe 10 auf S. 250. $\square$

Ein Verfahren zur Berechnung inverser Matrizen

Die Berechnung inverser Matrizen von 3×3 -Matrizen durch das Lösen drei linearer Gleichungssysteme mit drei Gleichungen und drei Variablen ist, wie die Beispiele 6.9 und 6.10 gezeigt haben, recht aufwändig – bei „größeren“ Matrizen wird dieses Vorgehen unpraktikabel. Wir demonstrieren daher im Folgenden anhand der Matrix aus dem Beispiel 6.9 ein etwas einfacheres Verfahren. Dazu schreiben wir die gegebene Matrix $\mathbf{A}$ und die Einheitsmatrix $\mathbf{E}_3$ nebeneinander und formen beide Matrizen nach dem Gauß-Jordan-Algorithmus so lange um, bis die ursprüngliche Matrix $\mathbf{A}$ die Gestalt der Einheitsmatrix angenommen hat.

$$\begin{array}{ccc|ccc|cl}
 2 & -5 & 3 & 1 & 0 & 0 & \cdot 1 & \cdot (-3) & 2 & -5 & 3 & 1 & 0 & 0 & & \\
 -1 & 2 & -2 & 0 & 1 & 0 & \cdot 2 & & 0 & -1 & -1 & 1 & 2 & 0 & \cdot 15 & \rightarrow \\
 3 & 0 & -1 & 0 & 0 & 1 & \cdot 2 & & 0 & 15 & -11 & -3 & 0 & 2 & \cdot 1 & \\
 \\
 2 & -5 & 3 & 1 & 0 & 0 & & \cdot 26 & 52 & -130 & 0 & 62 & 90 & 6 & \cdot 1 & \\
 0 & -1 & -1 & 1 & 2 & 0 & \cdot 26 & & 0 & -26 & 0 & 14 & 22 & -2 & \cdot (-5) & \rightarrow \\
 0 & 0 & -26 & 12 & 30 & 2 & \cdot (-1) & \cdot 3 & 0 & 0 & -26 & 12 & 30 & 2 & & \\
 \\
 52 & 0 & 0 & -8 & -20 & 16 & : 52 & & 1 & 0 & 0 & -\frac{2}{13} & -\frac{5}{13} & \frac{4}{13} & & \\
 0 & -26 & 0 & 14 & 22 & -2 & : (-26) & \rightarrow & 0 & 1 & 0 & -\frac{7}{13} & -\frac{11}{13} & \frac{1}{13} & & \\
 0 & 0 & -26 & 12 & 30 & 2 & : (-26) & & 0 & 0 & 1 & -\frac{6}{13} & -\frac{15}{13} & -\frac{1}{13} & &
 \end{array}$$

Nach dem Abschluss des Verfahrens steht links die Einheitsmatrix und rechts die Inverse der Ausgangsmatrix $\mathbf{A}$. Wir haben also erneut die inverse Matrix

$$\mathbf{A}^{-1} = \begin{pmatrix} -\frac{2}{13} & -\frac{5}{13} & \frac{4}{13} \\ -\frac{7}{13} & -\frac{11}{13} & \frac{1}{13} \\ -\frac{6}{13} & -\frac{15}{13} & -\frac{1}{13} \end{pmatrix}$$

erhalten. Obwohl der rechnerische Aufwand etwas geringer war als bei dem in dem Beispiel 6.9 durchgeführten Verfahren, nehmen die Umformungen auch bei diesem Vorgehen bereits bei 4×4 -Matrizen sehr viel Zeit in Anspruch, weshalb für die Bestimmung inverser Matrizen die Verwendung des Computers empfehlenswert ist, siehe hierzu den Abschnitt 6.4.

Matrizeninversion und Lösen linearer Gleichungssysteme

Multipliziert man eine $m \times n$ -Matrix $\mathbf{A}$ mit einem Spaltenvektor (d. h. einer einspaltigen Matrix) $\vec{x} \in \mathbb{R}^n$, so ergibt sich nach der Definition 6.6 (S. 238):

$$\mathbf{A} \circ \vec{x} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \circ \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix}.$$

Das Ergebnis hat die Gestalt der linken Seite eines linearen Gleichungssystems mit der einfachen Koeffizientenmatrix $\mathbf{A}$ (siehe Abschnitt 1.3.1):

$$\begin{array}{lcl}
 a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n & = & b_1 \\
 a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n & = & b_2 \\
 \vdots & & \vdots \\
 a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n & = & b_m
 \end{array} \quad \text{bzw.} \quad \left(\begin{array}{cccc|c}
 a_{11} & a_{12} & \dots & a_{1n} & b_1 \\
 a_{21} & a_{22} & \dots & a_{2n} & b_2 \\
 \vdots & \vdots & & \vdots & \vdots \\
 a_{m1} & a_{m2} & \dots & a_{mn} & b_m
 \end{array} \right).$$

Somit lässt sich ein lineares Gleichungssystem mit einer einfachen Koeffizientenmatrix $\mathbf{A}$ und einem Spaltenvektor $\vec{b}$ der absoluten Glieder auch in der Form

$$\mathbf{A} \circ \vec{x} = \vec{b}$$

schreiben. Ist $\mathbf{A}$ eine reguläre Matrix, so lässt sich ihre Inverse $\mathbf{A}^{-1}$ bestimmen. Durch Multiplizieren beider Seiten der obigen Gleichung mit $\mathbf{A}^{-1}$ ergibt sich

$$\vec{x} = \mathbf{A}^{-1} \circ \vec{b},$$

also eine direkte Darstellung des Lösungsvektors $\vec{x}$ durch $\mathbf{A}^{-1}$ und $\vec{b}$. Damit lassen sich Lösungen von LGS mit regulären Koeffizientenmatrizen ermitteln.

Beispiel 6.11

Das LGS

$$\begin{array}{rcl} 2x - 5y + 3z & = & 3 \\ -x + 2y - 2z & = & 4 \\ 3x & - & z = -7 \end{array} \quad \text{bzw.} \quad \mathbf{A} \circ \vec{x} = \begin{pmatrix} 2 & -5 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 3 \\ 4 \\ -7 \end{pmatrix}$$

löst man mithilfe der in dem Beispiel 6.9 bestimmten inversen Matrix $\mathbf{A}^{-1}$.

$$\vec{x} = \mathbf{A}^{-1} \circ \vec{b} = \begin{pmatrix} -\frac{2}{13} & -\frac{5}{13} & \frac{4}{13} \\ -\frac{7}{13} & -\frac{11}{13} & \frac{1}{13} \\ -\frac{6}{13} & -\frac{15}{13} & -\frac{1}{13} \end{pmatrix} \circ \begin{pmatrix} 3 \\ 4 \\ -7 \end{pmatrix} = \begin{pmatrix} -\frac{54}{13} \\ -\frac{72}{13} \\ -\frac{71}{13} \end{pmatrix}.$$

6.2.5 Aufgaben zu Abschnitt 6.2

1. Aus vier verschiedenen Rohstoffen R_1, R_2, R_3 und R_4 werden Zwischenprodukte Z_1, Z_2, Z_3 hergestellt, welche zu Endprodukten E_1, E_2, E_3 verarbeitet werden. Wie viele Mengeneinheiten der Rohstoffe zur Herstellung der jeweiligen Zwischenprodukte und wie viele Zwischenprodukte zur Herstellung der verschiedenen Endprodukte benötigt werden, geht aus der Abb. 6.3 hervor.
- Stellen Sie die Verflechtungsmatrizen $\mathbf{A}_{Z \rightarrow E}$ und $\mathbf{B}_{R \rightarrow Z}$ auf (siehe das Beispiel 6.2 auf S. 235f.).
 - Berechnen Sie für den Outputvektor $\vec{e} = \begin{pmatrix} 15 \\ 20 \\ 40 \end{pmatrix}$ den Bedarf an Zwischenprodukten (Zwischenvektor) und dann den Rohstoffbedarf (Inputvektor).
 - Ermitteln Sie die „Gesamtverflechtungsmatrix“ $\mathbf{C}_{R \rightarrow E}$, die den Rohstoffbedarf direkt in Abhängigkeit von dem Outputvektor angibt. Berechnen Sie mithilfe dieser Matrix erneut den Inputvektor und vergleichen Sie das Ergebnis mit dem, welches sie zuvor erhalten haben.

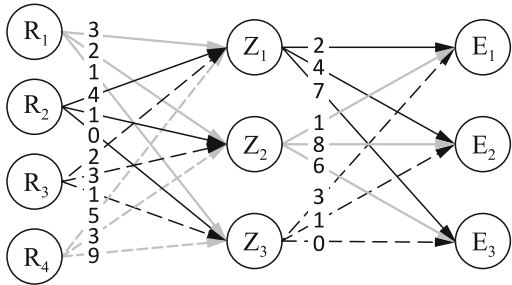


Abb. 6.3:
Materialverflechtung
(Schema zu Aufgabe 1)

2. Gegeben sind die folgenden Matrizen:

$$\mathbf{A} = \begin{pmatrix} 3 & 4 & 6 \\ 1 & -3 & 1 \\ 9 & 0 & -13 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 2 & -7 & 3 \\ -1 & 2 & 2 \\ 3 & 0 & -1 \end{pmatrix}, \quad \mathbf{C} = \begin{pmatrix} -1 & -9 & 0 & -3 \\ 0 & 5 & -5 & 11 \\ 3 & -2 & 0 & 6 \end{pmatrix}.$$

- Begründen Sie, dass sich die Produkte $\mathbf{C} \circ \mathbf{A}$ und $\mathbf{C} \circ \mathbf{B}$ nicht bilden lassen.
- Bilden Sie die Produkte $\mathbf{A} \circ \mathbf{B}$, $\mathbf{B} \circ \mathbf{A}$, $\mathbf{B} \circ \mathbf{C}$ und $\mathbf{A} \circ \mathbf{C}$ sowie $(\mathbf{A} \circ \mathbf{B}) \circ \mathbf{C}$ und $\mathbf{A} \circ (\mathbf{B} \circ \mathbf{C})$.

3. Prüfen Sie, ob für die folgenden Matrizen $\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A}$ gilt.

a) $\mathbf{A} = \begin{pmatrix} 3 & -4 \\ -4 & 3 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 2 & 7 \\ 7 & 2 \end{pmatrix}$

b) $\mathbf{A} = \begin{pmatrix} a & b \\ -b & a \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} c & d \\ -d & c \end{pmatrix} \quad (a, b, c, d \in \mathbb{R})$

4. Bestimmen Sie alle Matrizen $\mathbf{B}$, die mit der Matrix $\mathbf{A} = \begin{pmatrix} 2 & 1 \\ 3 & 4 \end{pmatrix}$ vertauschbar sind, d. h. für die $\mathbf{A} \circ \mathbf{B} = \mathbf{B} \circ \mathbf{A}$ gilt.

5. Bilden Sie die transponierten Matrizen $\mathbf{B}^T$ und $\mathbf{C}^T$ zu den in der Aufgabe 2 gegebenen Matrizen $\mathbf{B}$ und $\mathbf{C}$. Berechnen Sie dann $(\mathbf{B} \circ \mathbf{C})^T$ sowie $\mathbf{C}^T \circ \mathbf{B}^T$ und bestätigen Sie damit, dass der Satz 6.6 (S. 240) für dieses Beispiel zutrifft.

6. Beweisen Sie den Teil b des Satzes 6.7 auf S. 240.

7. Gegeben ist die folgende Situation:

Aus den Eiern eines Käfers schlüpfen nach einem Monat Larven. Nach einem weiteren Monat werden diese zu Käfern, die nach einem Monat jeweils 8 Eier legen und von denen dann 75% sterben und 25% noch einen Monat als „alte Käfer“ leben, dann noch vier Eier legen und sterben. Aus einem Viertel der Eier werden Larven, die anderen Eier werden gefressen oder verenden. Von den Larven werden vier Neuntel zu Käfern, fünf Neuntel sterben.

- Stellen Sie eine Populationsmatrix für diese Situation auf (siehe hierzu die Beispiele 6.7 und 6.8 auf S. 242ff.).
- Berechnen Sie für einen Anfangsbestand von 1000 Eiern, 1000 Larven, 1000 Käfern und 1000 alten Käfern die Bestände nach 1, 2, 3, 4 Monaten.
- Bestimmen Sie eine Anfangspopulation von q_E Eiern, q_L Larven, q_K Käfern und q_{aK} alten Käfern, deren Anzahlen konstant bleiben.

8. Berechnen Sie die Inversen der folgenden Matrizen:

$$\mathbf{A} = \begin{pmatrix} 3 & 5 \\ -3 & 2 \end{pmatrix}, \quad \mathbf{B} = \begin{pmatrix} 5 & -2 & 1 \\ 2 & -1 & 1 \\ -4 & 2 & -1 \end{pmatrix}, \quad \mathbf{C} = \begin{pmatrix} 2 & 1 & -2 & 0 \\ 2 & 0 & -1 & 4 \\ 4 & 2 & 0 & -3 \\ 0 & 1 & -3 & 2 \end{pmatrix}.$$

9. Gegeben ist eine 2×2 -Matrix $\mathbf{A} = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$. Geben Sie eine Bedingung (als Gleichung bzw. Ungleichung in den Koeffizienten a, b, c, d) dafür an, dass $\mathbf{A}$ regulär ist und berechnen Sie $\mathbf{A}^{-1}$.

10. Beweisen Sie den Teil c) des Satzes 6.12 auf S. 247.

Hinweis: Wenden Sie den Satz 6.6 auf S. 240 an.

6.3 Ein Ausblick auf Determinanten

Mithilfe von Determinanten („Bestimmenden“) ist es möglich, einige zentrale Eigenschaften von Matrizen durch einzelne Zahlen zum Ausdruck zu bringen. Sie geben zudem wichtige Eigenschaften von linearen Abbildungen (die im nächsten Kapitel behandelt werden) an, können für Berechnungen von verallgemeinerten Volumina genutzt werden und beschreiben Orientierungen.

Die Theorie der Determinanten ist ein grundlegender „Baustein“ weiter führender Gebiete der linearen Algebra. In diesem Buch kann nur eine kurze Einführung in Eigenschaften und Anwendungen von Determinanten gegeben werden. Wir beschränken uns daher auf Berechnungen recht einfacher Determinanten (wobei kompliziertere Determinanten leicht mithilfe des Computers berechnet werden können) und verzichten auf Beweise einer Reihe von Eigenschaften bzw. führen diese nur für die Determinanten „kleiner“ Matrizen.

Wir beginnen unsere Betrachtungen zu Determinanten mit der 2×2 -Matrix

$$\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}.$$

Diese Matrix ist genau dann regulär (und somit invertierbar), wenn sie den Rang 2 hat, d. h. wenn ihre Zeilenvektoren (und damit wegen des Satzes 6.3 auf S. 234 auch ihre Spaltenvektoren) linear unabhängig sind. Die Zeilenvektoren $\vec{a}_{Z1} = (a_{11}; a_{12})$ und $\vec{a}_{Z2} = (a_{21}; a_{22})$ von $\mathbf{A}$ sind genau dann linear abhängig, wenn $\vec{a}_{Z1} = (0; 0)$ ist oder $\lambda \in \mathbb{R}$ mit $\vec{a}_{Z2} = \lambda \vec{a}_{Z1}$ existiert. Letzteres bedeutet $a_{21} = \lambda a_{11}$ und $a_{22} = \lambda a_{12}$ bzw. (für $a_{11} \neq 0$ und $a_{12} \neq 0$) $\frac{a_{21}}{a_{11}} = \lambda$ und $\frac{a_{22}}{a_{12}} = \lambda$, somit also $\frac{a_{21}}{a_{11}} = \frac{a_{22}}{a_{12}}$ bzw. $a_{11} a_{22} = a_{12} a_{21}$. Auch für $a_{11} = 0$ führt diese Gleichung zur linearen Abhängigkeit der Zeilen- oder Spaltenvektoren, denn in diesem Falle würde daraus $a_{12} = 0$ oder $a_{21} = 0$ folgen, ein Zeilen- oder ein Spaltenvektor wäre also der Nullvektor. Ebenso lässt sich für den Fall $a_{12} = 0$ argumentieren und auch die Umkehrung, dass aus $a_{11} a_{22} \neq a_{12} a_{21}$ die lineare Unabhängigkeit der Zeilen- sowie der Spaltenvektoren resultiert, ergibt sich aus einfachen Schlüssen. Wir fassen die geführten Überlegungen zu einem Satz zusammen.

Satz 6.13

Eine 2×2 -Matrix $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ ist genau dann regulär, wenn

$$a_{11} a_{22} - a_{12} a_{21} \neq 0$$

gilt.

Der Satz 6.13 besagt nicht weniger, als dass sich die Regularität (und somit Invertierbarkeit) einer 2×2 -Matrix durch Berechnung einer einzigen Zahl ermitteln lässt, die somit eine „bestimmende“ Bedeutung für die Matrix hat.

Definition 6.8

Als Determinante einer 2×2 -Matrix $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ bezeichnet man die Zahl

$$\det \mathbf{A} = a_{11} a_{22} - a_{12} a_{21}.$$

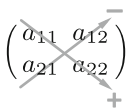


Abb. 6.4:

Merkhilfe zur Berechnung der Determinanten von 2×2 -Matrizen

Wir formulieren im Folgenden einige Eigenschaften von Determinanten, die nicht nur für die Determinanten von 2×2 -Matrizen, sondern allgemein für Determinanten beliebiger quadratischer Matrizen gelten.

Eigenschaften von Determinanten

1. Die Determinante einer Matrix $\mathbf{A}$ ist genau dann Null, wenn
 - $\mathbf{A}$ nicht regulär ist,
 - $\text{rg } \mathbf{A} < n$ gilt (wobei n die Zeilen- und Spaltenanzahl von $\mathbf{A}$ ist),
 - die Zeilenvektoren von $\mathbf{A}$ linear abhängig sind,
 - die Spaltenvektoren von $\mathbf{A}$ linear abhängig sind.

Dies ist gleichbedeutend damit, dass jede quadratische Matrix $\mathbf{A}$ mit $\det \mathbf{A} \neq 0$ regulär ist und somit den „vollen Rang“ n , also linear unabhängige Zeilenvektoren sowie Spaltenvektoren besitzt.
2. Falls eine Matrix $\mathbf{A}$ eine Zeile oder eine Spalte besitzt, die nur aus Nullen besteht, so ist $\det \mathbf{A} = 0$ (Spezialfall von 1.).
3. Ist $\mathbf{A}$ eine quadratische Matrix und $\mathbf{A}^T$ ihre Transponierte, so gilt $\det \mathbf{A}^T = \det \mathbf{A}$.
4. Vertauscht man zwei Zeilen einer Matrix, so wechselt die Determinante ihr Vorzeichen. Dasselbe gilt beim Vertauschen zweier Spalten.
5. Entsteht eine Matrix $\mathbf{A}'$ aus einer Matrix $\mathbf{A}$ durch *Multiplikation einer Zeile oder einer Spalte* von $\mathbf{A}$ mit $\lambda \in \mathbb{R}$, so gilt $\det \mathbf{A}' = \lambda \det \mathbf{A}$.
6. Addiert man ein beliebiges Vielfaches einer Zeile/Spalte einer Matrix $\mathbf{A}$ zu einer anderen Zeile/Spalte von $\mathbf{A}$, so bleibt die Determinante gleich.

Wir werden diese Eigenschaften nun für 2×2 -Matrizen bestätigen. Mit dem Satz 6.13 wurde bereits die Gültigkeit der Eigenschaften 1. und 2. für 2×2 -Matrizen gezeigt, es bleibt also die Gültigkeit von 3.-6. zu bestätigen.

3. Für $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ ist $\mathbf{A}^T = \begin{pmatrix} a'_{11} & a'_{12} \\ a'_{21} & a'_{22} \end{pmatrix} = \begin{pmatrix} a_{11} & a_{21} \\ a_{12} & a_{22} \end{pmatrix}$. Somit gilt

$$\det \mathbf{A}^T = a'_{11}a'_{22} - a'_{12}a'_{21} = a_{11}a_{22} - a_{21}a_{12} = \det \mathbf{A}.$$
4. Durch Zeilenvertauschung entsteht $\mathbf{A}' = \begin{pmatrix} a'_{11} & a'_{12} \\ a'_{21} & a'_{22} \end{pmatrix} = \begin{pmatrix} a_{21} & a_{22} \\ a_{11} & a_{12} \end{pmatrix}$. Es gilt

$$\det \mathbf{A}' = a'_{11}a'_{22} - a'_{12}a'_{21} = a_{21}a_{12} - a_{22}a_{11} = -(a_{11}a_{22} - a_{12}a_{21}) = -\det \mathbf{A}.$$
5. Für $\mathbf{A}' = \begin{pmatrix} \lambda a_{11} & \lambda a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ ist $\det \mathbf{A}' = (\lambda a_{11})a_{22} - (\lambda a_{12})a_{21} = \lambda \det \mathbf{A}$.
Dasselbe Ergebnis erhält man bei Multiplikation der zweiten Zeile mit $\lambda \in \mathbb{R}$.
6. Durch Addieren des λ -fachen der ersten Zeile zur zweiten Zeile entsteht die Matrix $\mathbf{A}' = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} + \lambda a_{11} & a_{22} + \lambda a_{12} \end{pmatrix}$. Ihre Determinante ist

$$\begin{aligned} \det \mathbf{A}' &= a_{11}(a_{22} + \lambda a_{12}) - a_{12}(a_{21} + \lambda a_{11}) \\ &= a_{11}a_{22} + \lambda a_{11}a_{12} - a_{12}a_{21} - \lambda a_{12}a_{11} = a_{11}a_{22} - a_{12}a_{21} \\ &= \det \mathbf{A}. \end{aligned}$$

Es wurde darauf verzichtet, die Eigenschaften 4.-6. auch für die entsprechenden Spaltenmanipulationen zu bestätigen, siehe hierzu die Aufgabe 1 auf S. 255.

Eine interessante Folgerung ergibt sich aus der Eigenschaft 5. für die Determinante der Matrix $\lambda \cdot \mathbf{A}$, die durch die Multiplikation einer Matrix $\mathbf{A}$ mit einer reellen Zahl λ entsteht. Dabei werden alle Koeffizienten, bei einer 2×2 -Matrix also beide Zeilen, mit λ multipliziert. Somit wird die Determinante der Ausgangsmatrix $\mathbf{A}$ zweifach mit λ multipliziert; es gilt daher $\det(\lambda \cdot \mathbf{A}) = \lambda^2 \det \mathbf{A}$. Wir überprüfen dies noch anhand der Definition 6.8:

$$\det(\lambda \cdot \mathbf{A}) = \det \begin{pmatrix} \lambda a_{11} & \lambda a_{12} \\ \lambda a_{21} & \lambda a_{22} \end{pmatrix} = \lambda a_{11} \lambda a_{22} - \lambda a_{12} \lambda a_{21} = \lambda^2 \det \mathbf{A}.$$

Für $n \times n$ -Matrizen folgt aus der Eigenschaft 5. der folgende Satz.

Satz 6.14

Für eine beliebige $n \times n$ -Matrix $\mathbf{A}$ und $\lambda \in \mathbb{R}$ gilt $\det(\lambda \cdot \mathbf{A}) = \lambda^n \det \mathbf{A}$.

Determinanten von 3×3 -Matrizen

Definition 6.9

Als Determinante einer 3×3 -Matrix $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$ wird die Zahl

$\det \mathbf{A} = a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{13}a_{22}a_{31} - a_{12}a_{21}a_{33} - a_{11}a_{23}a_{32}$ bezeichnet. ◆

Bemerkung: Eine Merkhilfe für diese Berechnungsvorschrift gibt die Abb. 6.5. Man schreibt die ersten beiden Spalten der Matrix rechts neben die Matrix und bildet Produkte von je 3 Koeffizienten, die auf beiden Diagonalen der Matrix und ihren Parallelen liegen. Die nach rechts unten verlaufenden Produkte werden addiert und die nach rechts oben verlaufenden Produkte davon subtrahiert.

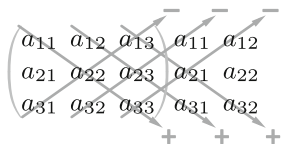


Abb. 6.5: Merkhilfe zur Berechnung der Determinanten von 3×3 -Matrizen (Regel von Sarrus)

Beispiel 6.12

Die Matrix $\mathbf{A} = \begin{pmatrix} 2 & -5 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix}$ (siehe Beispiel 6.9, S. 245) hat die Determinante

$$\det \mathbf{A} = 2 \cdot 2 \cdot (-1) + (-5) \cdot (-2) \cdot 3 + 0 - 3 \cdot 2 \cdot 3 - (-1) \cdot (-1) \cdot (-5) - 0 = 13. \quad \blacksquare$$

Beispiel 6.13

Die Matrix $\mathbf{A} = \begin{pmatrix} 5 & -4 & 3 \\ -1 & 2 & -2 \\ 3 & 0 & -1 \end{pmatrix}$ ist nicht regulär, siehe Beispiel 6.10 auf S. 246.

Ihre Determinante ist

$$\det \mathbf{A} = 5 \cdot 2 \cdot (-1) + (-4) \cdot (-2) \cdot 3 + 0 - 3 \cdot 2 \cdot 3 - (-1) \cdot (-1) \cdot (-4) - 0 = 0. \quad \blacksquare$$

Wir vermerken, dass für die nach der Definition 6.9 bestimmten Determinanten von 3×3 -Matrizen die auf S. 252 aufgeführten Eigenschaften 1.-6. gelten, verzichten aber auf Beweise.

Beispiel 6.14

Wir betrachten die Matrix $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ 0 & a_{22} & a_{23} \\ 0 & 0 & a_{33} \end{pmatrix}$, die unterhalb der Hauptdiagonalen nur Nullen enthält und deshalb als *obere Dreiecksmatrix* bezeichnet wird. Die Determinante dieser Matrix ist

$$\det \mathbf{A} = a_{11}a_{22}a_{33} + a_{12}a_{23} \cdot 0 + a_{13} \cdot 0 \cdot 0 - a_{13}a_{22} \cdot 0 - a_{12} \cdot 0 \cdot a_{33} - a_{11}a_{23} \cdot 0 \\ = a_{11}a_{22}a_{33},$$

also das Produkt der Komponenten der Hauptdiagonalen. Dies gilt natürlich insbesondere für die Diagonalmatrix $\begin{pmatrix} a_{11} & 0 & 0 \\ 0 & a_{22} & 0 \\ 0 & 0 & a_{33} \end{pmatrix}$, die ja eine besondere obere

Dreiecksmatrix ist. Als noch spezielleren Fall erhält man für die Determinante der dreizeiligen und dreispaltigen Einheitsmatrix $\det \mathbf{E}_3 = 1$. ■

Ohne Beweis vermerken wir, dass die in dem Beispiel 6.14 für 3×3 -Matrizen gemachten Feststellungen auch für beliebige $n \times n$ -Matrizen gelten:

Satz 6.15

Die Determinante der oberen Dreiecksmatrix $\mathbf{A} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \ddots & \vdots \\ \vdots & \ddots & \ddots & a_{n-1,n} \\ 0 & \dots & 0 & a_{nn} \end{pmatrix}$ ist das

Produkt der Komponenten ihrer Hauptdiagonalen: $\det \mathbf{A} = a_{11} \cdot a_{22} \cdot \dots \cdot a_{nn}$.

Speziell ist $\det \mathbf{E}_n = \det \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & 1 \end{pmatrix} = 1$.

Ebenfalls ohne Beweis seien noch zwei weitere zentrale Eigenschaften von Determinanten angegeben.

Satz 6.16

- a) Die Determinante des Produkts zweier Matrizen $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times n}$ ist gleich dem Produkt der Determinanten beider Matrizen: $\det(\mathbf{A} \circ \mathbf{B}) = \det \mathbf{A} \cdot \det \mathbf{B}$.
- b) Für die Determinante der Inversen einer Matrix $\mathbf{A}$ gilt: $\det \mathbf{A}^{-1} = \frac{1}{\det \mathbf{A}}$.

Determinanten und Volumina bzw. Flächeninhalte

In dem Abschnitt 3.6 (S. 133ff.) wurde das Vektorprodukt $\vec{u} \times \vec{v}$ zweier Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^3$ mit der Festlegung eingeführt, dass der Betrag von $\vec{u} \times \vec{v}$ gleich dem Flächeninhalt des von $\vec{u}$ und $\vec{v}$ aufgespannten Parallelogramms ist. Betrachten wir zwei Vektoren $\vec{u} = \begin{pmatrix} x_u \\ y_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \end{pmatrix}$ von $\mathbb{R}^2$, so können wir diese durch

jeweilige Hinzunahme der Null als z -Komponente zu Vektoren $\begin{pmatrix} x_u \\ y_u \\ 0 \end{pmatrix}, \begin{pmatrix} x_v \\ y_v \\ 0 \end{pmatrix} \in \mathbb{R}^3$ „erweitern“ und ihr Vektorprodukt berechnen (siehe den Satz 3.21 auf S. 135):

$$\begin{pmatrix} x_u \\ y_u \\ 0 \end{pmatrix} \times \begin{pmatrix} x_v \\ y_v \\ 0 \end{pmatrix} = \begin{pmatrix} y_u \cdot 0 - 0 \cdot y_v \\ 0 \cdot x_v - x_u \cdot 0 \\ x_u y_v - y_u x_v \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ x_u y_v - y_u x_v \end{pmatrix}.$$

Der Betrag dieses Vektorproduktes ist $\sqrt{(x_u y_v - y_u x_v)^2} = |x_u y_v - y_u x_v|$ und somit gleich dem Absolutbetrag der Determinante der Matrix $\begin{pmatrix} x_u & x_v \\ y_u & y_v \end{pmatrix}$.

Satz 6.17

Der Flächeninhalt des von zwei Vektoren $\vec{u} = \begin{pmatrix} x_u \\ y_u \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} x_v \\ y_v \end{pmatrix}$ aufgespannten Parallelogramms ist gleich dem Betrag der Determinante der Matrix $\begin{pmatrix} x_u & x_v \\ y_u & y_v \end{pmatrix}$ mit den Spaltenvektoren $\vec{u}$ und $\vec{v}$.

Um zu einer analogen Aussage über Volumina von Spaten (Parallelepipeden) zu gelangen, erinnern wir uns an den Abschnitt 3.6.2 (S. 136) über das Spatprodukt $(\vec{u} \times \vec{v}) \cdot \vec{w}$ dreier Vektoren $\vec{u}, \vec{v}, \vec{w} \in \mathbb{R}^3$. Es gilt

$$\begin{aligned} (\vec{u} \times \vec{v}) \cdot \vec{w} &= \left(\begin{pmatrix} x_u \\ y_u \\ z_u \end{pmatrix} \times \begin{pmatrix} x_v \\ y_v \\ z_v \end{pmatrix} \right) \cdot \begin{pmatrix} x_w \\ y_w \\ z_w \end{pmatrix} = \begin{pmatrix} y_u z_v - z_u y_v \\ z_u x_v - x_u z_v \\ x_u y_v - y_u x_v \end{pmatrix} \cdot \begin{pmatrix} x_w \\ y_w \\ z_w \end{pmatrix} \\ &= (y_u z_v - z_u y_v) \cdot x_w + (z_u x_v - x_u z_v) \cdot y_w + (x_u y_v - y_u x_v) \cdot z_w \\ &= y_u z_v x_w - z_u y_v x_w + z_u x_v y_w - x_u z_v y_w + x_u y_v z_w - y_u x_v z_w \\ &= x_u y_v z_w + x_v y_w z_u + x_w y_u z_v - z_u y_v x_w - z_v y_w x_u - z_w y_u x_v \\ &= \det \begin{pmatrix} x_u & x_v & x_w \\ y_u & y_v & y_w \\ z_u & z_v & z_w \end{pmatrix}. \end{aligned}$$

Satz 6.18

Das Spatprodukt dreier Vektoren $\vec{u}, \vec{v}, \vec{w} \in \mathbb{R}^3$ ist gleich der Determinante der 3×3 -Matrix mit den Spaltenvektoren $\vec{u}, \vec{v}$ und $\vec{w}$. Das Volumen des von $\vec{u}, \vec{v}$ und $\vec{w}$ aufgespannten Spates ist gleich dem Absolutbetrag dieser Determinante.

Der Begriff des Volumens ist auf n -dimensionale Räume erweiterbar. Durch n Vektoren $\vec{u}_1, \dots, \vec{u}_n$ werden verallgemeinerte Spate bzw. *Parallelotope* aufgespannt, deren Volumina sich mithilfe der Determinanten der aus den Spaltenvektoren $\vec{u}_1, \dots, \vec{u}_n$ gebildeten Matrizen berechnen lassen.

Der Begriff der Orientierung

In dem Abschnitt 3.6 wurde bereits die Orientierung eines geordneten Tripels von Vektoren erwähnt, wobei noch keine exakte Definition angegeben werden konnte. Dies ist mithilfe von Determinanten nun möglich.

Definition 6.10

Ein geordnetes Paar $(\vec{u}, \vec{v})$ von Vektoren $\vec{u}, \vec{v} \in \mathbb{R}^2$ heißt *positiv orientiert*, *Rechtssystem* bzw. *gleichorientiert* zu der Standardbasis von $\mathbb{R}^2$, falls die Determinante der Matrix $\begin{pmatrix} x_u & x_v \\ y_u & y_v \end{pmatrix}$ mit den Spaltenvektoren $\vec{u}$ und $\vec{v}$ positiv ist.

Ein geordnetes Tripel $(\vec{u}, \vec{v}, \vec{w})$ von Vektoren $\vec{u}, \vec{v}, \vec{w} \in \mathbb{R}^3$ heißt *Rechtssystem* bzw. *gleichorientiert* zu der Standardbasis von $\mathbb{R}^3$, falls die Determinante der von $\vec{u}, \vec{v}$ und $\vec{w}$ (in dieser Reihenfolge) gebildeten Matrix positiv ist. $\blacklozenge$

6.3.1 Aufgaben zu Abschnitt 6.3

1. Zeigen Sie, dass die auf S. 252 aufgeführten Eigenschaften 4.-6. von Determinanten auch für die entsprechenden *Spaltenmanipulationen* gelten.
2. Berechnen Sie die Determinanten der in der Aufgabe 2 auf S. 250 gegebenen Matrizen **A** und **B**.
3. Zeigen Sie für eine beliebige 3×3 -Matrix **A** und $\lambda \in \mathbb{R}$: $\det(\lambda \cdot \mathbf{A}) = \lambda^3 \det \mathbf{A}$.

6.4 Matrizenrechnung mithilfe des Computers

Mithilfe des CAS Maxima (siehe dazu auch die Abschnitte 1.5 und 3.7) lassen sich alle in diesem Kapitel beschriebenen Berechnungen mit Matrizen ausführen.

Eingabe von Matrizen

Matrizen werden durch `matrix([Zeile 1], [Zeile 2], ..., [Zeile m]);` eingegeben, wobei [Zeile 1],... Zeilenvektoren sind.

Beispiel: `A:matrix([1,3,5], [0,2,3], [1,-3,-4], [4,6,11]);`

Rangbestimmung

Der Rang einer Matrix **A** wird (nach deren Eingabe) mittels

`rank(A);`

berechnet. Es gibt noch andere Möglichkeiten der Rangbestimmung, z. B. durch die Überprüfung der linearen Abhängigkeit von Zeilen- und Spaltenvektoren. Beispiele enthält eine Maxima-Datei auf der Internetseite zu diesem Buch. Diese sind aber komplizierter als die einfache Rangberechnung mittels `rank`.

Matrizenmultiplikation

Für das Produkt zweier Matrizen (bzw. einer Matrix und eines Spaltenvektors) wird (wie für das Skalarprodukt von Vektoren) ein Punkt `.` gesetzt.

Beispiele:

```
A:matrix([9,3,8,2],[5,1,8,1]);    C:matrix([3,2,8],[1,0,9],[3,2,1]);
B:matrix([4,2],[5,7],[3,3]);      x:matrix([3],[2],[-7]);
B.A;                               C.x;
```

Bildung transponierter Matrizen

Nach vorheriger Eingabe einer Matrix **A** wird ihre Transponierte durch

`transpose(A);`

gebildet.

Berechnung inverser Matrizen

Nach Eingabe einer regulären Matrix **A** wird ihre Inverse durch

`invert(A);`

berechnet. Bei einer nicht regulären Matrix erscheint die Meldung

`Division by 0 -- an error`

Die Lösung eines LGS mit einer regulären einfachen Koeffizientenmatrix **A** und dem Spaltenvektor $\vec{b}$ der absoluten Glieder ist $\vec{x} = \mathbf{A}^{-1} \circ \vec{b}$ (siehe S. 248):

```
A:matrix([2,-5,3],[-1,2,-2],[3,0,-1]);
b:matrix([3],[4],[-7]);
invert(A).b;
```

Berechnung von Determinanten

Die Berechnung der Determinante einer Matrix **A** erfolgt durch

`determinant(A);`

Beispiele zu allen hier angegebenen Berechnungen sind in einer Maxima-Datei enthalten, die auf der Internetseite zu diesem Buch zur Verfügung steht.

7 Lineare und affine Abbildungen

Übersicht

7.1	Abbildungen: Definition und einige Beispiele	258
7.2	Lineare Abbildungen	264
7.3	Affine Abbildungen	279
7.4	Kongruenz- und Ähnlichkeitsabbildungen.....	287

In diesem Kapitel treffen geometrische Abbildungen wie Verschiebungen, Drehungen, Spiegelungen und Streckungen sowie Strukturen der linearen Algebra aufeinander. Dabei werden sich Vektoren und Matrizen als äußerst leistungsfähige Hilfsmittel zur Beschreibung und Untersuchung von Abbildungen erweisen.

Bei der Behandlung von Abbildungen mithilfe der linearen Algebra ist zwischen Punkt- und Vektorabbildungen zu unterscheiden. Nach der exemplarischen Betrachtung einiger einfacher geometrischer Abbildungen als Punktabbildungen (Abschnitt 7.1) und der Herstellung von Bezügen zu entsprechenden Vektorabbildungen werden spezielle Vektorabbildungen (lineare Abbildungen) näher untersucht. Diese bilden dann die Grundlage für die Behandlung wichtiger Punktabbildungen, der affinen Abbildungen. Der letzte Abschnitt befasst sich mit besonderen Abbildungen: den Isometrien als Vektorabbildungen und auf deren Grundlage dann mit Kongruenz- und Ähnlichkeitsabbildungen.

Das vorliegende Kapitel weist enge Bezüge zur Schulgeometrie auf – die im Schulunterricht auftretenden geometrischen Abbildungen werden von einem höheren Standpunkt aus betrachtet. Es werden – ohne Anspruch auf Vollständigkeit – u. a. Nacheinanderausführungen von Abbildungen betrachtet und Fixpunkte bestimmt. Dabei wird exemplarisch verdeutlicht, wie die Untersuchung geometrischer Abbildungen mit den Mitteln der Vektor- und Matrizenrechnung auf elegante und übersichtliche Weise vorgenommen werden kann.

Analytische Beschreibungen von Abbildungen bilden die Grundlage für ihre graphische Darstellung mithilfe des Computers. Die Verwendung des Computers zur Veranschaulichung wird wärmstens empfohlen. Auf der Internetseite des Buches stehen Dateien zur Verfügung, unter deren Nutzung dies leicht mit dem Computeralgebrasystem Maxima erfolgen kann.

7.1 Abbildungen: Definition und einige Beispiele

7.1.1 Der Begriff „Abbildung“, Eigenschaften von Abbildungen

Vor der Behandlung linearer Abbildungen werden hier kurz der allgemeine Abbildungsbegriff und einige Eigenschaften von Abbildungen zusammengestellt.

Definition 7.1

Es seien A und B zwei Mengen. Eine *Abbildung* bzw. *Funktion* f von A in B ist eine Zuordnung,¹ bei der jedem Element x von A *genau* ein Element von B zugeordnet wird, man bezeichnet dieses mit $f(x)$.

- Für „ f ist eine Abbildung von A in B “ schreibt man kurz: $f: A \rightarrow B$.
- Die Zuordnung eines einzelnen Elements $x \in A$ zu seinem „Bildelement“ $f(x)$ kennzeichnet man durch einen etwas anderen Pfeil: $x \mapsto f(x)$. ♦

Bemerkung: Die Forderung, dass jedem $x \in A$ *genau* ein Element von B zugeordnet wird, beinhaltet, dass *kein* Element von A *zwei* verschiedene Bildelemente haben kann, d. h. es muss für beliebige $x \in A$ gelten:

Wenn $f(x) = y_1$ und $f(x) = y_2$, so ist $y_1 = y_2$.

Beispiele: Die durch Pfeile dargestellte Zuordnung f_1 in der Abb. 7.1 ist keine Abbildung, da dem Element $a_1 \in A$ zwei verschiedene Elemente $b_1, b_2 \in B$ zugeordnet werden. Hingegen handelt es sich bei f_2, f_3 und f_4 um Abbildungen.

Definition 7.2

Es sei $f: A \rightarrow B$ eine Abbildung von einer Menge A in eine Menge B .

- f heißt *injektiv* (bzw. *eindeutig*), wenn keine zwei Elemente in A existieren, die auf dasselbe Element von B abgebildet werden, d. h. wenn für beliebige $a_1, a_2 \in A$ aus $f(a_1) = f(a_2)$ stets $a_1 = a_2$ folgt.
- f heißt *surjektiv*, wenn jedes Element von B mindestens ein Urbild in A besitzt, d. h. wenn zu jedem Element $b \in B$ ein $a \in A$ mit $f(a) = b$ existiert.
- f heißt *bijektiv*, wenn f injektiv und surjektiv ist. ♦

Beispiele: Die in der Abb. 7.1 dargestellte Abbildung f_2 ist weder injektiv noch surjektiv; f_3 ist injektiv aber nicht surjektiv, da b_6 kein Urbild besitzt. f_4 ist injektiv und surjektiv, also eine bijektive Abbildung.

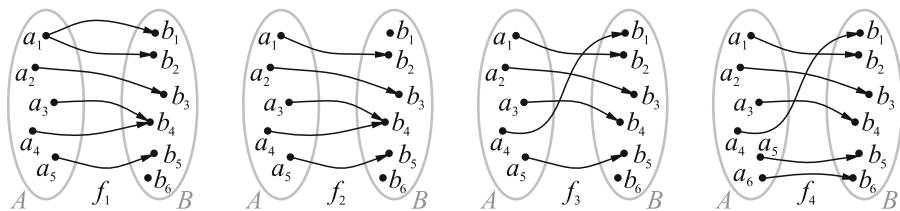


Abb. 7.1: Zuordnungen und Abbildungen

¹Der Begriff der Zuordnung wird hier in einem umgangssprachlichen Sinne verwendet. Für eine tiefer gehende mengentheoretische Fundierung des Abbildungs- bzw. Funktionsbegriffs siehe Lehmann/Schulz (2007), S. 86ff.

7.1.2 Koordinatenbeschreibungen geometrischer Abbildungen

Es werden im Folgenden einige bekannte Punktabbildungen der Ebene sowie des Raumes auf eine Ebene durch Koordinaten beschrieben.

Beispiel 7.1

Die Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} -4 \\ 1 \end{pmatrix}$ für alle $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ ist eine *Verschiebung* mit dem Verschiebungsvektor $\vec{v} = \begin{pmatrix} -4 \\ 1 \end{pmatrix}$. In der Abb. 7.2 sind die Eckpunkte $P = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$, $Q = \begin{pmatrix} 5 \\ 1 \end{pmatrix}$, $R = \begin{pmatrix} 5 \\ 3 \end{pmatrix}$ und $S = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$ eines Rechtecks sowie ihre Bildpunkte $P' = f(P)$, $Q' = f(Q)$, $R' = f(R)$ und $S' = f(S)$ dargestellt. ■

Bemerkungen:

- Punktabbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden immer *alle* Punkte von $\mathbb{R}^2$ auf Punkte von $\mathbb{R}^2$ ab. Dies lässt sich natürlich nicht graphisch darstellen, weshalb in zeichnerischen Darstellungen wie Abb. 7.2 und Abb. 7.3 stellvertretend einige Punkte, ihre Bildpunkte sowie ggf. Verbindungsstrecken zwischen Punkten und zwischen Bildpunkten dargestellt werden.
- Wie in dem Abschnitt 5.5 ausgeführt wurde, wird $\mathbb{R}^2$ als Punkt- und zugleich als zugehöriger Vektorraum betrachtet. Punkte und Vektoren werden dabei in gleicher Weise in Spaltenform dargestellt, siehe die Anmerkung „Komponenten- und Koordinatenschreibweise“ auf S. 206. In dem Beispiel 7.1 bezeichnet $\begin{pmatrix} x \\ y \end{pmatrix}$ einen Punkt, $\begin{pmatrix} -4 \\ 1 \end{pmatrix}$ hingegen einen Vektor.
- Statt $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} -4 \\ 1 \end{pmatrix}$ schreibt man auch $f\left(\begin{pmatrix} x \\ y \end{pmatrix}\right) = \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} -4 \\ 1 \end{pmatrix}$ bzw. $f(P) = P + \vec{v}$. Bildpunkte werden oft mit Strichen bezeichnet, z. B. $P' = f(P)$.
- Die in dem Beispiel 7.1 beschriebene Verschiebung ist eine *affine Abbildung*, näher wird auf diesen Begriff in dem Abschnitt 7.3 eingegangen.

Beispiel 7.2

Die Abbildung $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} -x \\ y \end{pmatrix}$ für alle $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ ist eine *Spiegelung* an der y -Achse (siehe Abb. 7.3). ■

Beispiel 7.3

Wir führen die Abbildungen g und f aus den Beispielen 7.1 und 7.2 nacheinander aus. Führt man zuerst f und dann g aus, so wird die resultierende „zusammengesetzte“ Abbildung mit $g \circ f$ bezeichnet (man beachte die Reihenfolge).

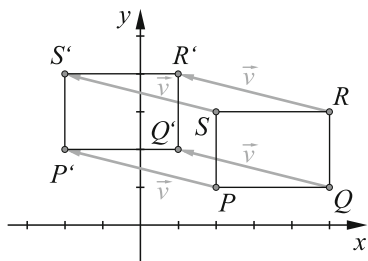


Abb. 7.2: Verschiebung (Beispiel 7.1)

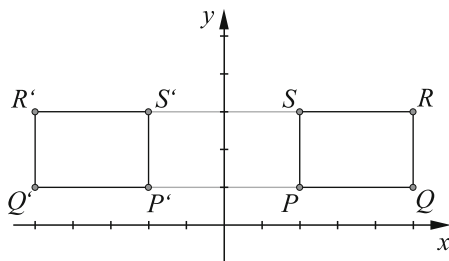


Abb. 7.3: Spiegelung (Beispiel 7.2)

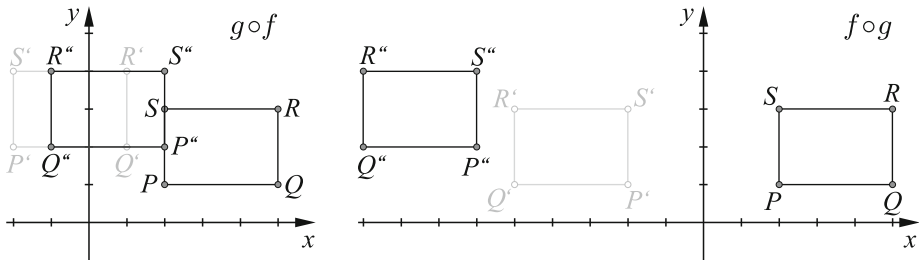


Abb. 7.4: Nacheinanderausführungen einer Verschiebung f und einer Spiegelung g

Wir erhalten in diesem Falle

$$f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} x-4 \\ y+1 \end{pmatrix}, \quad g: \begin{pmatrix} x' \\ y' \end{pmatrix} \mapsto \begin{pmatrix} x'' \\ y'' \end{pmatrix} = \begin{pmatrix} -x' \\ y' \end{pmatrix} = \begin{pmatrix} -x+4 \\ y+1 \end{pmatrix}$$

Bei der umgekehrten Reihenfolge der Nacheinanderausführung (wenn man $f \circ g$ betrachtet, also zuerst g und dann f ausführt) ergibt sich

$$g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} -x \\ y \end{pmatrix}, \quad f: \begin{pmatrix} x' \\ y' \end{pmatrix} \mapsto \begin{pmatrix} x'' \\ y'' \end{pmatrix} = \begin{pmatrix} x'-4 \\ y'+1 \end{pmatrix} = \begin{pmatrix} -x-4 \\ y+1 \end{pmatrix}.$$

Es ist also $g \circ f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} -x+4 \\ y+1 \end{pmatrix}$ und $f \circ g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} -x-4 \\ y+1 \end{pmatrix}$; d. h. $g \circ f \neq f \circ g$, siehe auch Abb. 7.4, in der die in dem Beispiel 7.1 angegebenen Punkte und ihre Bildpunkte dargestellt sind. Es zeigt sich somit, dass die *Nacheinanderausführung von Abbildungen im Allgemeinen nicht kommutativ* ist. ■

Beispiel 7.4

Wir überlegen im Folgenden, wie eine *Drehung* von $\mathbb{R}^2$ um den Koordinatenursprung mit einem Drehwinkel α beschrieben werden kann. Dazu betrachten wir einen Punkt $P = \begin{pmatrix} x \\ y \end{pmatrix}$ und seinen Bildpunkt $P' = \begin{pmatrix} x' \\ y' \end{pmatrix}$, siehe Abb. 7.5 a). Es ist $x' = |OP'| \cos(\varphi + \alpha) = |OP| \cos(\varphi + \alpha)$, $y' = |OP'| \sin(\varphi + \alpha) = |OP| \sin(\varphi + \alpha)$. Durch Anwendung der Additionstheoreme der Kosinus- und der Sinusfunktion ergibt sich daraus:

$$x' = |OP| (\cos \varphi \cos \alpha - \sin \varphi \sin \alpha), \quad y' = |OP| (\sin \varphi \cos \alpha + \cos \varphi \sin \alpha).$$

In diesen Gleichungen lassen sich $\cos \varphi$, $\sin \varphi$ und gleichzeitig $|OP|$ ersetzen, denn es gilt $x = |OP| \cos \varphi$ und $y = |OP| \sin \varphi$ (siehe Abb. 7.5 a). Man erhält $x' = x \cos \alpha - y \sin \alpha$, $y' = y \cos \alpha + x \sin \alpha$ bzw. $d: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \cos \alpha - y \sin \alpha \\ x \sin \alpha + y \cos \alpha \end{pmatrix}$ als Komponenten-/Koordinatenbeschreibung der Drehung um den Ursprung mit dem Drehwinkel α , Abb. 7.5 b) zeigt ein Beispiel mit $\alpha = 110^\circ$. ■

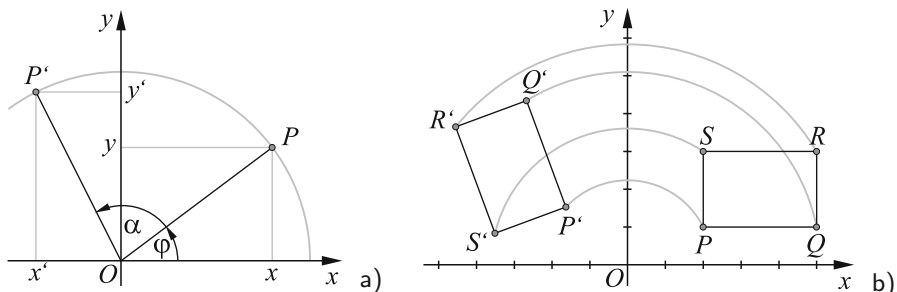


Abb. 7.5: Drehungen

Bemerkungen zu Punkt- und Vektorabbildungen:

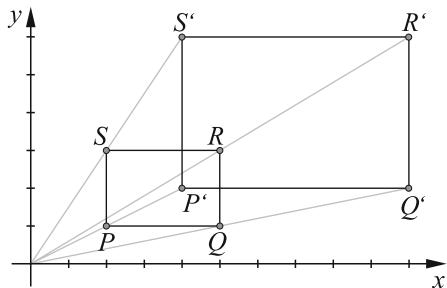
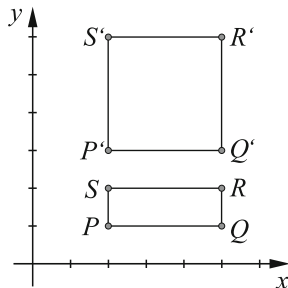
1. *Drehungen* lassen sich auch als *Vektorabbildungen* auffassen. Statt Punkten P lassen sich auch ihre Ortsvektoren gedreht denken. Dies kann durch Pfeile veranschaulicht werden, welche die Ortsvektoren $\overrightarrow{OP}$ und $\overrightarrow{OP}'$ von Punkten und ihren Bildpunkten repräsentieren und deren Anfangspunkte im Drehzentrum liegen. Hingegen ist eine *Auffassung von Verschiebungen als Vektorabbildungen* nicht sinnvoll. Da Vektoren geometrisch als Klassen gleich langer, paralleler und gleich orientierter Pfeile zu interpretieren sind (siehe Abschnitt 3.1) führt eine Verschiebung eines Pfeils zu einem Repräsentanten desselben Vektors (während bei einer Drehung und auch bei einer Spiegelung eine neue Pfeilkategorie entsteht). Eine klare Abgrenzung von Vektorabbildungen und Punktabbildungen sowie eine Exaktifizierung ihrer Zusammenhänge erfolgt in den folgenden Abschnitten.
2. Unter Verwendung der in dem Abschnitt 6.2 eingeführten Multiplikation von Matrizen und Spaltenvektoren erhalten wir für die in dem Beispiel 7.4 hergeleitete Darstellung von Drehungen $\begin{pmatrix} x \cos \alpha - y \sin \alpha \\ x \sin \alpha + y \cos \alpha \end{pmatrix} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix}$. Die Drehung von $\mathbb{R}^2$ um den Koordinatenursprung mit einem Drehwinkel α lässt sich somit als Vektorabbildung auch durch $d: \vec{x} \mapsto \mathbf{D} \circ \vec{x}$ mit der Drehmatrix $\mathbf{D} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}$ beschreiben. Auch Spiegelungen lassen sich (als Vektorabbildungen) matriziell beschreiben (siehe die Aufgabe 7 auf S. 263). Hingegen ist es nicht möglich, Verschiebungen auf diese Weise durch eine Matrix darzustellen.

Beispiel 7.5

Die Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} 2x \\ 2y \end{pmatrix}$ ist eine *zentrische Streckung* mit dem Zentrum im Koordinatenursprung und dem Streckfaktor 2 (Abb. 7.6). Fasst man f als Vektorabbildung auf, so lässt sich auch hierfür eine matrizielle Darstellung $f: \vec{x} \mapsto \mathbf{A} \circ \vec{x}$ mit der Streckungsmatrix $\mathbf{A} = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix}$ angeben. ■

Beispiel 7.6

Wird nur eine Komponente eines Vektors bzw. Punktes mit einer reellen Zahl multipliziert, so ergibt sich eine Streckung entlang einer der Koordinatenachsen. So ist die Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ 3y \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 3 \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix}$ eine *axiale*

**Abb. 7.6:** Zentrische Streckung**Abb. 7.7:** Axiale Streckung

Streckung in y -Richtung (siehe Abb. 7.7, aus Platzgründen wurde ein anderes Rechteck $PQRS$ abgebildet als in den vorherigen Beispielen). ■

Beispiel 7.7

Es wird die Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x-2y \\ y \end{pmatrix}$ betrachtet und es werden dabei wieder die Bildpunkte der in dem Beispiel 7.1 aufgeführten Punkte P, Q, R, S bestimmt. Die durch die angegebene Zuordnungsvorschrift beschriebene Abbildung ist eine *Scherung*. ■

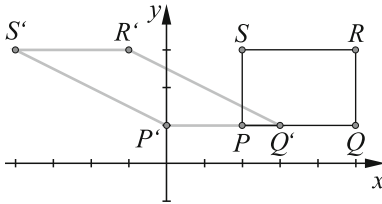


Abb. 7.8: Scherung

Beispiel 7.8

Wir betrachten die Abbildung $p: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ mit $p: \begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto \begin{pmatrix} x \\ y \end{pmatrix}$, die sich auch durch die Matrix $\mathbf{P} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}$ darstellen lässt. p bildet alle Punkte bzw. zugehörigen Ortsvektoren von $\mathbb{R}^3$ in $\mathbb{R}^2$ ab. Fasst man $\mathbb{R}^2$ als die x - y -Ebene von $\mathbb{R}^3$ auf, so lässt sich p geometrisch als *Parallelprojektion* des Raumes auf diese Ebene interpretieren. In der Abb. 7.9 wird dies durch die Eckpunkte $A = \begin{pmatrix} 4 \\ -7 \\ 3 \end{pmatrix}$, $B = \begin{pmatrix} 7 \\ -7 \\ 3 \end{pmatrix}$, $C = \begin{pmatrix} 7 \\ -3 \\ 3 \end{pmatrix}$, $D = \begin{pmatrix} 4 \\ -3 \\ 3 \end{pmatrix}$, $E = \begin{pmatrix} 4 \\ -7 \\ 5 \end{pmatrix}$, $F = \begin{pmatrix} 7 \\ -7 \\ 5 \end{pmatrix}$, $G = \begin{pmatrix} 7 \\ -3 \\ 5 \end{pmatrix}$, $H = \begin{pmatrix} 4 \\ -3 \\ 5 \end{pmatrix}$ eines Quaders und die zugehörigen Bildpunkte veranschaulicht. ■

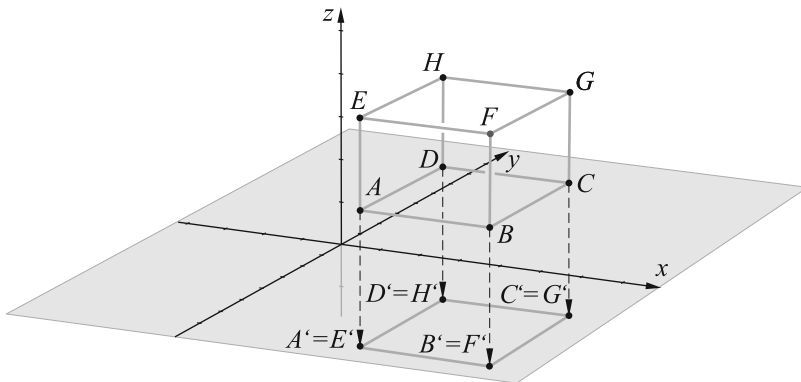


Abb. 7.9: Parallelprojektion

Nachdem nun einige Beispiele geometrischer Abbildungen behandelt wurden, erfolgen in den nächsten Abschnitten systematischere Untersuchungen von Vektor- und Punktabbildungen. Wir untersuchen dazu zunächst spezielle Vektorabbildungen (lineare Abbildungen) und dann auf dieser Grundlage Abbildungen affiner Punkträume.

7.1.3 Aufgaben zu Abschnitt 7.1

- Welche der folgenden Funktionen sind injektiv, surjektiv, bijektiv?
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = 3x+5$
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^3$
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = \sin x$
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = e^x$
 - $f: \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^3 - 5x^2$
- Welche der in den Beispielen 7.1-7.8 behandelten geometrischen Abbildungen sind injektiv, surjektiv, bijektiv?
- Geben Sie Koordinatenbeschreibungen folgender Abbildungen von $\mathbb{R}^2$ an:
 - Spiegelung an der Geraden mit der Gleichung $y=x$,
 - Punktspiegelung am Koordinatenursprung,
 - Spiegelung an der Geraden mit der Gleichung $y=6$.
- Gegeben sind eine Verschiebung $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 2 \\ 3 \end{pmatrix}$ und eine Spiegelung an der x -Achse $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ -y \end{pmatrix}$. Geben Sie Koordinatenbeschreibungen der Nacheinanderausführungen $g \circ f$ und $f \circ g$ an.
- Geben Sie Koordinatendarstellungen der beiden Nacheinanderausführungen der in dem Beispiel 7.1 behandelten Verschiebung und der zentrischen Streckung aus dem Beispiel 7.5 an.
- Beschreiben Sie die Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} 2x \\ y+1 \end{pmatrix}$ als Nacheinanderausführung zweier Abbildungen.
- Geben Sie Matrizen an, welche die Spiegelung aus dem Beispiel 7.2 (S. 259), die Scherung aus dem Beispiel 7.7 (S. 262) sowie die Abbildungen aus der Aufgabe 3 a) und b) beschreiben (siehe dazu die Bemerkung 2 auf S. 261).
- Geben Sie graphische Veranschaulichungen der folgenden Abbildungen, indem Sie das Rechteck $PQRS$ mit $P = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$, $Q = \begin{pmatrix} 5 \\ 1 \end{pmatrix}$, $R = \begin{pmatrix} 5 \\ 3 \end{pmatrix}$ und $S = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$ sowie das Rechteck, dessen Eckpunkte die Bildpunkte dieser Punkte sind, in einem Koordinatensystem darstellen:²
 - $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} 2x \\ 2y \end{pmatrix} - \begin{pmatrix} 4 \\ 4 \end{pmatrix}$,
 - $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x \\ y - \frac{3}{2}x \end{pmatrix}$,
 - $h: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x-3y \\ y-x \end{pmatrix}$,
 - Drehung von $\mathbb{R}^2$ um den Koordinatenursprung mit dem Drehwinkel 30° .

²Es können Zeichnungen mit der Hand angefertigt werden, aber es lässt sich auch der Computer dafür nutzen. Auf der Internetseite zu diesem Buch steht eine Datei für das Computeralgebrasystem Maxima zur Verfügung, unter deren Verwendung sich, nach Eingabe der Abbildungsgleichungen, sehr leicht die angegebenen Punkte und ihre Bildpunkte (sowie die entsprechenden Rechtecke) darstellen lassen.

7.2 Lineare Abbildungen

7.2.1 Definition und einige Beispiele

Definition 7.3

Es seien V und W Vektorräume über dem Körper der reellen Zahlen. Eine Abbildung $f: V \rightarrow W$, gegeben durch die Zuordnung $\vec{u} \mapsto f(\vec{u})$ (für alle $\vec{u} \in V$), heißt *lineare Abbildung* bzw. *Vektorraumhomomorphismus*, falls für beliebige $\vec{u}_1, \vec{u}_2 \in V$, $\lambda \in \mathbb{R}$ gilt:

- a) $f(\vec{u}_1 + \vec{u}_2) = f(\vec{u}_1) + f(\vec{u}_2)$ (Additivität),
- b) $f(\lambda \cdot \vec{u}_1) = \lambda \cdot f(\vec{u}_1)$ (Homogenität).

Die Menge aller linearen Abbildungen $f: V \rightarrow W$ wird mit $\text{Hom}(V, W)$ bezeichnet. ◆

Bemerkung: Die Definition 7.3 besagt, dass lineare Abbildungen *strukturtreu* bezüglich der Verknüpfungen „+“ und „ $\cdot$ “ sind, die Vektorräume kennzeichnen (siehe Def. 5.1 auf S. 168). Damit ist gemeint, dass das Bild $f(\vec{u}_1 + \vec{u}_2)$ der Summe zweier beliebiger Vektoren $\vec{u}_1, \vec{u}_2$ gleich der Summe der Bilder $f(\vec{u}_1)$ und $f(\vec{u}_2)$ dieser Vektoren sowie dass das Bild $f(\lambda \cdot \vec{u}_1)$ des Produkts eines Vektors mit einer reellen Zahl gleich dem Produkt des Bildes $f(\vec{u}_1)$ mit dieser Zahl ist.

Die Symbole „+“ und „ $\cdot$ “ treten in der Definition 7.3 in doppelter Bedeutung auf. In $f(\vec{u}_1 + \vec{u}_2)$ und $f(\lambda \cdot \vec{u}_1)$ stehen sie für die Addition und die skalare Multiplikation des Vektorraumes V . Hingegen werden in $f(\vec{u}_1) + f(\vec{u}_2)$ und in $\lambda \cdot f(\vec{u}_1)$ Vektoren von W miteinander bzw. mit reellen Zahlen verknüpft. Lineare Abbildungen sind nach der Definition 7.3 also *mit den Verknüpfungen auf V und W verträglich*.

Beispiele und Gegenbeispiele für lineare Abbildungen

Beispiel 7.9

Proportionale Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = a \cdot x$ (wobei $a \in \mathbb{R}$ ist) sind lineare Abbildungen (mit $V = W = \mathbb{R}$). Es gilt für beliebige $x_1, x_2 \in \mathbb{R}$, $\lambda \in \mathbb{R}$:

- a) $f(x_1 + x_2) = a(x_1 + x_2) = ax_1 + ax_2 = f(x_1) + f(x_2)$ sowie
- b) $f(\lambda x_1) = a \lambda x_1 = \lambda a x_1 = \lambda f(x_1)$. ■

Beispiel 7.10

Lineare Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = m \cdot x + n$ ($m, n \in \mathbb{R}$) sind für $n \neq 0$ *keine linearen Abbildungen* des Vektorraumes $\mathbb{R}$ in den Vektorraum $\mathbb{R}$:

$$\begin{aligned} f(x_1 + x_2) &= m(x_1 + x_2) + n = mx_1 + mx_2 + n, \\ f(x_1) + f(x_2) &= mx_1 + n + mx_2 + n = mx_1 + mx_2 + 2n. \end{aligned}$$

Somit ist $f(x_1 + x_2) \neq f(x_1) + f(x_2)$ für $n \neq 0$. Damit ist bereits gezeigt, dass lineare Funktionen mit $f(x) = m \cdot x + n$ für $n \neq 0$ keine linearen Abbildungen sind. (Leicht überprüft man, dass auch die Homogenität verletzt ist, obwohl dies nicht mehr notwendig ist, da lineare Abbildungen nach der Definition 7.3 additiv und homogen sein müssen.) ■

In dem Abschnitt 7.1.2 wurden einige geometrische Abbildungen betrachtet. Die meisten der dort untersuchten Abbildungen lassen sich auch als Vektorabbildungen auffassen und sind lineare Abbildungen (vgl. Aufgabe 1, S. 277). Hingegen sind Verschiebungen (Beispiel 7.1) sowie Nacheinanderausführungen von Verschiebungen und anderen Abbildungen (Beispiel 7.3) keine linearen Abbildungen. Wir betrachten im Folgenden weitere, z. T. verallgemeinerte Beispiele.

Beispiel 7.11

$f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x \\ y-z \end{pmatrix}$ ist eine lineare Abbildung, denn es gilt für beliebige $\begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix}, \begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix} \in \mathbb{R}^3$ und $\lambda \in \mathbb{R}$:

$$\begin{aligned} \text{a) } f\left(\begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix} + \begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix}\right) &= f\begin{pmatrix} x_1+x_2 \\ y_1+y_2 \\ z_1+z_2 \end{pmatrix} = \begin{pmatrix} x_1+x_2 \\ (y_1+y_2)-(z_1+z_2) \end{pmatrix} \\ &= \begin{pmatrix} x_1 \\ y_1-z_1 \end{pmatrix} + \begin{pmatrix} x_2 \\ y_2-z_2 \end{pmatrix} = f\begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix} + f\begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix} \text{ und} \end{aligned}$$

$$\text{b) } f\left(\lambda \begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix}\right) = f\begin{pmatrix} \lambda x_1 \\ \lambda y_1 \\ \lambda z_1 \end{pmatrix} = \begin{pmatrix} \lambda x_1 \\ \lambda y_1 - \lambda z_1 \end{pmatrix} = \lambda \begin{pmatrix} x_1 \\ y_1 - z_1 \end{pmatrix} = \lambda f\begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix}.$$

In der Definition 7.2 auf S. 258 wurden die *Injektivität* und die *Surjektivität* von Abbildungen definiert. Die lineare Abbildung f wäre genau dann injektiv, wenn für beliebige $\vec{x}_1 = \begin{pmatrix} x_1 \\ y_1 \\ z_1 \end{pmatrix}, \vec{x}_2 = \begin{pmatrix} x_2 \\ y_2 \\ z_2 \end{pmatrix} \in \mathbb{R}^3$ aus $f(\vec{x}_1) = f(\vec{x}_2)$ folgen würde $\vec{x}_1 = \vec{x}_2$. Dies hieße, aus $\begin{pmatrix} x_1 \\ y_1 - z_1 \end{pmatrix} = \begin{pmatrix} x_2 \\ y_2 - z_2 \end{pmatrix}$ müsste $\vec{x}_1 = \vec{x}_2$ resultieren. Offensichtlich ist dies nicht der Fall, denn das Beispiel $\vec{x}_1 = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \vec{x}_2 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$ belegt bereits das Gegenteil: Es ist $f(\vec{x}_1) = f(\vec{x}_2) = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$, aber $\vec{x}_1 \neq \vec{x}_2$. f ist also *nicht injektiv*.

f ist *surjektiv*, wenn für jeden Vektor $\vec{v} \in \mathbb{R}^2$ ein Vektor $\vec{x} \in \mathbb{R}^3$ mit $f(\vec{x}) = \vec{v}$ existiert. Dies ist der Fall, denn setzt man für einen beliebig vorgegebenen Vektor $\vec{v} = \begin{pmatrix} v_x \\ v_y \end{pmatrix} \in \mathbb{R}^2$ z. B. $\vec{x} = \begin{pmatrix} v_x \\ v_y \\ 0 \end{pmatrix}$, so ist $f(\vec{x}) = \begin{pmatrix} v_x \\ v_y - 0 \end{pmatrix} = \begin{pmatrix} v_x \\ v_y \end{pmatrix}$. ■

Beispiel 7.12

$f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} ax+by \\ cx+dy \end{pmatrix}$ ist für beliebige $a, b, c, d \in \mathbb{R}$ eine lineare Abbildung, denn es gilt für beliebige $\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}, \begin{pmatrix} x_2 \\ y_2 \end{pmatrix} \in \mathbb{R}^2, \lambda \in \mathbb{R}$:

$$\begin{aligned} \text{a) } f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + \begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right) &= f\begin{pmatrix} x_1+x_2 \\ y_1+y_2 \end{pmatrix} = \begin{pmatrix} a(x_1+x_2) + b(y_1+y_2) \\ c(x_1+x_2) + d(y_1+y_2) \end{pmatrix} \\ &= \begin{pmatrix} ax_1+ax_2+by_1+by_2 \\ cx_1+cx_2+dy_1+dy_2 \end{pmatrix}, \end{aligned}$$

$$f\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + f\begin{pmatrix} x_2 \\ y_2 \end{pmatrix} = \begin{pmatrix} ax_1+by_1 \\ cx_1+dy_1 \end{pmatrix} + \begin{pmatrix} ax_2+by_2 \\ cx_2+dy_2 \end{pmatrix} = \begin{pmatrix} ax_1+ax_2+by_1+by_2 \\ cx_1+cx_2+dy_1+dy_2 \end{pmatrix};$$

$$\text{b) } f\left(\lambda \begin{pmatrix} x_1 \\ y_1 \end{pmatrix}\right) = f\begin{pmatrix} \lambda x_1 \\ \lambda y_1 \end{pmatrix} = \begin{pmatrix} a \lambda x_1 + b \lambda y_1 \\ c \lambda x_1 + d \lambda y_1 \end{pmatrix}$$

$$\lambda f\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} = \lambda \begin{pmatrix} ax_1+by_1 \\ cx_1+dy_1 \end{pmatrix} = \begin{pmatrix} \lambda ax_1 + \lambda by_1 \\ \lambda cx_1 + \lambda dy_1 \end{pmatrix} = \begin{pmatrix} a \lambda x_1 + b \lambda y_1 \\ c \lambda x_1 + d \lambda y_1 \end{pmatrix}.$$

Wir untersuchen nun f in Abhängigkeit von a, b, c und d auf *Injektivität und Surjektivität*. f ist injektiv, wenn für zwei beliebige Vektoren $\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}, \begin{pmatrix} x_2 \\ y_2 \end{pmatrix} \in \mathbb{R}^2$

aus $f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}\right) = f\left(\begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right)$ folgt, dass $\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} = \begin{pmatrix} x_2 \\ y_2 \end{pmatrix}$ ist. Dies bedeutet, dass aus

$$\begin{pmatrix} ax_1 + by_1 \\ cx_1 + dy_1 \end{pmatrix} = \begin{pmatrix} ax_2 + by_2 \\ cx_2 + dy_2 \end{pmatrix} \quad \text{bzw.} \quad \begin{pmatrix} ax_1 + by_1 \\ cx_1 + dy_1 \end{pmatrix} - \begin{pmatrix} ax_2 + by_2 \\ cx_2 + dy_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

die Gleichheit der Vektoren $\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}$ und $\begin{pmatrix} x_2 \\ y_2 \end{pmatrix}$ folgt. Setzt man $x = x_1 - x_2$ und $y = y_1 - y_2$, so wird deutlich, dass diese Bedingung genau dann erfüllt und somit f genau dann injektiv ist, wenn das lineare Gleichungssystem

$$\begin{aligned} ax + by &= 0 \\ cx + dy &= 0 \end{aligned} \quad \text{bzw.} \quad \begin{pmatrix} a & b \\ c & d \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

nur die triviale Lösung $x = y = 0$ besitzt. Das ist genau dann der Fall, wenn die Koeffizientenmatrix $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ den Rang 2 hat. Unter Nutzung der in dem Abschnitt 6.1.2 herausgearbeiteten Zusammenhänge erhalten wir also das Ergebnis, dass f genau dann injektiv ist, wenn die Vektoren $\begin{pmatrix} a \\ c \end{pmatrix}$ und $\begin{pmatrix} b \\ d \end{pmatrix}$ sowie die Zeilenvektoren $(a; b)$ und $(c; d)$ jeweils linear unabhängig sind.

Die Abbildung f ist surjektiv, wenn für jeden Vektor $\vec{v} = \begin{pmatrix} v_x \\ v_y \end{pmatrix} \in \mathbb{R}^2$ ein Vektor $\vec{x} = \begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$ mit $f(\vec{x}) = \vec{v}$ existiert. Dies trifft genau dann zu, wenn das LGS

$$\begin{aligned} ax + by &= v_x \\ cx + dy &= v_y \end{aligned} \quad \text{bzw.} \quad \begin{pmatrix} a & b \\ c & d \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} v_x \\ v_y \end{pmatrix}$$

für beliebige $v_x, v_y \in \mathbb{R}$ lösbar ist, was wiederum genau dann zutrifft, wenn die Matrix $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ den Rang 2 hat. Die in diesem Beispiel betrachtete lineare Abbildung f ist also genau dann surjektiv, wenn sie injektiv ist. Dieser Zusammenhang ist aber nicht verallgemeinerbar, wie das Beispiel 7.11 zeigt. ■

Beispiel 7.13

$f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\left(\begin{pmatrix} x \\ y \end{pmatrix}\right) = \begin{pmatrix} x+a \\ y+b \end{pmatrix}$ ist für beliebige $a, b \in \mathbb{R}$ mit Ausnahme von $a = b = 0$ keine lineare Abbildung. Es gilt

$$\begin{aligned} f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + \begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right) &= f\left(\begin{pmatrix} x_1+x_2 \\ y_1+y_2 \end{pmatrix}\right) = \begin{pmatrix} x_1+x_2+a \\ y_1+y_2+b \end{pmatrix} \quad \text{und} \\ f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}\right) + f\left(\begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right) &= \begin{pmatrix} x_1+a \\ y_1+b \end{pmatrix} + \begin{pmatrix} x_2+a \\ y_2+b \end{pmatrix} = \begin{pmatrix} x_1+x_2+2a \\ y_1+y_2+2b \end{pmatrix}. \end{aligned}$$

Die Additivität $f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix} + \begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right) = f\left(\begin{pmatrix} x_1 \\ y_1 \end{pmatrix}\right) + f\left(\begin{pmatrix} x_2 \\ y_2 \end{pmatrix}\right)$ ist also nur dann gegeben, wenn $a = 2a$ und $b = 2b$ gilt, also nur im Falle $a = b = 0$. Dasselbe trifft, wie man leicht nachprüft, auch für die Homogenität zu. ■

Beispiel 7.14

Sind V und W beliebige Vektorräume, so ist die Abbildung $f: V \rightarrow W$ mit $f(\vec{u}) = \vec{o}$ für alle $\vec{u} \in V$, die also jedem Vektor des Vektorraumes V den Nullvektor des Vektorraumes W zuordnet, eine lineare Abbildung. Es gilt:

- a) $f(\vec{u}_1 + \vec{u}_2) = \vec{o} = \vec{o} + \vec{o} = f(\vec{u}_1) + f(\vec{u}_2)$ für alle $\vec{u}_1, \vec{u}_2 \in V$ sowie
- b) $f(\lambda \cdot \vec{u}) = \vec{o} = \lambda \vec{o} = \lambda f(\vec{u})$ für alle $\vec{u} \in V, \lambda \in \mathbb{R}$. ■

Beispiel 7.15

In dem Beispiel 5.10 (siehe S. 174) wurde erwähnt, dass die Menge aller auf einem Intervall $I = [a; b]$ (mit $a, b \in \mathbb{R}$, $a < b$) differenzierbaren reellwertigen Funktionen ein Vektorraum ist (wir nennen ihn D_I).

Die Abbildung $\frac{d}{dx} : D_I \rightarrow D_I$, gegeben durch die Zuordnung $f \mapsto f'$, die jeder Funktion $f \in D_I$ ihre erste Ableitung zuordnet, ist eine lineare Abbildung (die auch *Differentialoperator* genannt wird). Es gilt für beliebige $f, g \in D_I$:

- a) $\frac{d}{dx}(f+g) = \frac{d}{dx}f + \frac{d}{dx}g$, da für alle $x \in I$: $(f+g)'(x) = f'(x) + g'(x)$, sowie
 b) $\frac{d}{dx}(\lambda f) = \lambda \frac{d}{dx}f$ für beliebige $\lambda \in \mathbb{R}$ (wegen $(\lambda f)'(x) = \lambda f'(x)$ f. a. $x \in I$). ■

Beispiel 7.16

Die Abbildung $f: P_2 \rightarrow P_0$ von dem Vektorraum der Polynome höchstens 2. Grades in den Vektorraum der Polynome 0. Grades (siehe das Beispiel 5.5, S. 170), die jedem Polynom die Summe seiner Koeffizienten zuordnet, d. h.

$$f: p \mapsto q \quad \text{mit} \quad p(x) = a_2x^2 + a_1x + a_0, \quad q(x) = a_2 + a_1 + a_0$$

ist, wie man leicht nachprüft, eine lineare Abbildung. ■

7.2.2 Einige Eigenschaften linearer Abbildungen**Satz 7.1**

Für beliebige Vektorräume V, W und lineare Abbildungen $f: V \rightarrow W$ gilt:

- (i) $f(\lambda \vec{u} + \mu \vec{v}) = \lambda f(\vec{u}) + \mu f(\vec{v})$ (für beliebige $\lambda, \mu \in \mathbb{R}$; $\vec{u}, \vec{v} \in V$).
 (ii) $f\left(\sum_{i=1}^k \lambda_i \vec{u}_i\right) = \sum_{i=1}^k \lambda_i f(\vec{u}_i)$ (für beliebige $\lambda_1, \dots, \lambda_k \in \mathbb{R}$; $\vec{u}_1, \dots, \vec{u}_k \in V$).
 (iii) Der Nullvektor von V wird auf den Nullvektor von W abgebildet: $f(\vec{0}) = \vec{0}$.
 (iv) Das Bild des Gegenvektors eines beliebigen Vektors $\vec{u} \in V$ ist der Gegenvektor des Bildes $f(\vec{u})$, d. h. für alle $\vec{u} \in V$ gilt $f(-\vec{u}) = -f(\vec{u})$.
 (v) Eine linear abhängige Vektormenge von V wird stets auf eine linear abhängige Vektormenge von W abgebildet; ist also $\{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\} \subseteq V$ linear abhängig, so ist auch $\{f(\vec{u}_1); f(\vec{u}_2); \dots; f(\vec{u}_k)\} \subseteq W$ linear abhängig.
 (vi) Jeder lineare Unterraum U von V wird auf einen linearen Unterraum von W abgebildet, d. h. für jeden linearen Unterraum $U \subseteq V$ ist die Menge $f(U) = \{f(\vec{u}) \mid \vec{u} \in U\}$ ein linearer Unterraum von W .

Bemerkung: Eine zu (v) analoge Aussage über linear unabhängige Vektormengen gilt nicht. Es gibt durchaus lineare Abbildungen, die gewisse linear unabhängige Mengen auf linear abhängige Mengen von Vektoren abbilden (siehe die Aufgabe 4 auf S. 277).

Beweis:

- (i) Wegen der Definition 7.3 b) ist $f(\lambda \vec{u}) = \lambda f(\vec{u})$ und $f(\mu \vec{v}) = \mu f(\vec{v})$. Nach Definition 7.3 a) folgt daraus die Behauptung $f(\lambda \vec{u} + \mu \vec{v}) = \lambda f(\vec{u}) + \mu f(\vec{v})$.
 (ii) Diese Behauptung ergibt sich durch mehrfache Anwendung von (i).

- (iii) Für beliebige Vektoren $\vec{u} \in V$ ist $\vec{u} + \vec{o} = \vec{u}$ und somit $f(\vec{u} + \vec{o}) = f(\vec{u})$. Nach der Definition 7.3 a) folgt daraus $f(\vec{u}) + f(\vec{o}) = f(\vec{u})$. Damit ist nach der Definition 5.1 A3 (S. 168) $f(\vec{o})$ der Nullvektor von W .
- (iv) Der Gegenvektor $-\vec{u}$ eines beliebigen Vektors $\vec{u} \in V$ ist durch die Eigenschaft $\vec{u} + (-\vec{u}) = \vec{o}$ bestimmt (Definition 5.1 A4). Nach der Definition 7.3 a) folgt daraus $f(\vec{u}) + f(-\vec{u}) = f(\vec{o})$. Somit gilt wegen der oben bewiesenen Eigenschaft (iii): $f(\vec{u}) + f(-\vec{u}) = \vec{o}$. $f(-\vec{u})$ ist also der Gegenvektor von $f(\vec{u})$.
- (v) Eine Vektormenge $\{\vec{u}_1; \vec{u}_2; \dots; \vec{u}_k\} \subseteq W$ ist genau dann linear abhängig, wenn $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{R}$ existieren, von denen für mindestens ein i ($1 \leq i \leq k$) $\lambda_i \neq 0$ ist, so dass gilt: $\sum_{i=1}^k \lambda_i \vec{u}_i = \vec{o}$. Nach den bereits bewiesenen Teilen (ii) und (iii) folgt daraus $\sum_{i=1}^k \lambda_i f(\vec{u}_i) = f\left(\sum_{i=1}^k \lambda_i \vec{u}_i\right) = f(\vec{o}) = \vec{o}$. Der Nullvektor von W ist damit auf nicht triviale Weise als Linearkombination der Vektoren $f(\vec{u}_1), \dots, f(\vec{u}_k)$ darstellbar; die Menge $\{f(\vec{u}_1); f(\vec{u}_2); \dots; f(\vec{u}_k)\}$ ist also linear abhängig.

Der Beweis von (vi) wird als Aufgabe gestellt (Aufgabe 6, S. 277). $\square$

Der folgende Satz ist von fundamentaler Bedeutung für die Beschreibung linearer Abbildungen und insbesondere für ihre Darstellung durch Matrizen. Er wird daher mitunter als *Hauptsatz über lineare Abbildungen* bezeichnet.

Satz 7.2

Es seien V und W Vektorräume, $B = \{\vec{b}_1; \vec{b}_2; \dots; \vec{b}_n\}$ eine Basis von V sowie $\vec{w}_1, \vec{w}_2, \dots, \vec{w}_n$ beliebige Vektoren von W . Dann gibt es genau eine lineare Abbildung $f: V \rightarrow W$ mit $f(\vec{b}_1) = \vec{w}_1$, $f(\vec{b}_2) = \vec{w}_2$, $\dots$, $f(\vec{b}_n) = \vec{w}_n$.

Bemerkung: Der Satz 7.2 enthält zwei Teilaussagen. Er besagt erstens, dass für eine beliebige Basis eines n -dimensionalen Vektorraumes V und beliebige n Vektoren eines (ebenfalls beliebigen) Vektorraumes W stets eine lineare Abbildung existiert, welche die Basisvektoren von V (in beliebiger Zuordnung) auf die gegebenen n Vektoren von W abbildet.

Die zweite (und noch bedeutsamere) Aussage des Satzes 7.2 besteht darin, dass eine lineare Abbildung durch die Bilder der Vektoren einer Basis ihres Urbildraumes eindeutig bestimmt ist.

Beweis: Es seien eine Basis $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ von V und Vektoren $\vec{w}_1, \dots, \vec{w}_n$ von W gegeben. Jeder Vektor $\vec{u} \in V$ lässt sich mit eindeutig bestimmten $\lambda_1, \dots, \lambda_n$ als Linearkombination der gegebenen Basisvektoren darstellen:

$$\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i.$$

Wir legen nun für jeden Vektor $\vec{u} \in V$ sein Bild $f(\vec{u}) \in W$ folgendermaßen fest:

$$(*) \quad f(\vec{u}) = \sum_{i=1}^n \lambda_i \vec{w}_i.$$

Es ist zu zeigen, dass durch diese Festlegung tatsächlich eine lineare Abbildung nach Definition 7.3 gegeben ist.

- a) Sind $\vec{u}^{(1)} = \sum_{i=1}^n \lambda_i^{(1)} \vec{b}_i$ und $\vec{u}^{(2)} = \sum_{i=1}^n \lambda_i^{(2)} \vec{b}_i$ beliebige Vektoren von V , so ist $\vec{u}^{(1)} + \vec{u}^{(2)} = \sum_{i=1}^n (\lambda_i^{(1)} + \lambda_i^{(2)}) \vec{b}_i$. Nach der Abbildungsvorschrift (*) ist

$$f\left(\vec{u}^{(1)} + \vec{u}^{(2)}\right) = \sum_{i=1}^n \left(\lambda_i^{(1)} + \lambda_i^{(2)}\right) \vec{w}_i.$$

Da ebenfalls nach (*) die Bilder von $\vec{u}^{(1)}$ und $\vec{u}^{(2)}$ durch

$f(\vec{u}^{(1)}) = \sum_{i=1}^n \lambda_i^{(1)} \vec{w}_i$ und $f(\vec{u}^{(2)}) = \sum_{i=1}^n \lambda_i^{(2)} \vec{w}_i$ festgelegt sind, gilt

$$f\left(\vec{u}^{(1)}\right) + f\left(\vec{u}^{(2)}\right) = \sum_{i=1}^n \lambda_i^{(1)} \vec{w}_i + \sum_{i=1}^n \lambda_i^{(2)} \vec{w}_i = \sum_{i=1}^n \left(\lambda_i^{(1)} + \lambda_i^{(2)}\right) \vec{w}_i$$

und somit $f(\vec{u}^{(1)} + \vec{u}^{(2)}) = f(\vec{u}^{(1)}) + f(\vec{u}^{(2)})$; daher ist f additiv.

- b) Für einen beliebigen Vektor $\vec{u} \in V$, $\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i$, seinen durch (*) gegebenen Bildvektor $f(\vec{u}) = \sum_{i=1}^n \lambda_i \vec{w}_i$ sowie eine beliebige reelle Zahl r gilt

$$f(r \vec{u}) = f\left(\sum_{i=1}^n r \lambda_i \vec{b}_i\right) = \sum_{i=1}^n r \lambda_i \vec{w}_i \quad \text{ sowie } \quad r f(\vec{u}) = r \sum_{i=1}^n \lambda_i \vec{w}_i.$$

Offensichtlich gilt also $f(r \vec{u}) = r f(\vec{u})$ (Homogenität).

Wir haben somit konstruktiv anhand der Abbildungsvorschrift (*) gezeigt, dass eine lineare Abbildung f existiert, welche die Bedingung des Satzes 7.2 erfüllt. Es bleibt die *Eindeutigkeit* zu zeigen, d. h. dass es nicht zwei verschiedene lineare Abbildungen f_1 und f_2 geben kann, die $\vec{b}_i$ auf $\vec{w}_i$ (für $i = 1 \dots n$) abbilden. Wäre dies der Fall, so müsste wegen der Additivität und der Homogenität von f_1 und f_2 bzw. wegen des Satzes 7.1 (ii) für jeden Vektor $\vec{u} \in V$, $\vec{u} = \sum_{i=1}^n \lambda_i \vec{b}_i$ gelten:

$$\begin{aligned} f_1(\vec{u}) &= f_1\left(\sum_{i=1}^n \lambda_i \vec{b}_i\right) = \sum_{i=1}^n \lambda_i f_1(\vec{b}_i) = \sum_{i=1}^n \lambda_i \vec{w}_i \quad \text{ sowie } \\ f_2(\vec{u}) &= f_2\left(\sum_{i=1}^n \lambda_i \vec{b}_i\right) = \sum_{i=1}^n \lambda_i f_2(\vec{b}_i) = \sum_{i=1}^n \lambda_i \vec{w}_i. \end{aligned}$$

Somit gilt also $f_1(\vec{u}) = f_2(\vec{u})$ für alle $\vec{u} \in V$; f_1 und f_2 sind also identisch. Auch die Eindeutigkeitsaussage des Satzes 7.2 ist damit nachgewiesen. $\square$

Beispiel 7.17

Es wird eine lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ durch Angabe der Bilder der Vektoren $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ und $\vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$ der Standardbasis $B_0 = \{\vec{e}_1; \vec{e}_2; \vec{e}_3\}$ beschrieben:

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

Aus diesen Angaben lässt sich das Bild eines beliebigen Vektors $\vec{x} \in \mathbb{R}^3$ bestimmen:

$$\begin{aligned} f(\vec{x}) &= f\left(\begin{pmatrix} x \\ y \\ z \end{pmatrix}\right) = f(x \vec{e}_1 + y \vec{e}_2 + z \vec{e}_3) = x f(\vec{e}_1) + y f(\vec{e}_2) + z f(\vec{e}_3) \\ &= x \begin{pmatrix} 1 \\ 0 \end{pmatrix} + y \begin{pmatrix} 0 \\ 1 \end{pmatrix} + z \begin{pmatrix} 0 \\ -1 \end{pmatrix} = \begin{pmatrix} x \\ y - z \end{pmatrix} \end{aligned}$$

Durch die gegebenen Bilder der Basisvektoren $\vec{e}_1, \vec{e}_2, \vec{e}_3$ wird also dieselbe lineare Abbildung beschrieben wie in dem Beispiel 7.11 auf S. 265. $\blacksquare$

Beispiel 7.18

Gegeben ist eine lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^4$ durch die Bilder der Basisvektoren

$\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$, $\vec{b}_3 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$ (siehe das Beispiel 5.23 auf S. 187):

$$f(\vec{b}_1) = \begin{pmatrix} -1 \\ 0 \\ 2 \\ 3 \end{pmatrix}, f(\vec{b}_2) = \begin{pmatrix} 2 \\ -2 \\ 0 \\ 1 \end{pmatrix}, f(\vec{b}_3) = \begin{pmatrix} 1 \\ 3 \\ 0 \\ -1 \end{pmatrix}.$$

Um für f eine Abbildungsgleichung der Form $f(\vec{x}) = f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} x'_1 \\ x'_2 \\ x'_3 \\ x'_4 \end{pmatrix}$ auf-

zustellen, müssen zunächst beliebige Vektoren $\vec{x} \in \mathbb{R}^3$ als Linearkombinationen bezüglich der Basis $B = \{\vec{b}_1; \vec{b}_2; \vec{b}_3\}$ dargestellt werden, d. h. für

$$\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \lambda_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$$

die Koeffizienten $\lambda_1, \lambda_2, \lambda_3$ bestimmt werden. Durch Lösung des entsprechenden LGS erhält man $\lambda_1 = x_1 - x_2$, $\lambda_2 = x_2 - x_3$, $\lambda_3 = x_3$. Damit lässt sich nun das Bild eines beliebigen Vektors $\vec{x} \in \mathbb{R}^3$ angeben:

$$\begin{aligned} f(\vec{x}) &= f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \lambda_1 \begin{pmatrix} -1 \\ 0 \\ 2 \\ 3 \end{pmatrix} + \lambda_2 \begin{pmatrix} 2 \\ -2 \\ 0 \\ 1 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ 3 \\ 0 \\ -1 \end{pmatrix} \\ &= (x_1 - x_2) \begin{pmatrix} -1 \\ 0 \\ 2 \\ 3 \end{pmatrix} + (x_2 - x_3) \begin{pmatrix} 2 \\ -2 \\ 0 \\ 1 \end{pmatrix} + x_3 \begin{pmatrix} 1 \\ 3 \\ 0 \\ -1 \end{pmatrix} = \begin{pmatrix} -x_1 + 3x_2 - x_3 \\ -2x_2 + 5x_3 \\ 2x_1 - 2x_2 \\ 3x_1 - 2x_2 - 2x_3 \end{pmatrix}. \end{aligned}$$

7.2.3 Matrizen Darstellung linearer Abbildungen

Bereits in dem Abschnitt 7.1.2 wurden spezielle geometrische Abbildungen in $\mathbb{R}^2$ und $\mathbb{R}^3$ durch Matrizen beschrieben, siehe die entsprechende Bemerkung auf S. 261 und die Beispiele 7.5, 7.6 und 7.8. Im Folgenden verallgemeinern wir das Konzept der matriziellen Beschreibung auf der Grundlage des Satzes 7.2.

Beispiel 7.19

Wir betrachten die lineare Abbildung f aus dem Beispiel 7.18 und ordnen ihr eine Matrix zu, indem wir die Bilder der Basisvektoren von B nebeneinander schreiben:

$$\mathbf{A}_{BB_0}^f = \begin{pmatrix} -1 & 2 & 1 \\ 0 & -2 & 3 \\ 2 & 0 & 0 \\ 3 & 1 & -1 \end{pmatrix}.$$

Diese Matrix wird mit $\mathbf{A}_{BB_0}^f$ bezeichnet, da sie die lineare Abbildung f bezüglich der (geordneten) Basen B und der Standardbasis B_0^4 von $\mathbb{R}^4$ beschreibt. Letzteres kommt dadurch zustande, dass die Vektoren $f(\vec{b}_1)$, $f(\vec{b}_2)$ und $f(\vec{b}_3)$ als Spaltenvektoren, d. h. als Linearkombinationen der Standardbasis von $\mathbb{R}^4$

gegeben sind, z. B. $f(\vec{b}_1) = \begin{pmatrix} -1 \\ 0 \\ 2 \\ 3 \end{pmatrix} = (-1) \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix} + 0 \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix} + 2 \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix} + 3 \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}$. Es

stehen also in den Spalten der Matrix $\mathbf{A}_{BB_0}^f$ die Bilder der Vektoren einer Basis

des Urbildraumes $\mathbb{R}^3$. Diese sind als Linearkombinationen der Basisvektoren der Standardbasis des Bildraumes $\mathbb{R}^4$ angegeben, welche damit den Zeilen der Matrix zuzuordnen sind. Das folgende Schema verdeutlicht dies.

$$\begin{array}{c} \text{Bilder der Vektoren der Basis } B \\ \left[\begin{array}{ccc} f(\vec{b}_1) & f(\vec{b}_2) & f(\vec{b}_3) \end{array} \right] \\ \mathbf{A}_{BB_0}^f = \left(\begin{array}{ccc} -1 & 2 & 1 \\ 0 & -2 & 3 \\ 2 & 0 & 0 \\ 3 & 1 & -1 \end{array} \right) \left[\begin{array}{c} \vec{e}_1 \\ \vec{e}_2 \\ \vec{e}_3 \\ \vec{e}_4 \end{array} \right] \end{array} \left. \vphantom{\begin{array}{c} \text{Bilder der Vektoren der Basis } B \\ \left[\begin{array}{ccc} f(\vec{b}_1) & f(\vec{b}_2) & f(\vec{b}_3) \end{array} \right] \\ \mathbf{A}_{BB_0}^f = \left(\begin{array}{ccc} -1 & 2 & 1 \\ 0 & -2 & 3 \\ 2 & 0 & 0 \\ 3 & 1 & -1 \end{array} \right) \left[\begin{array}{c} \vec{e}_1 \\ \vec{e}_2 \\ \vec{e}_3 \\ \vec{e}_4 \end{array} \right]} \right\} \begin{array}{l} \text{Vektoren der} \\ \text{Standardbasis} \\ B_0^4 \text{ von } \mathbb{R}^4 \end{array}$$

Mithilfe der Matrix $\mathbf{A}_{BB_0}^f$ lassen sich Bilder von Vektoren $\vec{x}$ des Urbildraumes berechnen, wenn diese als Linearkombinationen $\vec{x} = \lambda_1 \vec{b}_1 + \lambda_2 \vec{b}_2 + \lambda_3 \vec{b}_3$ der Basis B , d. h. in Koordinaten $\lambda_1, \lambda_2, \lambda_3$ bezüglich B gegeben sind (Abschnitt 5.4.2):

$$f(x) = \mathbf{A}_{BB_0}^f \circ \begin{pmatrix} \lambda_1 \\ \lambda_2 \\ \lambda_3 \end{pmatrix} = \begin{pmatrix} -1 & 2 & 1 \\ 0 & -2 & 3 \\ 2 & 0 & 0 \\ 3 & 1 & -1 \end{pmatrix} \circ \begin{pmatrix} \lambda_1 \\ \lambda_2 \\ \lambda_3 \end{pmatrix} = \begin{pmatrix} -\lambda_1 + 2\lambda_2 + \lambda_3 \\ -2\lambda_2 + 3\lambda_3 \\ 2\lambda_1 \\ 3\lambda_1 + \lambda_2 - \lambda_3 \end{pmatrix}.$$

Ist f durch eine Abbildungsvorschrift wie $f(\vec{x}) = f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -x_1 + 3x_2 - x_3 \\ -2x_2 + 5x_3 \\ 2x_1 - 2x_2 \\ 3x_1 - 2x_2 - 2x_3 \end{pmatrix}$ bzw. durch die (daraus ablesbaren) Bilder der Vektoren der Standardbasis B_0^3 des Urbildraumes $\mathbb{R}^3$: $f \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} = \begin{pmatrix} -1 \\ 0 \\ 2 \\ 3 \end{pmatrix}$, $f \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 3 \\ -2 \\ -2 \\ -2 \end{pmatrix}$, $f \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = \begin{pmatrix} -1 \\ 5 \\ 0 \\ -2 \end{pmatrix}$ gegeben, so lässt sich daraus eine *Abbildungsmatrix von f bezüglich der Standardbasen* aufstellen.

$$\begin{array}{c} \text{Bilder der Vektoren der Standardbasis } B_0^3 \text{ von } \mathbb{R}^3 \\ \left[\begin{array}{ccc} f(\vec{e}_1) & f(\vec{e}_2) & f(\vec{e}_3) \end{array} \right] \\ \mathbf{A}_{B_0^3 B_0^4}^f = \left(\begin{array}{ccc} -1 & 3 & -1 \\ 0 & -2 & 5 \\ 2 & -2 & 0 \\ 3 & -2 & -2 \end{array} \right) \left[\begin{array}{c} \vec{e}_1 \\ \vec{e}_2 \\ \vec{e}_3 \\ \vec{e}_4 \end{array} \right] \end{array} \left. \vphantom{\begin{array}{c} \text{Bilder der Vektoren der Standardbasis } B_0^3 \text{ von } \mathbb{R}^3 \\ \left[\begin{array}{ccc} f(\vec{e}_1) & f(\vec{e}_2) & f(\vec{e}_3) \end{array} \right] \\ \mathbf{A}_{B_0^3 B_0^4}^f = \left(\begin{array}{ccc} -1 & 3 & -1 \\ 0 & -2 & 5 \\ 2 & -2 & 0 \\ 3 & -2 & -2 \end{array} \right) \left[\begin{array}{c} \vec{e}_1 \\ \vec{e}_2 \\ \vec{e}_3 \\ \vec{e}_4 \end{array} \right]} \right\} \begin{array}{l} \text{Vektoren der} \\ \text{Standardbasis} \\ B_0^4 \text{ von } \mathbb{R}^4 \end{array}$$

Mithilfe dieser Matrix lassen sich Bilder von Vektoren, die in Spaltenschreibweise gegeben sind, berechnen:

$$f(x) = \mathbf{A}_{B_0^3 B_0^4}^f \circ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -1 & 3 & -1 \\ 0 & -2 & 5 \\ 2 & -2 & 0 \\ 3 & -2 & -2 \end{pmatrix} \circ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -x_1 + 3x_2 - x_3 \\ -2x_2 + 5x_3 \\ 2x_1 - 2x_2 \\ 3x_1 - 2x_2 - 2x_3 \end{pmatrix}. \quad \blacksquare$$

Bemerkungen:

- In dem Beispiel 7.19 ist von *geordneten Basen* die Rede. Während bei unseren bisherigen Betrachtungen zu Basen die Reihenfolge von Basisvektoren keine Bedeutung hatte, muss sie bei der Aufstellung von Matrizen linearer Abbildungen unbedingt beachtet werden.
- Dieselbe lineare Abbildung kann durch unterschiedliche Matrizen beschrieben werden, da sich Matrizen linearer Abbildungen stets auf Basen beziehen.

- Die Anzahl der Spalten der Matrix einer linearen Abbildung $f: V \rightarrow W$ ist gleich der Anzahl der Vektoren einer Basis des Urbildraumes V und somit gleich der Dimension des Urbildraumes; die Anzahl der Zeilen entspricht der Dimension des Bildraumes W .
- Da wir lineare Abbildungen im Folgenden hauptsächlich durch Matrizen bezüglich der Standardbasen beschreiben, verzichten wir auf die Angabe der Basen in Bezeichnungen wie $\mathbf{A}_{B_0^3 B_0^4}^f$ und schreiben einfach $\mathbf{A}^f$.

Matrizen linearer Abbildungen $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ bezüglich der Standardbasen

Eine lineare Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$, die durch die Bilder der Vektoren der Standardbasis des Urbildraumes $\mathbb{R}^n$ bzw. durch eine Abbildungsvorschrift

$$f \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{pmatrix}$$

gegeben ist, lässt sich durch eine Matrix beschreiben, deren Spaltenvektoren die Bilder der Vektoren der Standardbasis des Urbildraumes sind.

Bilder der Vektoren der Standardbasis B_0^n von $\mathbb{R}^n$

$$\mathbf{A}^f = \left(\begin{array}{c|c|c|c|c} f(\vec{e}_1) & f(\vec{e}_2) & \dots & f(\vec{e}_n) & \\ \hline a_{11} & a_{12} & \dots & a_{1n} & \vec{e}_1 \\ a_{21} & a_{22} & \dots & a_{2n} & \vec{e}_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} & \vec{e}_m \end{array} \right) \left. \begin{array}{l} \\ \\ \\ \\ \end{array} \right\} \begin{array}{l} \text{Vektoren der} \\ \text{Standardbasis} \\ B_0^m \text{ von } \mathbb{R}^m \end{array}$$

Das Bild eines beliebigen Vektors $\vec{x}$ ergibt sich durch sein Produkt mit der Abbildungsmatrix: $f(\vec{x}) = \mathbf{A}^f \circ \vec{x}$.

Die Abbildungsgleichung $f(\vec{x}) = \mathbf{A}^f \circ \vec{x}$ weist auf einen *Zusammenhang zwischen linearen Abbildungen und linearen Gleichungssystemen* hin. Auf S. 248f. wurde herausgearbeitet, dass sich jedes LGS in der Form $\mathbf{A} \circ \vec{x} = \vec{b}$ schreiben lässt. $f(\vec{x})$ entspricht also dem Spaltenvektor $\vec{b}$ der absoluten Glieder des betrachteten LGS. Bei einer linearen Abbildung $f: \mathbb{R}^n \rightarrow \mathbb{R}^m$ zu einem gegebenen Vektor $\vec{b} \in \mathbb{R}^m$ Urbilder zu suchen, bedeutet somit, das LGS $\mathbf{A}^f \circ \vec{x} = \vec{b}$ zu lösen.

Beispiel 7.20

Gesucht werden alle Vektoren, die durch die in den Beispielen 7.18 und 7.19

behandelte lineare Abbildung f auf den Vektor $\vec{b} = \begin{pmatrix} 4 \\ -2 \\ -3 \end{pmatrix}$ abgebildet werden.

Dazu ist das LGS

$$\mathbf{A}_{B_0^3 B_0^4}^f \circ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} -1 & 3 & -1 \\ 0 & -2 & 5 \\ 2 & -2 & 0 \\ 3 & -2 & -2 \end{pmatrix} \circ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 4 \\ 1 \\ -2 \\ -3 \end{pmatrix}$$

zu lösen. Man erhält die (eindeutige) Lösung $\vec{x} = \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix}$; dieser Vektor ist also der (einzige) Vektor, der durch f auf $\vec{b}$ abgebildet wird. ■

7.2.4 Nacheinanderausführung linearer Abbildungen

Definition 7.4

Es seien A , B und C beliebige Mengen sowie $f: A \rightarrow B$ und $g: B \rightarrow C$ Abbildungen. Als *Nacheinanderausführung* bzw. *Verkettung* bzw. *Komposition* der Abbildungen f und g bezeichnet man die Abbildung $g \circ f$, die durch die Zuordnung $(g \circ f)(x) = g(f(x))$ für alle $x \in A$ gegeben ist.

$$A \xrightarrow{f} B \xrightarrow{g} C \quad x \xrightarrow{f} f(x) \xrightarrow{g} g(f(x)) \quad \blacklozenge$$

In dem Beispiel 7.3 auf S. 259 wurden bereits Nacheinanderausführungen von Spiegelungen und Verschiebungen untersucht. Im Folgenden betrachten wir Nacheinanderausführungen linearer Abbildungen.

Satz 7.3

Sind $f: U \rightarrow V$ und $g: V \rightarrow W$ lineare Abbildungen, so ist die Nacheinanderausführung $g \circ f: U \rightarrow W$ ebenfalls eine lineare Abbildung.

Der Beweis des Satzes wird als Übungsaufgabe gestellt (Aufgabe 16, S. 278).

Beispiel 7.21

Wir betrachten die in dem Beispiel 7.2 (S. 259) behandelte Spiegelung an der y -Achse $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} -x \\ y \end{pmatrix}$ sowie die in dem Beispiel 7.7 untersuchte Scherung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x-2y \\ y \end{pmatrix}$ als Vektorabbildungen.

$$f \circ g: g: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} -x \\ y \end{pmatrix}, f: \begin{pmatrix} x' \\ y' \end{pmatrix} \mapsto \begin{pmatrix} x'' \\ y'' \end{pmatrix} = \begin{pmatrix} x' - 2y' \\ y' \end{pmatrix} = \begin{pmatrix} -x - 2y \\ y \end{pmatrix}$$

$$g \circ f: f: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} x - 2y \\ y \end{pmatrix}, g: \begin{pmatrix} x' \\ y' \end{pmatrix} \mapsto \begin{pmatrix} x'' \\ y'' \end{pmatrix} = \begin{pmatrix} -x' \\ y' \end{pmatrix} = \begin{pmatrix} -x + 2y \\ y \end{pmatrix}$$

Somit ist $(f \circ g) \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} -x - 2y \\ y \end{pmatrix}$ und $(g \circ f) \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} -x + 2y \\ y \end{pmatrix}$ für alle $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$.

Eine Vertauschung der Reihenfolge führt also zu unterschiedlichen Ergebnissen. Die Abb. 7.10 zeigt $g \circ f$ und $f \circ g$ als Punktabbildungen: dargestellt werden Punkte P, Q, R, S , die durch Ortsvektoren $\vec{p}, \vec{q}, \vec{r}, \vec{s}$ beschrieben sind, sowie Punkte P', Q', R', S' und P'', Q'', R'', S'' mit den durch Abbildung der Vektoren $\vec{p}, \vec{q}, \vec{r}, \vec{s}$ entstehenden Vektoren $\vec{p}', \vec{q}', \vec{r}', \vec{s}'$ bzw. $\vec{p}'', \vec{q}'', \vec{r}'', \vec{s}''$ als Ortsvektoren.

Wir betrachten nun die Matrizen der linearen Abbildungen f und g bezüglich der Standardbasis von $\mathbb{R}^2$: $\mathbf{A}^g = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$, $\mathbf{A}^f = \begin{pmatrix} 1 & -2 \\ 0 & 1 \end{pmatrix}$. Es gilt

$$\mathbf{A}^f \circ \mathbf{A}^g = \begin{pmatrix} 1 & -2 \\ 0 & 1 \end{pmatrix} \circ \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} -1 & -2 \\ 0 & 1 \end{pmatrix}, \quad \mathbf{A}^g \circ \mathbf{A}^f = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix} \circ \begin{pmatrix} 1 & -2 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} -1 & 2 \\ 0 & 1 \end{pmatrix}.$$

Diese Produktmatrizen entsprechen, wie man leicht nachprüft, den oben hergeleiteten Abbildungsvorschriften von $g \circ f$ und $f \circ g$. ■

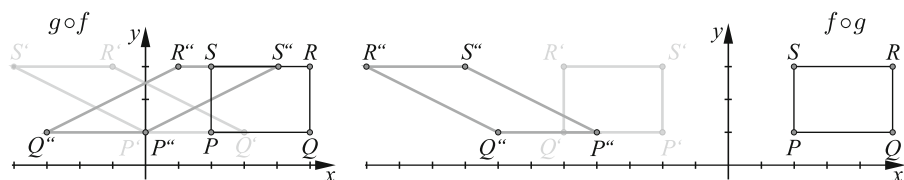


Abb. 7.10: Nacheinanderausführung einer Scherung f und einer Spiegelung g

Das Beispiel 7.21 gibt Anlass zu der Vermutung, dass sich Nacheinanderausführungen linearer Abbildungen stets durch die Produkte zugehöriger Matrizen beschreiben lassen. Mit dem folgenden Satz (auf dessen Beweis wir verzichten) wird dies festgeschrieben und verallgemeinert.

Satz 7.4

Es seien $f: U \rightarrow V$ und $g: V \rightarrow W$ lineare Abbildungen, $\mathbf{A}^f$ die Matrix von f bezüglich zweier Basen B_U von U und B_V von V sowie $\mathbf{A}^g$ die Matrix von g bezüglich der Basis B_V und einer Basis B_W von W . Dann beschreibt die Matrix $\mathbf{A}^g \circ \mathbf{A}^f$ die Abbildung $g \circ f$ bezüglich der Basen B_U und B_W .

Beispiel 7.22

Wir betrachten die in den Beispielen auf S. 265 untersuchten linearen Abbildungen $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ mit $f \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x \\ y-z \end{pmatrix}$ und $g: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $g \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} ax+by \\ cx+dy \end{pmatrix}$. Es gibt nur eine Möglichkeit, f und g nacheinander auszuführen, denn die zuerst ausgeführte Abbildung muss in den Vektorraum abbilden, welcher der Urbildraum der danach ausgeführten Abbildung ist. Daher kann in diesem Beispiel $f \circ g$ nicht gebildet werden: g bildet in $\mathbb{R}^2$ ab, aber der Urbildraum von f ist $\mathbb{R}^3$. Hingegen ist es möglich $g \circ f$ zu bilden: $\mathbb{R}^3 \xrightarrow{f} \mathbb{R}^2 \xrightarrow{g} \mathbb{R}^2$.

Matrizen, die f und g bezüglich der Standardbasen von $\mathbb{R}^3$ bzw. $\mathbb{R}^2$ beschreiben, sind $\mathbf{A}^f = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & -1 \end{pmatrix}$ und $\mathbf{A}^g = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ (vgl. die Aufgabe 9 auf S. 277). Das Produkt dieser Matrizen ist

$$\mathbf{A}^g \circ \mathbf{A}^f = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \circ \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & -1 \end{pmatrix} = \begin{pmatrix} a & b & -b \\ c & d & -d \end{pmatrix}.$$

Daraus erhält man eine Abbildungsvorschrift für $g \circ f$:

$$(g \circ f)(\vec{x}) = \mathbf{A}^g \circ \mathbf{A}^f \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} a & b & -b \\ c & d & -d \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} ax + by - bz \\ cx + dy - dz \end{pmatrix}. \quad \blacksquare$$

7.2.5 Isomorphismen

Definition 7.5

Eine bijektive lineare Abbildung heißt *Isomorphismus*. ◆

Mit der wahrscheinlich kürzesten Definition dieses Buches wird zugleich einer der „stärksten“ Strukturbegriffe der linearen Algebra beschrieben. Dies liegt daran, dass lineare Abbildungen an sich bereits zentrale strukturerhaltende Eigenschaften besitzen und die Bijektivität, wie sich herausstellen wird, zusätzlich ihre Umkehrbarkeit ermöglicht und die Erhaltung der linearen Unabhängigkeit von Vektormengen gewährleistet (dies ist für lineare Abbildungen nicht in jedem Falle gegeben). Isomorphismen werden somit die „Strukturgleichheit“ von Vektorräumen beschreiben.

Eine wichtige Bedeutung für die Untersuchung von Isomorphismen haben die Matrizen, welche diese speziellen linearen Abbildungen beschreiben, wie das folgende Beispiel verdeutlicht.

Beispiel 7.23

In dem Beispiel 7.12 auf S. 265 haben wir festgestellt, dass die lineare Abbildung $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} ax+by \\ cx+dy \end{pmatrix}$ (mit $a, b, c, d \in \mathbb{R}$) genau dann surjektiv ist, wenn ihre Abbildungsmatrix $\begin{pmatrix} a & b \\ c & d \end{pmatrix}$ den Rang 2 hat (also regulär ist), und dass f auch genau in diesem Falle injektiv ist. Nach der Definition 7.5 ist f also auch genau dann ein Isomorphismus, wenn diese Matrix regulär ist. ■

Der in dem Beispiel 7.23 konstatierte Zusammenhang zwischen der Eigenschaft einer linearen Abbildung f , ein Isomorphismus zu sein, und der Regularität einer Abbildungsmatrix von f besteht allgemein.

Satz 7.5

Es sei $f: V \rightarrow W$ eine lineare Abbildung. Ist $\mathbf{A}_{B_V B_W}^f$ eine Abbildungsmatrix von f bezüglich einer Basis B_V von V und einer Basis B_W von W , so ist f genau dann ein Isomorphismus, wenn $\mathbf{A}_{B_V B_W}^f$ eine reguläre Matrix ist.

Wir verzichten auf einen Beweis des Satzes 7.5, weisen aber auf einen wichtigen Sachverhalt hin: Da jeder Vektorraum (außer dem hier nicht weiter interessierenden Raum, der nur aus dem Nullvektor besteht) mehrere Basen besitzt, kann jede lineare Abbildung durch verschiedene Matrizen beschrieben werden. Der Satz 7.5 besagt zunächst, dass eine lineare Abbildung ein Isomorphismus ist, wenn *eine* ihrer Abbildungsmatrizen regulär ist. Da der Satz aber auf alle Basen von V und W und damit auf alle Matrizen angewendet werden kann, die f beschreiben, beinhaltet er auch, dass *alle* Abbildungsmatrizen eines Isomorphismus $f: V \rightarrow W$ (bezüglich beliebiger Basen von V bzw. W) regulär sind.

Umkehrbarkeit von Isomorphismen

Eine Abbildung $f^{-1}: B \rightarrow A$ heißt *Umkehrabbildung* einer Abbildung $f: A \rightarrow B$, wenn die Nacheinanderausführungen $f^{-1} \circ f$ und $f \circ f^{-1}$ die identischen Abbildungen von A bzw. B sind, also alle Elemente der Mengen A bzw. B auf sich selbst abbilden. Dies ist der Fall, wenn für beliebige Elemente genau dann $f(a) = b$ ist, wenn $f^{-1}(b) = a$ gilt, f^{-1} also gegenüber f „vertauschte“ Zuordnungen vornimmt. Veranschaulichen lässt sich dies in der Abb. 7.1 auf S. 258 durch Vertauschung der Pfeilrichtungen bei der Abbildung f_4 – dort erkennt man auch, dass f_1 und f_2 keine Umkehrabbildungen besitzen und dass durch die Umkehrung der Zuordnungen von f_3 keine Abbildung von B in A gegeben ist, denn b_6 besäße dabei kein Bild. Allgemein gilt für Abbildungen zwischen beliebigen Mengen, dass bijektive Abbildungen Umkehrabbildungen besitzen und diese ebenfalls bijektiv sind.

Satz 7.6

Jeder Isomorphismus $f: V \rightarrow W$ besitzt eine Umkehrabbildung $f^{-1}: W \rightarrow V$, die ebenfalls ein Isomorphismus ist.

Dass Isomorphismen Umkehrabbildungen besitzen (oder kurz: umkehrbar sind), folgt aus ihrer Bijektivität (siehe die Bemerkungen oberhalb des Satzes 7.6),

ebenso die Tatsache, dass diese bijektiv sind. Man könnte nun die Linearität der Umkehrabbildungen nachweisen, siehe z. B. Jänich (2008), S. 82. Die Gültigkeit des Satzes 7.6 wird aber auch anhand der Matrizendarstellung von linearen Abbildungen plausibel. Ist $\mathbf{A}$ eine Abbildungsmatrix eines Isomorphismus f , so ist $\mathbf{A}$ nach dem Satz 7.5 regulär. Damit besitzt $\mathbf{A}$ nach dem Satz 6.9 (S. 247) eine inverse Matrix $\mathbf{A}^{-1}$, diese ist eine Abbildungsmatrix von f^{-1} .

Beispiel 7.24

Wir betrachten erneut den Isomorphismus $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ aus dem Beispiel 7.23 mit der Abbildungsmatrix $\mathbf{A} = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ und berechnen die inverse Matrix

$$\mathbf{A}^{-1} = \begin{pmatrix} \frac{d}{a d - b c} & -\frac{b}{a d - b c} \\ -\frac{c}{a d - b c} & \frac{a}{a d - b c} \end{pmatrix}.$$

Es muss $ad - bc = \det \mathbf{A} \neq 0$ sein, da $\mathbf{A}$ regulär ist, vgl. Satz 6.13. $\mathbf{A}^{-1}$ ist die Matrix des Umkehrisomorphismus f^{-1} von f bezüglich der Standardbasis von $\mathbb{R}^2$; die Abbildungsvorschrift von f^{-1} ist somit $f^{-1} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \frac{d}{a d - b c} x - \frac{b}{a d - b c} y \\ -\frac{c}{a d - b c} x + \frac{a}{a d - b c} y \end{pmatrix}$.

Führt man f und f^{-1} nacheinander aus, so erhält man für $f^{-1} \circ f$:

$$\begin{aligned} f \begin{pmatrix} x \\ y \end{pmatrix} &= \begin{pmatrix} x' \\ y' \end{pmatrix} = \begin{pmatrix} ax + by \\ cx + dy \end{pmatrix}, \\ f^{-1} \begin{pmatrix} x' \\ y' \end{pmatrix} &= \begin{pmatrix} \frac{d}{a d - b c} x' - \frac{b}{a d - b c} y' \\ -\frac{c}{a d - b c} x' + \frac{a}{a d - b c} y' \end{pmatrix} = \begin{pmatrix} \frac{d \cdot (ax + by)}{a d - b c} - \frac{b \cdot (cx + dy)}{a d - b c} \\ -\frac{c \cdot (ax + by)}{a d - b c} + \frac{a \cdot (cx + dy)}{a d - b c} \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}, \end{aligned}$$

also $(f^{-1} \circ f) \begin{pmatrix} x \\ y \end{pmatrix} = f^{-1} \left(f \begin{pmatrix} x \\ y \end{pmatrix} \right) = \begin{pmatrix} x \\ y \end{pmatrix}$.

Für $f \circ f^{-1}$ erhält man dasselbe Ergebnis. Durch beide möglichen Nacheinanderführungen des Isomorphismus f und seiner Umkehrabbildung f^{-1} ergibt sich die identische Abbildung, welche jeden Vektor von $\mathbb{R}^2$ auf sich selbst abbildet. ■

Isomorphe Vektorräume

Definition 7.6

Zwei Vektorräume V und W heißen *isomorph* (zueinander), wenn ein Isomorphismus $f: V \rightarrow W$ existiert. ♦

Satz 7.7

Zwei (endlichdimensionale) Vektorräume V und W gleicher Dimension sind stets isomorph zueinander.

Bemerkung: Sind zwei Vektorräume V und W isomorph zueinander, so sagt man auch, *es besteht Isomorphie zwischen V und W* . Das Wort Isomorphie (griech.) bedeutet in der wörtlichen Übersetzung „Gestaltgleichheit“, besser wird sein Inhalt jedoch durch „Strukturgleichheit“ zum Ausdruck gebracht.

Beispiele

- Bereits in dem Abschnitt 3.3.2 wurde die Isomorphie zwischen der Menge der Pfeilklassen der Ebene bzw. des Raumes und $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ herausgearbeitet. Diese Isomorphie lag bereits an vielen Stellen dieses Buches der Beschreibung geometrischer Sachverhalte durch Vektoren von $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ zu Grunde.

- Die Isomorphie zwischen dem *Vektorraum der $m \times n$ -Matrizen* (siehe das Beispiel 5.6 auf S. 171) und $\mathbb{R}^{m \cdot n}$ wurde in dem Beispiel 5.29 auf S. 197 genutzt, um eine Basis des Vektorraumes aller magischen Quadrate der Kantenlänge 3 zu ermitteln.
- In dem Beispiel 5.26 auf S. 189 wurde eine Basis des *Vektorraumes der Polynome höchstens n -ten Grades* aufgestellt, wobei die „Verwandschaft“ dieses Vektorraumes zu $\mathbb{R}^{n+1}$ auffiel. In der Tat lässt sich leicht zeigen, dass diese Vektorräume zueinander isomorph sind.

7.2.6 Aufgaben zu Abschnitt 7.2

1. Weisen Sie nach, dass die in den Beispielen 7.2 und 7.4-7.8 (siehe S. 259-262) betrachteten Abbildungen (wenn man sie als Vektorabbildungen auffasst) lineare Abbildungen sind.
2. Welche der folgenden Abbildungen sind lineare Abbildungen? Begründen Sie Ihre Antworten.

a) $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x+y \\ x \end{pmatrix}$

b) $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x^2 \\ y^2 \end{pmatrix}$

c) $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \cdot y \\ x \end{pmatrix}$

d) $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $f\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} y \\ -y \end{pmatrix}$
3. Welche der in den Beispielen 7.9, 7.14, 7.15 und 7.18 (siehe S. 264-267 und S. 270) behandelten linearen Abbildungen sind injektiv, surjektiv, bijektiv?
4. Begründen Sie anhand eines Gegenbeispiels, dass folgende Aussage falsch ist: Eine linear unabhängige Vektormenge von V wird durch eine lineare Abbildung $f: V \rightarrow W$ stets auf eine linear unabhängige Menge von W abgebildet.
5. Weisen Sie nach: Ist $f: V \rightarrow W$ eine injektive lineare Abbildung, so wird durch f eine beliebige linear unabhängige Vektormenge von V auf eine linear unabhängige Menge von W abgebildet.
6. Beweisen Sie den Satz 7.1 (vi): Durch eine lineare Abbildung $f: V \rightarrow W$ wird jeder lineare Unterraum von V auf einen linearen Unterraum von W abgebildet, d. h. für jeden linearen Unterraum $U \subseteq V$ ist $f(U) = \{f(\vec{u}) \mid \vec{u} \in U\}$ ein linearer Unterraum von W .
7. Weisen Sie nach, dass für beliebige lineare Abbildungen $f: V \rightarrow W$ gilt: Sind $\vec{u}_1, \dots, \vec{u}_k$ Vektoren von V , so ist die lineare Hülle $\langle f(\vec{u}_1), \dots, f(\vec{u}_k) \rangle$ ihrer Bildvektoren gleich dem Bild $f(\langle \vec{u}_1, \dots, \vec{u}_k \rangle) = \{f(\vec{u}) \mid \vec{u} \in \langle \vec{u}_1, \dots, \vec{u}_k \rangle\}$ der linearen Hülle $\langle \vec{u}_1, \dots, \vec{u}_k \rangle$ von $\vec{u}_1, \dots, \vec{u}_k$.
8. Begründen Sie, dass es keine lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ geben kann, die folgende Zuordnungen trifft: $f\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $f\begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$, $f\begin{pmatrix} 7 \\ 8 \\ 9 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$.
9. Stellen Sie Matrizen für die in den Beispielen 7.11 und 7.12 auf S. 265 gegebenen linearen Abbildungen auf.

10. Geben Sie eine Matrix der linearen Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ an, die jeden Vektor von $\mathbb{R}^3$ auf den Nullvektor von $\mathbb{R}^2$ abbildet.
11. Gegeben ist eine lin. Abb. $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ durch die Bilder der Basisvektoren $\vec{b}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\vec{b}_2 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$, $\vec{b}_3 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$: $f(\vec{b}_1) = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$, $f(\vec{b}_2) = \begin{pmatrix} 4 \\ 5 \\ 6 \end{pmatrix}$, $f(\vec{b}_3) = \begin{pmatrix} 7 \\ 8 \\ 9 \end{pmatrix}$.
- a) Geben Sie für f eine Abbildungsgleichung der Form $f \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} x'_1 \\ x'_2 \\ x'_3 \end{pmatrix}$ an.
- b) Stellen Sie für f eine Matrix $\mathbf{A}_{B_0^3 B_0^3}^f$ auf (vgl. das Beispiel 7.19 auf S. 270).
- c) Stellen Sie für f eine Matrix $\mathbf{A}_{B_0^3 B_0^3}^f$ auf.
- d) Ist f injektiv, surjektiv bzw. sogar bijektiv?
12. Bestimmen Sie die Menge aller Vektoren, die durch die in der Aufgabe 11 gegebene lineare Abbildung f auf den Nullvektor abgebildet werden.
13. Weisen Sie nach: Ist $f: V \rightarrow W$ eine lineare Abbildung, so ist die Menge aller Vektoren, die durch f auf den Nullvektor von W abgebildet werden, ein linearer Unterraum von V . (Dieser Unterraum wird als *Kern* der linearen Abbildung f bezeichnet.)
14. Bestimmen Sie die Menge aller Vektoren $\vec{w} \in \mathbb{R}^3$, die bei der in der Aufgabe 11 gegebenen linearen Abbildung f als Bilder von Vektoren $\vec{v} \in \mathbb{R}^3$ auftreten, also die Menge $\{\vec{w} \mid \exists \vec{v} \in \mathbb{R}^3: \vec{w} = f(\vec{v})\}$.
15. Weisen Sie nach: Ist $f: V \rightarrow W$ eine lineare Abbildung, so ist die Menge $\{\vec{w} \mid \exists \vec{v} \in V: \vec{w} = f(\vec{v})\}$ aller Vektoren $\vec{w} \in W$, die bei f als Bilder von Vektoren $\vec{v} \in V$ auftreten, ein linearer Unterraum von W . (Dieser Unterraum wird als *Bild* der linearen Abbildung f bezeichnet.)
16. Beweisen Sie, dass die Nacheinanderausführung $g \circ f$ linearer Abbildungen $f: U \rightarrow V$ und $g: V \rightarrow W$ eine lineare Abbildung ist (Satz 7.3, S. 273).
17. Bilden Sie die beiden Nacheinanderausführungen der in dem Beispiel 7.4 auf S. 260 behandelten Drehung und der in dem Beispiel 7.5 (S. 261) betrachteten zentrischen Streckung. Geben Sie jeweils die Abbildungsmatrix an.
18. In welcher Reihenfolge lassen sich die Drehung aus dem Beispiel 7.4 (S. 260) und die Projektion aus dem Beispiel 7.8 (S. 262) nacheinander ausführen? Geben Sie für diese Nacheinanderausführung eine Abbildungsmatrix an.
19. Welche der in den Beispielen 7.2 und 7.4-7.8 (S. 259-262) behandelten linearen Abbildungen sind Isomorphismen? Geben Sie für diese die Umkehrabbildungen an.
20. Begründen Sie: Wird eine lineare Abbildung f durch eine Matrix $\mathbf{A}^f$ beschrieben, so ist f genau dann ein Isomorphismus, wenn die Determinante von $\mathbf{A}^f$ ungleich Null ist.
21. Beweisen Sie den Satz 7.7 auf S. 276.

7.3 Affine Abbildungen

7.3.1 Definition und Beispiele

In dem Abschnitt 7.1.2 wurden geometrische Abbildungen wie Verschiebungen, Drehungen, Spiegelungen, Streckungen sowie eine Scherung und eine Projektion betrachtet. In dem Abschnitt 7.2 stellte sich heraus, dass sich die meisten dieser – zunächst als Punktabbildungen eingeführten – Abbildungen auch als Vektorabbildungen auffassen lassen und bei dieser Betrachtungsweise lineare Abbildungen sind. Im Folgenden wird das Verhältnis von Punkt- und Vektorabbildungen durch die Einführung des Begriffs der affinen Abbildung präzisiert.

Definition 7.7

Es seien A und A' affine Punkträume mit zugehörigen Vektorräumen V bzw. V' . Eine Abbildung $\phi: A \rightarrow A'$ heißt *affine Abbildung*, wenn die durch die Zuordnung $f: \overrightarrow{PQ} \mapsto \overrightarrow{\phi(P)\phi(Q)}$ (für beliebige Punkte $P, Q \in A$) festgelegte Vektorabbildung $f: V \rightarrow V'$ eine lineare Abbildung ist. $\blacklozenge$

Beispiel 7.25

In dem Beispiel 7.3 auf S. 259 wurden *Schubspiegelungen* als Nacheinander- ausführungen von Verschiebungen und Spiegelungen betrachtet. Wir bezeichnen die dort untersuchte Abbildung $g \circ f$ mit ϕ . Es ist also $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $\phi: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} -x+4 \\ y+1 \end{pmatrix}$. Wir zeigen im Folgenden, dass ϕ eine affine Abbildung ist.

Für beliebige Punkte $P = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix}$ und $Q = \begin{pmatrix} x_2 \\ y_2 \end{pmatrix}$ ist $\overrightarrow{PQ} = \begin{pmatrix} x_2 - x_1 \\ y_2 - y_1 \end{pmatrix}$. Weiterhin ist $\phi(P) = \begin{pmatrix} -x_1+4 \\ y_1+1 \end{pmatrix}$, $\phi(Q) = \begin{pmatrix} -x_2+4 \\ y_2+1 \end{pmatrix}$ und somit

$$\overrightarrow{\phi(P)\phi(Q)} = \begin{pmatrix} -x_2+4 \\ y_2+1 \end{pmatrix} - \begin{pmatrix} -x_1+4 \\ y_1+1 \end{pmatrix} = \begin{pmatrix} -x_2+x_1 \\ y_2-y_1 \end{pmatrix}.$$

Bezeichnet man $x_2 - x_1$ mit ξ und $y_2 - y_1$ mit η , so erhält man $\overrightarrow{PQ} = \begin{pmatrix} \xi \\ \eta \end{pmatrix}$ sowie $\overrightarrow{\phi(P)\phi(Q)} = \begin{pmatrix} -\xi \\ \eta \end{pmatrix}$. Die der Abbildung ϕ durch $f: \overrightarrow{PQ} \mapsto \overrightarrow{\phi(P)\phi(Q)}$ zugeordnete Vektorabbildung f mit $f: \begin{pmatrix} \xi \\ \eta \end{pmatrix} \mapsto \begin{pmatrix} -\xi \\ \eta \end{pmatrix}$ (Spiegelung an der y -Achse) ist offensichtlich eine lineare Abbildung, siehe die Aufgabe 1 auf S. 277. Somit ist die Schubspiegelung ϕ nach der Definition 7.7 eine affine Abbildung. $\blacksquare$

Beispiel 7.26

Lineare Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = mx + n$ ($m, n \in \mathbb{R}$) sind, wie in dem Beispiel 7.10 (S. 264) überraschenderweise festgestellt wurde, für $n \neq 0$ keine linearen Abbildungen; wir zeigen, dass es sich dabei um affine Abbildungen handelt. Um den Bezug zu der Definition 7.7 hervorzuheben, bezeichnen wir f mit ϕ . Sowohl Punkte als auch Vektoren sind bei diesem Beispiel reelle Zahlen. Für beliebige $P = x_1$, $Q = x_2$ ist $\overrightarrow{PQ} = x_2 - x_1$, $\phi(P) = mx_1 + n$, $\phi(Q) = mx_2 + n$ sowie $\overrightarrow{\phi(P)\phi(Q)} = m(x_2 - x_1)$. Die durch $g: \overrightarrow{PQ} \mapsto \overrightarrow{\phi(P)\phi(Q)}$ gegebene Vektorabbildung ist die proportionale Funktion $g: \mathbb{R} \rightarrow \mathbb{R}$ mit $g(x) = mx$. Diese ist eine lineare Abbildung (Beispiel 7.9, S. 264); somit ist ϕ bzw. f eine affine Abbildung. $\blacksquare$

Schreibweisen, Abbildungsvorschriften für affine Abbildungen

Bei der Einführung affiner Punkträume haben wir festgestellt, dass die Schreibweisen $Q = P + \vec{v}$ und $\vec{v} = \overrightarrow{PQ}$ dieselbe Bedeutung haben (S. 203). Wir verwenden dies nun, um aus dem in der Definition 7.7 durch $f: \overrightarrow{PQ} \mapsto \overrightarrow{\phi(P)\phi(Q)}$ bzw.

$$f(\overrightarrow{PQ}) = \overrightarrow{\phi(P)\phi(Q)}$$

gegebenen Zusammenhang zwischen affinen und linearen Abbildungen eine Abbildungsvorschrift für affine Abbildungen abzuleiten. Mit $\vec{v} = \overrightarrow{PQ}$ folgt daraus:

$$\phi(P) + f(\vec{v}) = \phi(P) + f(\overrightarrow{PQ}) = \phi(P) + \overrightarrow{\phi(P)\phi(Q)} = \phi(Q) \quad \text{bzw.}$$

$$\phi(Q) = \phi(P) + f(\vec{v}).$$

Somit können Bilder beliebiger Punkte $Q \in A$ aus dem Bild $\phi(P)$ eines einzigen Punktes P und Bildern der Verbindungsvektoren $\vec{v} = \overrightarrow{PQ}$ bei der linearen Abbildung f bestimmt werden. Meist verwendet man für P den Koordinatenursprung O eines in dem affinen Raum A gegebenen Koordinatensystems.

Eine affine Abbildung $\phi: A \rightarrow A'$ ist durch das Bild $\phi(O)$ eines (festen) Punktes $O \in A$ sowie die zugehörige lineare Abbildung $f: V \rightarrow V'$ eindeutig bestimmt. Der Bildpunkt $\phi(X)$ eines beliebigen Punktes $X \in A$ berechnet sich aus $\phi(O)$ und dem Bild des Vektors $\overrightarrow{OX}$ bei der linearen Abbildung f durch $\phi(X) = \phi(O) + f(\overrightarrow{OX})$. Bezüglich eines Koordinatensystems mit dem Ursprung O und mit Ortsvektoren $\vec{x} = \overrightarrow{OX}$ von Punkten X ist

$$\phi(X) = \phi(O) + f(\vec{x}).$$

Wir werden hauptsächlich affine Abbildungen innerhalb der affinen Räume $\mathbb{R}^2$ und $\mathbb{R}^3$ behandeln, wobei $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ zugleich die zugehörigen Vektorräume sind. In Bezug auf die Definition 7.7 betrachten wir also Beispiele mit $A = A' = \mathbb{R}^2$ und $A = A' = \mathbb{R}^3$. Wenn $A = A'$ ist, liegen die Punkte O und $O' = \phi(O)$ in demselben Raum, und es gibt daher einen Vektor $\overrightarrow{OO'}$. Die oben herausgearbeitete Abbildungsvorschrift lässt sich damit folgendermaßen schreiben:

$$(*) \quad \phi(X) = O + \overrightarrow{OO'} + f(\vec{x}).$$

Hieran wird deutlich, dass sich eine affine Abbildung ϕ innerhalb desselben Raumes aus der zu ϕ gehörenden linearen (Vektor-)Abbildung f sowie einer Verschiebung mit dem Verschiebungsvektor $\overrightarrow{OO'}$ „zusammensetzt“. Bei der Schubspiegelung aus dem Beispiel 7.25 wird dies besonders in der Schreibweise

$$\phi(P) = \phi\left(\begin{pmatrix} x \\ y \end{pmatrix}\right) = \begin{pmatrix} 4 \\ 1 \end{pmatrix} + \begin{pmatrix} -x \\ y \end{pmatrix}$$

deutlich. Gegenüber der allgemeinen Abbildungsvorschrift $(*)$ kann dabei auf $O = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ verzichtet werden, da sich $\begin{pmatrix} 4 \\ 1 \end{pmatrix} + \begin{pmatrix} -x \\ y \end{pmatrix}$ als Punkt auffassen lässt, während im allgemeinen Fall $\overrightarrow{OO'} + f(\vec{x})$ Vektoren sind, die mit dem Koordinatenursprung addiert werden, um Punkte zu beschreiben.

Affine Abbildungen lassen sich aus linearen Abbildungen der zugehörigen Vektorräume sowie Verschiebungen zusammengesetzt auffassen.

Für affine Abbildungen $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ bzw. $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ lassen sich (bei Beschreibung der zugehörigen linearen Abbildungen durch Matrizen) allgemeine Abbildungsvorschriften aufstellen:

$$\begin{pmatrix} x' \\ y' \end{pmatrix} = \phi \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix} \quad (7.1)$$

bzw.

$$\begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} = \phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \\ v_z \end{pmatrix}. \quad (7.2)$$

Dabei sind $\begin{pmatrix} v_x \\ v_y \end{pmatrix}$ bzw. $\begin{pmatrix} v_x \\ v_y \\ v_z \end{pmatrix}$ die Verschiebungsvektoren der jeweiligen affinen

Abbildungen ϕ sowie $\begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ bzw. $\begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$ die Matrizen der zu ϕ gehörenden linearen Abbildungen bezüglich der Standardbasen von $\mathbb{R}^2$ bzw. $\mathbb{R}^3$.

In dem Abschnitt 7.1.2 (S. 259ff.) wurden einige Abbildungen in $\mathbb{R}^2$ betrachtet. Die dort behandelten Spiegelungen, Drehungen, Streckungen und Scherungen sind – als Punktabbildungen aufgefasst – affine Abbildungen mit dem Verschiebungsvektor $\vec{0}$. Dies trifft auch auf die Projektion $p: \mathbb{R}^3 \rightarrow \mathbb{R}^2$ (Beispiel 7.8) zu. Im Folgenden betrachten wir einige affine Abbildungen innerhalb von $\mathbb{R}^3$ anhand der Gleichung (7.2).

Beispiel 7.27

Durch die Abbildungsgleichung $\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 6 \\ -1 \\ -2 \end{pmatrix}$ wird eine affine Abbildung beschrieben, die sich aus einer zentrischen Streckung mit dem Streckfaktor 3 und dem Zentrum im Koordinatenursprung sowie einer Verschiebung zusammensetzt. In der Abb. 7.11 a) auf S. 282 ist diese Abbildung anhand eines Würfels und seiner Bildfigur dargestellt.

Alle Verbindungsgeraden zwischen beliebigen Punkten des Ursprungswürfels und ihren Bildpunkten schneiden sich in einem Punkt Z . Dies gilt sogar für alle Punkte $P \in \mathbb{R}^3$ und ihre Bildpunkte. Um dies zu bestätigen, suchen wir nach einem Punkt Z , der durch die Abbildung ϕ auf sich selbst abgebildet wird. ■

Definition 7.8

Es seien A ein affiner Raum und $\phi: A \rightarrow A$ eine affine Abbildung innerhalb von A . Ein Punkt $Z \in A$, der durch ϕ auf sich selbst abgebildet wird, für den also $\phi(Z) = Z$ gilt, wird als *Fixpunkt* von ϕ bezeichnet. ♦

Beispiel 7.28

Für die in dem Beispiel 7.27 betrachtete affine Abbildung wird nach Fixpunkten $Z = \begin{pmatrix} z_1 \\ z_2 \\ z_3 \end{pmatrix}$ mit $\phi(Z) = Z$ gesucht. Dies bedeutet für dieses Beispiel, dass

$$\begin{pmatrix} z_1 \\ z_2 \\ z_3 \end{pmatrix} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} z_1 \\ z_2 \\ z_3 \end{pmatrix} + \begin{pmatrix} 6 \\ -1 \\ -2 \end{pmatrix} \quad \text{bzw.} \quad \begin{aligned} z_1 &= 3z_1 + 6 \\ z_2 &= 3z_2 - 1 \\ z_3 &= 3z_3 - 2 \end{aligned}$$

gelten muss. Als (eindeutige) Lösung erhält man $z_1 = -3$, $z_2 = \frac{1}{2}$, $z_3 = 1$.

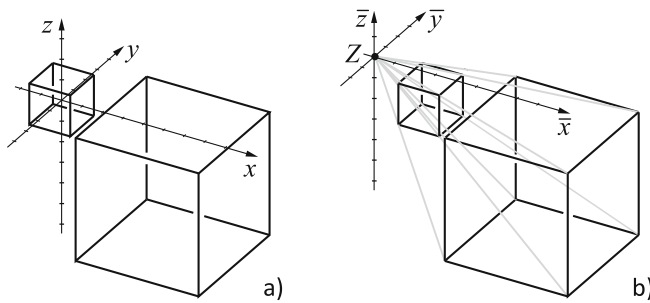


Abb. 7.11:
Zentrische Streckung
und Verschiebung

Die Abbildungen dieses Abschnitts können mithilfe einer Datei für das CAS Maxima, die auf der Internetseite zu diesem Buch zur Verfügung steht, interaktiv aus verschiedenen Perspektiven betrachtet werden.

Die Abbildung ϕ kann als zentrische Streckung mit dem Zentrum Z und dem Streckfaktor 3 aufgefasst werden. Dies zeigt sich anhand eines Koordinatensystems $\bar{K}$ mit Z als Koordinatenursprung und den unveränderten Basisvektoren der Standardbasis (siehe Abb. 7.11 b). Hierfür hat ϕ die Abbildungsvorschrift

$$\begin{pmatrix} \bar{x}' \\ \bar{y}' \\ \bar{z}' \end{pmatrix} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} \bar{x} \\ \bar{y} \\ \bar{z} \end{pmatrix}.$$

Die Koordinaten $\bar{x}, \bar{y}, \bar{z}$ eines beliebigen Punktes P bezüglich des Koordinatensystems $\bar{K}$ stehen, wegen der Koordinaten von Z , mit den „alten“ Koordinaten in der Beziehung $x = \bar{x} - 3, y = \bar{y} + \frac{1}{2}, z = \bar{z} + 1$; $\bar{x} = x + 3, \bar{y} = y - \frac{1}{2}, \bar{z} = z - 1$. Setzt man in die obige Gleichung die „alten“ Koordinaten ein, so erhält man

$$\begin{pmatrix} \bar{x}' \\ \bar{y}' \\ \bar{z}' \end{pmatrix} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} \bar{x} \\ \bar{y} \\ \bar{z} \end{pmatrix} = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} x+3 \\ y-\frac{1}{2} \\ z-1 \end{pmatrix} = \begin{pmatrix} 3x+9 \\ 3y-\frac{3}{2} \\ 3z-3 \end{pmatrix}.$$

Daraus ergibt sich

$$\begin{pmatrix} x' \\ y' \\ z' \end{pmatrix} = \begin{pmatrix} \bar{x}' - 3 \\ \bar{y}' + \frac{1}{2} \\ \bar{z}' + 1 \end{pmatrix} = \begin{pmatrix} 3x+9-3 \\ 3y-\frac{3}{2}+\frac{1}{2} \\ 3z-3+1 \end{pmatrix} = \begin{pmatrix} 3x+6 \\ 3y-1 \\ 3z-2 \end{pmatrix},$$

womit bestätigt ist, dass die zentrische Streckung mit dem Zentrum Z und dem Streckfaktor 3 dieselbe Abbildung ist, die in dem Beispiel 7.27 gegeben wurde. ■

Beispiel 7.29

Die Abbildungsgleichung $\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 4 & 0 \\ 0 & 0 & \frac{1}{2} \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 5 \\ -2 \\ -3 \end{pmatrix}$ beschreibt axiale Streckungen mit den Streckfaktoren 2, 4 und $\frac{1}{2}$ entlang der Koordinatenachsen und eine anschließende Verschiebung (siehe Abb. 7.12). Auch diese Abbildung besitzt genau einen Fixpunkt (siehe die Aufgabe 2 auf S. 286). ■

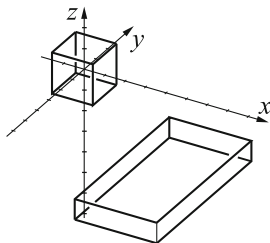


Abb. 7.12:
Axiale Streckung und Verschiebung

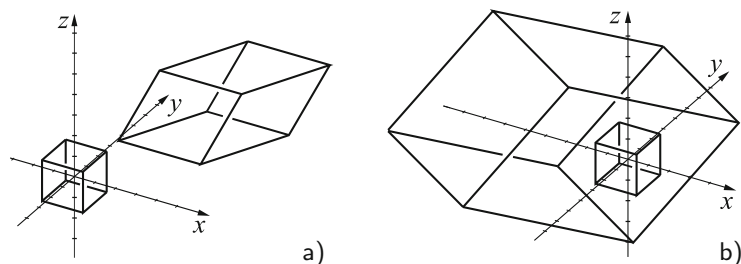


Abb. 7.13:
Scherungen
in verallgemeinertem
Sinne

Beispiel 7.30

Die in der Abb. 7.13 a) dargestellte affine Abbildung mit der Abbildungsmatrix $\mathbf{A} = \begin{pmatrix} 2 & 1 & 1 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix}$ und dem Verschiebungsvektor $\vec{v} = \begin{pmatrix} 5 \\ 4 \\ 3 \end{pmatrix}$ kombiniert Scherungen (siehe das Beispiel 7.7 auf S. 262), Streckungen und Verschiebungen. ■

In verallgemeinertem Sinne wird der Begriff „Scherung“ synonym zu „affine Abbildung des Raumes auf sich“ verwendet und umfasst u. a. auch Kongruenzabbildungen und Streckungen. Besonders zu beachten ist dabei das Wörtchen „auf“, welches gleichbedeutend mit der Surjektivität einer Abbildung ist – einer Eigenschaft, die alle bisher betrachteten Beispiele dieses Abschnittes besitzen, die aber dennoch nicht selbstverständlich ist. Gibt man willkürlich Abbildungsmatrizen affiner Abbildungen an, so werden in den meisten Fällen verallgemeinerte Scherungen dadurch beschrieben. Der Leserin bzw. dem Leser wird empfohlen, mithilfe des Computers unterschiedliche Abbildungsmatrizen „auszuprobieren“ und ihre Wirkung zu veranschaulichen. Dafür steht auf der Internetseite zu diesem Buch eine vorbereitete Datei für das Computeralgebrasystem Maxima zur Verfügung, die alle Beispiele dieses Abschnittes enthält und ihre Variation sowie die Eingabe eigener Matrizen und Verschiebungsvektoren ermöglicht. Das folgende Beispiel entstand auf diese Weise.

Beispiel 7.31

Die Wirkung der Abbildung ϕ mit $\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -3 & -3 & 2 \\ -2 & 2 & 1 \\ 1 & 0 & 3 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} -2 \\ -2 \\ 2 \end{pmatrix}$ auf einen Würfel zeigt die Abb. 7.13 b). ■

In den Beispielen 7.27-7.31 zeigten sich einige Gemeinsamkeiten aller bisher betrachteten affinen Abbildungen. Es scheint, dass Geraden auf Geraden sowie Ebenen auf Ebenen abgebildet werden und dass die Parallelität von Geraden und von Ebenen erhalten bleibt. Ein weiteres Beispiel weist allerdings auf einige Einschränkungen hin.

Beispiel 7.32

Durch die Abbildung ϕ mit $\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 & 2 & 0 \\ 2 & 3 & 1 \\ 1 & 1 & 1 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} -4 \\ 3 \\ 1 \end{pmatrix}$ wird der gesamte dreidimensionale Raum in eine einzige Ebene abgebildet, siehe Abb. 7.14. Die Abbildung ϕ unterscheidet sich somit fundamental von den bisher in diesem Abschnitt betrachteten affinen Abbildungen und weist gewisse Analogien

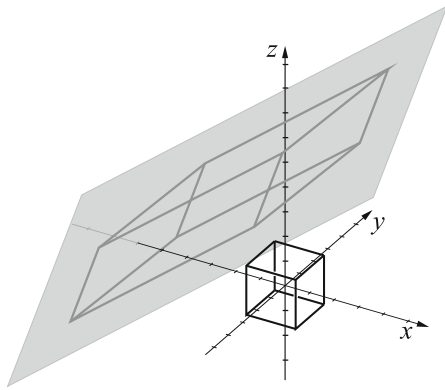


Abb. 7.14: Affine Abbildung des Raumes auf eine Ebene

zu der in dem Beispiel 7.8 auf S. 262 betrachteten Projektion auf. Ein genauerer Blick auf die Abbildungsmatrix legt die Ursachen dafür offen: Die Matrix $\begin{pmatrix} 1 & 2 & 0 \\ 2 & 3 & 1 \\ 1 & 1 & 1 \end{pmatrix}$ ist – im Gegensatz zu den Matrizen aus den Beispielen 7.27–7.31 – *nicht regulär*. Untersucht man ihre Spaltenvektoren, so stellt man fest, dass sich jeder davon als Linearkombination der beiden anderen darstellen lässt, z. B. $\begin{pmatrix} 2 \\ 3 \\ 1 \end{pmatrix} = 2 \cdot \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} - \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$. In dem Abschnitt 7.2.3 wurde festgestellt, dass die Spaltenvektoren der Abbildungsmatrix die Bilder der Basisvektoren des Raumes, der abgebildet wird, angeben. Ihre lineare Hülle beschreibt damit das Bild des abgebildeten Vektorraumes. In unserem Falle ist die lineare Hülle aller drei Spaltenvektoren gleich der linearen Hülle zweier Spaltenvektoren. Die Ebene, in welche ϕ den Raum abbildet (in Abb. 7.14 grau dargestellt), verläuft durch den Punkt, in den der Verschiebungsvektor von ϕ den Koordinatenursprung überführt. Sie hat zwei linear unabhängige Spaltenvektoren der Abbildungsmatrix als Richtungsvektoren. Eine Parameterdarstellung dieser Ebene ist also

$$\varepsilon: \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -4 \\ 3 \\ 1 \end{pmatrix} + \lambda \cdot \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} + \mu \cdot \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$$

7.3.2 Einige Eigenschaften affiner Abbildungen

Der Vergleich des letzten Beispiels mit den vorherigen Beispielen zeigt, dass sich affine Abbildungen stark voneinander unterscheiden können. Im Folgenden werden zunächst einige Gemeinsamkeiten verallgemeinert.

Satz 7.8

Für beliebige affine Räume A, A' und jede affine Abbildung $\phi: A \rightarrow A'$ gilt:

1. ϕ bildet jeden affinen Unterraum von A auf einen affinen Unterraum von A' ab, d. h. ist N ein affiner Unterraum von A , so ist $\phi(N) = \{Q \mid \exists P \in N: \phi(P) = Q\}$ ein affiner Unterraum von A' .
2. Sind N_1 und N_2 affine Unterräume von A , so ist das Bild der Schnittmenge $N_1 \cap N_2$ beider Unterräume eine Teilmenge der Schnittmenge der Bilder der Unterräume N_1 und N_2 : $\phi(N_1 \cap N_2) \subseteq \phi(N_1) \cap \phi(N_2)$.

Wir verzichten auf einen Beweis des Satzes 7.8, merken aber an, dass der Beweis des ersten Teils durch Rückführung auf analoge Eigenschaften linearer Abbildungen geführt werden kann (siehe dazu den Satz 7.1 auf S. 267).

Bemerkung: Die Aussage 2 des Satzes 7.8 klingt auf den ersten Blick etwas „schwach“ – warum ist nicht $\phi(N_1 \cap N_2) = \phi(N_1) \cap \phi(N_2)$? Dies erschließt sich anhand des Beispiels 7.32 oder der in dem Beispiel 7.8 (S. 262) behandelten Projektion. Bei derartigen Abbildungen können z. B. zwei Geraden, die parallel sind oder sich in einem Punkt schneiden, auf dieselbe Gerade abgebildet werden. Bezeichnet man derartige Geraden mit N_1 und N_2 , so kann z. B. $N_1 \cap N_2 = \{S\}$ und $\phi(N_1 \cap N_2) = \phi(S)$, also ein Punkt, aber $\phi(N_1) \cap \phi(N_2)$ eine Gerade sein.

Bijektive affine Abbildungen

Vergleicht man die in diesem Abschnitt behandelten Beispiele miteinander, so fällt auf, dass die Abbildungen aus den Beispielen 7.27–7.31 über einige Eigenschaften verfügen, welche die Abbildung aus dem Beispiel 7.32 nicht besitzt. Die Besonderheit dieser Abbildung besteht darin, dass sie den gesamten dreidimensionalen Raum auf eine einzige Ebene abbildet, sie ist also nicht surjektiv. Zudem bildet sie mehrere Punkte, sogar ganze Geraden, auf einzelne Punkte ab, ist also auch nicht injektiv. Affine Abbildungen, die injektiv und surjektiv sind, besitzen einige Eigenschaften, die bei anderen affinen Abbildungen nicht gegeben sind. Ohne Beweis vermerken wir zunächst den folgenden Satz.

Satz 7.9

Es seien A und A' affine Punkträume mit zugehörigen Vektorräumen V bzw. V' . Eine affine Abbildung $\phi: A \rightarrow A'$ ist genau dann bijektiv, wenn die zugehörige lineare Abbildung $f: V \rightarrow V'$ ein Isomorphismus ist.

Aus den Sätzen 7.9 und 7.5 (S. 275) folgt, dass eine affine Abbildung genau dann bijektiv ist, wenn eine Abbildungsmatrix der zugehörigen linearen Abbildung regulär ist; für affine Abbildungen $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ bzw. $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ betrifft dies die Matrizen in den Abbildungsgleichungen (7.1) bzw. (7.2) auf S. 281.

Bijektive affine Abbildungen „erben“ von den „starken“ Eigenschaften der Isomorphismen (vgl. Abschnitt 7.2.5) z. B. die Existenz von Umkehrabbildungen. Einige weitere Eigenschaften gibt der folgende Satz an.

Satz 7.10

Sind A, A' affine Räume, so gilt für jede bijektive affine Abbildung $\phi: A \rightarrow A'$:

1. *ϕ bildet jeden affinen Unterraum von A auf einen affinen Unterraum gleicher Dimension von A' ab.*
2. *Sind N_1 und N_2 affine Unterräume von A , so ist das Bild der Schnittmenge $N_1 \cap N_2$ beider Unterräume gleich der Schnittmenge der Bilder der Unterräume N_1 und N_2 : $\phi(N_1 \cap N_2) = \phi(N_1) \cap \phi(N_2)$.*
3. *Sind N_1 und N_2 parallele affine Unterräume von A , so sind $\phi(N_1)$ und $\phi(N_2)$ parallele affine Unterräume von A' .*

Da Geraden und Ebenen affine Unterräume sind, besagt der Teil 1 des Satzes 7.10, dass bei bijektiven affinen Abbildungen Geraden immer auf Geraden (und nicht auf Punkte, was in dem Beispiel 7.32 auftritt) sowie Ebenen auf „vollständige“ Ebenen abgebildet werden.

Wir verzichten darauf, den Satz 7.10 zu beweisen, weisen aber auch hier auf die Analogien zu entsprechenden Eigenschaften linearer Abbildungen (speziell Isomorphismen) hin, die aufgrund der Bezüge zwischen linearen und affinen Abbildungen (siehe die Definition 7.7 auf S. 279) sowie zwischen linearen und affinen Unterräumen (nach der Definition 5.12 auf S. 210) bestehen.

7.3.3 Aufgaben zu Abschnitt 7.3

1. Gegeben sind affine Abbildungen $\phi_1, \phi_2: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ und $\phi_3, \phi_4: \mathbb{R}^3 \rightarrow \mathbb{R}^3$:

$$\phi_1 \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} -4 \\ 2 \end{pmatrix}, \quad \phi_2 \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \sqrt{2} & -\sqrt{2} \\ \sqrt{2} & \sqrt{2} \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} 3 \\ -3 \end{pmatrix},$$

$$\phi_3 \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 3 \\ -4 \\ -1 \end{pmatrix}, \quad \phi_4 \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 5 \\ -4 \\ -1 \end{pmatrix}.$$

- Veranschaulichen Sie $\phi_1, \dots, \phi_4$ graphisch durch Punkte und Figuren bzw. Körper und ihre Bilder. Sie können dazu eine Datei für das CAS Maxima von der Internetseite dieses Buches herunterladen, die Vorlagen für die Darstellung affiner Abbildungen in $\mathbb{R}^2$ und in $\mathbb{R}^3$ enthält.
 - Überprüfen Sie ϕ_1, ϕ_2, ϕ_3 und ϕ_4 auf Fixpunkte und berechnen Sie diese (siehe dazu das Beispiel 7.28 auf S. 281).
2. Begründen Sie, dass die in dem Beispiel 7.29 auf S. 282 gegebene affine Abbildung genau einen Fixpunkt besitzt und bestimmen Sie diesen.
3. Untersuchen Sie die Scherung $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $\phi: \begin{pmatrix} x \\ y \end{pmatrix} \mapsto \begin{pmatrix} x-2y \\ y \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \end{pmatrix}$ (siehe das Beispiel 7.7 auf S. 262) auf Fixpunkte.
4. Untersuchen Sie die affine Abbildung ϕ aus dem Beispiel 7.31 auf Fixpunkte.
5. Untersuchen Sie, ob die affinen Abbildungen $\phi_1, \phi_2: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ bijektiv sind. Geben Sie für diejenige(n) dieser Abbildungen, für die das nicht der Fall ist, eine Gleichung der Ebene an, in welche ϕ_1 bzw. ϕ_2 den gesamten Raum abbildet (vgl. das Beispiel 7.32):
- $$\phi_1 \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 1 & 3 & 5 \\ 0 & 2 & 3 \\ 4 & 5 & 2 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 2 \\ -1 \\ 3 \end{pmatrix}, \quad \phi_2 \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} 3 & 1 & 2 \\ 4 & -1 & 7 \\ 2 & 3 & -3 \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} 5 \\ 7 \\ 4 \end{pmatrix}.$$
6. Begründen Sie: Gibt man drei beliebige nicht kollineare Punkte $P, Q, R \in \mathbb{R}^2$ und drei weitere beliebige Punkte $P', Q', R' \in \mathbb{R}^2$ vor, so ist durch $\phi(P) = P', \phi(Q) = Q', \phi(R) = R'$ eindeutig eine affine Abbildung $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ bestimmt.
7. Geben Sie drei Punkte $P, Q, R \in \mathbb{R}^2$ und drei Punkte $P', Q', R' \in \mathbb{R}^2$ an, so dass *keine* affine Abbildung $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ mit $\phi(P) = P', \phi(Q) = Q'$ und $\phi(R) = R'$ existiert.

7.4 Kongruenz- und Ähnlichkeitsabbildungen

Die Beispiele des vorangegangenen Abschnittes haben uns auf einige Eigenschaften affiner Abbildungen geführt. Sie haben aber auch gezeigt, dass affine Abbildungen über einige zentrale Eigenschaften in der Schulgeometrie auftretender Abbildungen im Allgemeinen *nicht* verfügen müssen (d. h. es gibt zwar affine Abbildungen, welche diese Eigenschaften haben, aber auch Gegenbeispiele).

- Affine Abbildungen sind i. Allg. *nicht abstandstreu*, Abstände von Punkten können sich von denen ihrer Bildpunkte unterscheiden.
- Affine Abbildungen können auch *Längenverhältnisse* von Strecken *verändern*, wenn diese nicht auf einer Geraden oder auf parallelen Geraden liegen.
- *Winkel* werden durch affine Abbildungen i. Allg. *nicht auf dazu kongruente Winkel* abgebildet. Auch *senkrechte Geraden* werden i. Allg. *nicht auf senkrechte Geraden* abgebildet.
- Auch *Flächeninhalte und Volumina* bleiben i. Allg. *nicht erhalten*.

Im Folgenden werden wir spezielle affine Abbildungen und die dazugehörigen linearen Abbildungen untersuchen, welche über diese Eigenschaften verfügen.

7.4.1 Isometrien

In der Elementargeometrie werden Kongruenzabbildungen als abstandserhaltende Abbildungen der Ebene oder des Raumes eingeführt. Aus der Abstandstreue lässt sich auch die Winkeltreue ableiten. Um auf der Grundlage der linearen Algebra überhaupt Eigenschaften wie Abstands- und Winkeltreue betrachten zu können, benötigen wir Punkträume (und zugehörige Vektorräume), in denen Längen und Winkel definiert sind. Dies ist in Euklidischen Vektor- und Punkträumen gegeben (vgl. Abschnitt 5.6), da in diesen ein Skalarprodukt existiert. Wir führen die folgenden Betrachtungen in $\mathbb{R}^2$ und $\mathbb{R}^3$ und verwenden das Standardskalarprodukt. Um Kongruenzabbildungen von Punkten beschreiben zu können, wenden wir uns zunächst den zugehörigen Vektorabbildungen zu.

Definition 7.9

Es sei V ein euklidischer Vektorraum. Eine lineare Abbildung $f: V \rightarrow V$ heißt *Isometrie*, falls alle Vektoren $\vec{u} \in V$ auf Vektoren gleicher Beträge abgebildet werden, also $|f(\vec{u})| = |\vec{u}|$ für alle $\vec{u} \in V$ gilt. ♦

Satz 7.11

Jede Isometrie $f: V \rightarrow V$ bildet zwei beliebige Vektoren $\vec{u}, \vec{v} \in V$ auf Vektoren mit demselben Skalarprodukt ab, d. h. $\vec{u} \cdot \vec{v} = f(\vec{u}) \cdot f(\vec{v})$.

Beweis: Es seien $\vec{u}, \vec{v}$ beliebige Vektoren von V . Bei jeder Isometrie $f: V \rightarrow V$ gilt $|f(\vec{u})| = |\vec{u}|$, $|f(\vec{v})| = |\vec{v}|$ und $|f(\vec{u}) + f(\vec{v})| = |f(\vec{u} + \vec{v})| = |\vec{u} + \vec{v}|$ und somit auch $|f(\vec{u})|^2 = |\vec{u}|^2$, $|f(\vec{v})|^2 = |\vec{v}|^2$ sowie $|f(\vec{u}) + f(\vec{v})|^2 = |\vec{u} + \vec{v}|^2$. Wegen $|\vec{u} + \vec{v}|^2 = (\vec{u} + \vec{v}) \cdot (\vec{u} + \vec{v})$ und $|f(\vec{u}) + f(\vec{v})|^2 = (f(\vec{u}) + f(\vec{v})) \cdot (f(\vec{u}) + f(\vec{v}))$ folgt daraus

$$\begin{aligned}
(\vec{u} + \vec{v}) \cdot (\vec{u} + \vec{v}) &= (f(\vec{u}) + f(\vec{v})) \cdot (f(\vec{u}) + f(\vec{v})) \\
\vec{u}^2 + \vec{v}^2 + 2\vec{u} \cdot \vec{v} &= f(\vec{u})^2 + f(\vec{v})^2 + 2f(\vec{u}) \cdot f(\vec{v}) \\
|\vec{u}|^2 + |\vec{v}|^2 + 2\vec{u} \cdot \vec{v} &= |f(\vec{u})|^2 + |f(\vec{v})|^2 + 2f(\vec{u}) \cdot f(\vec{v}) \\
2\vec{u} \cdot \vec{v} &= 2f(\vec{u}) \cdot f(\vec{v}),
\end{aligned}$$

und somit die Behauptung $\vec{u} \cdot \vec{v} = f(\vec{u}) \cdot f(\vec{v})$. $\square$

Als Folgerung aus dem Satz 7.11 ergibt sich nach den Definitionen des Winkels und der Orthogonalität von Vektoren (vgl. Abschnitt 5.6.2) der folgende Satz.

Satz 7.12

Isometrien sind winkeltreu, d. h. für beliebige Vektoren $\vec{u}, \vec{v} \in V$ und ihre Bilder $f(\vec{u}), f(\vec{v})$ bei einer beliebigen Isometrie $f: V \rightarrow V$ gilt $\angle(\vec{u}, \vec{v}) = \angle(f(\vec{u}), f(\vec{v}))$. Insbesondere werden zueinander orthogonale Vektoren auf wiederum zueinander orthogonale Vektoren abgebildet.

In den vorangegangenen Abschnitten haben wir bereits festgestellt, dass sich lineare (und infolge dessen auch affine) Abbildungen effektiv durch Matrizen beschreiben lassen. Wir befassen uns daher im Folgenden mit der besonderen Struktur der Matrizen, welche Isometrien beschreiben. Dazu betrachten wir für eine Isometrie $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Abbildungsmatrix $\mathbf{A}^f = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ bezüglich der Standardbasis $B_0 = \{\vec{e}_1, \vec{e}_2\}$ von $\mathbb{R}^2$ mit $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$. Die Spalten der Matrix $\mathbf{A}^f$ geben die Bilder der Basisvektoren an, siehe Abschnitt 7.2.3. Wegen $\vec{e}_1 \cdot \vec{e}_1 = 1$, $\vec{e}_2 \cdot \vec{e}_2 = 1$ und $\vec{e}_1 \cdot \vec{e}_2 = 0$ muss aufgrund des Satzes 7.11 gelten:

$$\begin{pmatrix} a_{11} \\ a_{21} \end{pmatrix} \cdot \begin{pmatrix} a_{11} \\ a_{21} \end{pmatrix} = 1, \quad \begin{pmatrix} a_{12} \\ a_{22} \end{pmatrix} \cdot \begin{pmatrix} a_{12} \\ a_{22} \end{pmatrix} = 1, \quad \begin{pmatrix} a_{11} \\ a_{21} \end{pmatrix} \cdot \begin{pmatrix} a_{12} \\ a_{22} \end{pmatrix} = 0$$

bzw.

$$\begin{aligned}
a_{11}^2 + a_{21}^2 &= 1 \\
a_{12}^2 + a_{22}^2 &= 1 \\
a_{11} \cdot a_{12} + a_{21} \cdot a_{22} &= 0.
\end{aligned}$$

Um hieraus verallgemeinerungsfähige Eigenschaften der Abbildungsmatrizen von Isometrien entnehmen zu können, betrachten wir das Produkt der Matrix $\mathbf{A}^f$ und ihrer transponierten Matrix (siehe hierzu die Definition 6.2 auf S. 228):

$$(\mathbf{A}^f)^T \circ \mathbf{A}^f = \begin{pmatrix} a_{11} & a_{21} \\ a_{12} & a_{22} \end{pmatrix} \circ \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = \begin{pmatrix} a_{11}^2 + a_{21}^2 & a_{11} a_{12} + a_{21} a_{22} \\ a_{11} a_{12} + a_{21} a_{22} & a_{12}^2 + a_{22}^2 \end{pmatrix}.$$

Wir stellen fest, dass für die Matrix einer Isometrie (für deren Koeffizienten die Gleichungen (*) erfüllt sein müssen), gilt:

$$(\mathbf{A}^f)^T \circ \mathbf{A}^f = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \mathbf{E}_2.$$

Dieses Ergebnis ist umkehr- und verallgemeinerbar.

Satz 7.13

Es sei V ein euklidischer Vektorraum. Eine lineare Abbildung $f: V \rightarrow V$ ist genau dann eine Isometrie, wenn für ihre Abbildungsmatrix $\mathbf{A}^f$ bezüglich einer Orthonormalbasis von V gilt:

$$(\mathbf{A}^f)^T \circ \mathbf{A}^f = \mathbf{A}^f \circ (\mathbf{A}^f)^T = \mathbf{E}.$$

Die eine Richtung des Satzes 7.13 wurde zuvor für Isometrien $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ gezeigt. In $\mathbb{R}^3$ ist der Nachweis in analoger Weise möglich, aber rechnerisch aufwändiger (siehe die Aufgabe 2 auf S. 296). Nicht eingegangen wurde bislang auf den Beweis der anderen Richtung des Satzes 7.13, dass eine lineare Abbildung eine Isometrie ist, wenn sie bezüglich einer Orthonormalbasis eine Abbildungsmatrix $\mathbf{A}^f$ mit der Eigenschaft $(\mathbf{A}^f)^T \circ \mathbf{A}^f = \mathbf{E}$ besitzt. Dies folgt daraus, dass eine lineare Abbildung eine Isometrie ist, wenn sie eine Orthonormalbasis auf eine Orthonormalbasis abbildet, siehe die Aufgabe 1 auf S. 296.

Definition 7.10

Eine Matrix $\mathbf{A}$ mit $\mathbf{A}^T \circ \mathbf{A} = \mathbf{A} \circ \mathbf{A}^T = \mathbf{E}$ heißt *Orthogonalmatrix*. ◆

Satz 7.14

Die Inverse und die Transponierte einer Orthogonalmatrix $\mathbf{A}$ sind identisch: $\mathbf{A}^{-1} = \mathbf{A}^T$.

Der Satz 7.14 folgt unmittelbar aus der Definition 7.10 und der Definition inverser Matrizen, vgl. Abschnitt 6.2.4.

Satz 7.15

Für die Determinante jeder Orthogonalmatrix $\mathbf{A}$ gilt $\det \mathbf{A} = 1$ oder $\det \mathbf{A} = -1$.

Der (sehr kurze) Beweis des Satzes 7.15 sei der Leserin bzw. dem Leser überlassen. Die Determinanten von Abbildungsmatrizen werden sich als wichtiges Klassifikationsmerkmal von Isometrien und Kongruenzabbildungen erweisen.

Isometrien in $\mathbb{R}^2$

Wir ziehen nun aus den diskutierten Eigenschaften von Isometrien geometrische Schlussfolgerungen und betrachten zunächst Isometrien innerhalb von $\mathbb{R}^2$. Für deren Matrizen $\mathbf{A}^f = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$ gelten die Gleichungen (*) auf S. 288. Für die Koeffizienten a_{11} und a_{21} gilt $a_{11}^2 + a_{21}^2 = 1$. Da sich jede reelle Zahl von Null bis Eins als Kosinus und als Sinus einer reellen Zahl α (mit $-\pi < \alpha \leq \pi$) darstellen lässt und für beliebige $\alpha \in \mathbb{R}$ gilt $\sin^2 \alpha + \cos^2 \alpha = 1$, lassen sich die Koeffizienten a mit a_{11} und a_{21} folgendermaßen darstellen:

$$a_{11} = \cos \alpha, \quad a_{21} = \sin \alpha.$$

Analog dazu folgt aus $a_{12}^2 + a_{22}^2 = 1$, dass sich diese Koeffizienten in der Form

$$a_{12} = \sin \beta, \quad a_{22} = \cos \beta$$

darstellen lassen. Setzen wir dies in die Gleichung $a_{11} a_{12} + a_{21} a_{22} = 0$ ein, so erhalten wir:

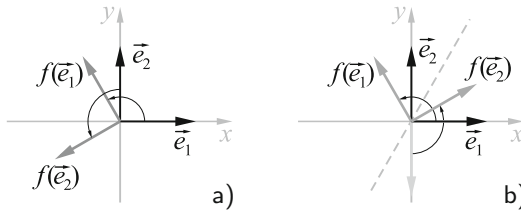
$$\cos \alpha \sin \beta + \sin \alpha \cos \beta = 0.$$

Nach dem Additionstheorem $\cos \alpha \sin \beta + \sin \alpha \cos \beta = \sin(\alpha + \beta)$ folgt daraus $\sin(\alpha + \beta) = 0$. Somit ist (in dem Intervall $]-\pi; \pi]$) $\alpha + \beta = 0$ oder $\alpha + \beta = \pi$. Aus diesen Fällen leiten sich die beiden Arten von Isometrien in $\mathbb{R}^2$ ab.

1. Für $\alpha + \beta = 0$ ist $\beta = -\alpha$ und somit $a_{12} = \sin \beta = -\sin \alpha$, $a_{22} = \cos \beta = \cos \alpha$.

Die Abbildungsmatrix erhält damit die Gestalt

$$\mathbf{A}_1^f = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}.$$

**Abb. 7.15:** Isometrien in $\mathbb{R}^2$

- a) Fall 1: $\mathbf{A}_1^f$
 b) Fall 2: $\mathbf{A}_2^f$

Die Abbildung f ist damit (in geometrischer Interpretation) eine Drehung mit dem Drehwinkel α , siehe auch das Beispiel 7.4 auf S. 260.

2. Für $\alpha + \beta = \pi$ ist $\beta = \pi - \alpha$ und somit $\cos \beta = -\cos \alpha$, $\sin \beta = \sin \alpha$. Die Abbildungsmatrix $\mathbf{A}^f$ nimmt damit die Gestalt

$$\mathbf{A}_2^f = \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix}$$

an. Für eine leichtere geometrische Interpretation können wir sie als Produkt zweier Matrizen darstellen:

$$\mathbf{A}_2^f = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \circ \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

Die durch $\mathbf{A}_2^f$ beschriebene Isometrie lässt sich somit als Spiegelung an der x -Achse und anschließende Drehung mit dem Drehwinkel α interpretieren.

In der Abb. 7.15 sind die beiden möglichen Fälle von Isometrien von $\mathbb{R}^2$ anhand der Vektoren der Standardbasis und ihrer Bilder für $\alpha = 120^\circ$ dargestellt. Zu dem zweiten Fall ist anzumerken, dass die durch $\mathbf{A}_2^f$ beschriebene Nacheinanderausführung einer Spiegelung an der x -Achse und einer anschließenden Drehung mit dem Drehwinkel α auch durch eine Spiegelung an der Geraden, die mit der x -Achse einen Winkel der Größe $\frac{\alpha}{2}$ (in positiver Richtung) einschließt, ersetzt werden kann (siehe die gestrichelte Spiegelachse in der Abb. 7.15 b).

Die beiden durch $\mathbf{A}_1^f$ und $\mathbf{A}_2^f$ beschriebenen Fälle von Isometrien unterscheiden sich hinsichtlich ihrer Orientierungstreue. In Abb. 7.15 ist ersichtlich, dass im Fall 1 die Orientierung der Basisvektoren $\vec{e}_1$ und $\vec{e}_2$ sowie die ihrer Bilder $f(\vec{e}_1)$ und $f(\vec{e}_2)$ gleich ist. Im Fall 2 hat sich die Orientierung hingegen verändert; $f(\vec{e}_1)$ und $f(\vec{e}_2)$ sind in dieser Reihenfolge negativ (also im Uhrzeigersinn) orientiert. Auf S. 255 wurde kurz auf die Orientierung eingegangen und festgestellt, dass diese durch Determinanten beschrieben wird. Wir berechnen daher die Determinanten der Abbildungsmatrizen $\mathbf{A}_1^f$ und $\mathbf{A}_2^f$:

$$\det \mathbf{A}_1^f = \det \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} = \cos^2 \alpha + \sin^2 \alpha = 1,$$

$$\det \mathbf{A}_2^f = \det \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} = -\cos^2 \alpha - \sin^2 \alpha = -1.$$

Dieses Beispiel ist, wie wir ohne Beweis festhalten, für beliebige Isometrien verallgemeinerbar.

Orthogonalmatrizen $\mathbf{A}$ mit $\det \mathbf{A} = 1$ beschreiben orientierungserhaltende (gleichsinnige) Isometrien; ist $\det \mathbf{A} = -1$, so beschreibt $\mathbf{A}^f$ eine ungleichsinnige (nicht orientierungserhaltende) Isometrie.

Isometrien in $\mathbb{R}^3$

Die Beschreibung von Isometrien $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ ist deutlich komplexer als in $\mathbb{R}^2$, grundsätzlich können jedoch auch hier zwei Fälle auftreten.

1. f ist eine Drehung um eine Drehachse im Raum; $\mathbf{A}_1^f = \mathbf{D}$, wobei $\mathbf{D}$ eine Drehmatrix ist, auf die noch näher einzugehen ist.
2. f setzt sich aus einer Spiegelung an einer Koordinatenebene und einer Drehung im Raum zusammen; $\mathbf{A}_2^f = \mathbf{D} \circ \mathbf{S}$.

Spiegelungen an Koordinatenebenen lassen sich einfach durch Matrizen beschreiben. So lässt die Matrix $\mathbf{S}_z = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}$ die x - und y -Komponenten von Vektoren unverändert und multipliziert die z -Komponenten mit -1 . $\mathbf{S}_z$ beschreibt also eine Spiegelung an der x - y -Ebene.

Beliebige *Drehungen im Raum* lassen sich durch Nacheinanderausführungen von Drehungen um die drei Koordinatenachsen kombinieren. Um diese durch Matrizen zu beschreiben, greifen wir auf Matrizen ebener Drehungen zurück und wenden diese jeweils auf zwei Komponenten von Vektoren $\begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathbb{R}^3$ an. Eine Drehung um die z -Achse lässt die z -Komponenten aller Vektoren unverändert und führt mit den x - und y -Komponenten eine Drehung wie in $\mathbb{R}^2$ aus, sie lässt sich also durch die folgende Matrix beschreiben:

$$\mathbf{D}_z = \begin{pmatrix} \cos \alpha_z & -\sin \alpha_z & 0 \\ \sin \alpha_z & \cos \alpha_z & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Analog lassen sich Drehungen um die beiden anderen Koordinatenachsen durch

$$\mathbf{D}_y = \begin{pmatrix} \cos \alpha_y & 0 & \sin \alpha_y \\ 0 & 1 & 0 \\ -\sin \alpha_y & 0 & \cos \alpha_y \end{pmatrix} \quad \text{und} \quad \mathbf{D}_x = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \alpha_x & -\sin \alpha_x \\ 0 & \sin \alpha_x & \cos \alpha_x \end{pmatrix}$$

beschreiben. Dabei sind $\alpha_x, \alpha_y, \alpha_z$ die (orientierten) Drehwinkel um die Koordinatenachsen. Die Vorzeichen in $\mathbf{D}_y$ rühren daher, dass von der y -Achse aus gesehen die z -Achse die erste und die x -Achse die zweite Koordinatenachse ist. Vertauscht man in $\mathbf{D}_y$ die erste und die dritte Zeile sowie die erste und die dritte Spalte, so erscheinen die Vorzeichen in der gewohnten Anordnung.

Die orientierten Drehwinkel α_x, α_y und α_z um die Koordinatenachsen werden als *Eulersche Winkel* bezeichnet.

Es existieren zwei Arten von Isometrien in $\mathbb{R}^3$:

1. Gleichsinnige Isometrien lassen sich als Nacheinanderausführungen von Drehungen um die drei Koordinatenachsen darstellen: $\mathbf{A}_1^f = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x$.
2. Ungleichsinnige Isometrien lassen sich als Nacheinanderausführungen einer Spiegelung an einer Koordinatenebene (z. B. an der x - y -Ebene) und von Drehungen um die drei Koordinatenachsen darstellen: $\mathbf{A}_2^f = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x \circ \mathbf{S}_z$. Statt an der x - y -Ebene könnte auch an einer der beiden anderen Koordinatenebenen gespiegelt werden, wobei dann andere Eulersche Winkel auftreten.

Es ist recht schwer, sich die Eulerschen Drehwinkel und ihre Wirkung auf Isometrien anschaulich vorzustellen. Die Internetseite zu diesem Buch enthält daher eine Datei für das CAS Maxima, in die Sie unterschiedliche Eulersche Drehwinkel eingeben und die dadurch beschriebenen (gleichsinnigen und ungleichsinnigen) Isometrien in interaktiven Darstellungen betrachten können.

7.4.2 Kongruenzabbildungen

Betrachtet man Isometrien in Punkträumen, so gelangt man zu den Kongruenzabbildungen, die mitunter auch als *Bewegungen* bezeichnet werden.

Definition 7.11

Es sei A ein euklidischer Punktraum. Als *Kongruenzabbildung* bezeichnet man eine affine Abbildung $\phi : A \rightarrow A$, deren zugehörige lineare Abbildung eine Isometrie ist. $\blacklozenge$

Bemerkungen:

- Eine alternative Definition für Kongruenzabbildungen wäre: *Eine Kongruenzabbildung ist eine affine Abbildung, die Abstände beliebiger Punktepaares unverändert lässt.* Die Äquivalenz zu der Definition 7.11 ergibt sich unmittelbar aus der Definition 7.9 (Isometrie) und der Tatsache, dass Abstände von Punkten als Beträge ihrer Verbindungsvektoren definiert sind.
- Mithilfe von Kongruenzabbildungen lässt sich nun der Begriff *Kongruenz* definieren: Zwei Punktmengen (bzw. geometrische Figuren) F_1 und F_2 heißen *kongruent*, falls eine Kongruenzabbildung existiert, die F_1 auf F_2 abbildet.

Allgemeine *Abbildungsgleichungen für Kongruenzabbildungen* in $\mathbb{R}^2$ erhält man, indem man in der Abbildungsgleichung (7.1) für affine Abbildungen $\phi : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ (siehe S. 281) die Abbildungsmatrizen beliebiger linearer Abbildungen durch die auf S. 289f. hergeleiteten Matrizen von Isometrien $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ ersetzt. Wir erhalten auf diese Weise zwei Arten von Kongruenzabbildungen:

- *Gleichsinnige Kongruenzabbildungen* werden durch Abbildungsvorschriften

$$\phi \begin{pmatrix} x \\ y \end{pmatrix} = \mathbf{A}_1^f \circ \vec{x} + \vec{v} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix}$$

beschrieben, Abb.7.16 a) zeigt ein Beispiel mit $\alpha = 60^\circ$ und $\vec{v} = \begin{pmatrix} 1 \\ 3 \end{pmatrix}$.

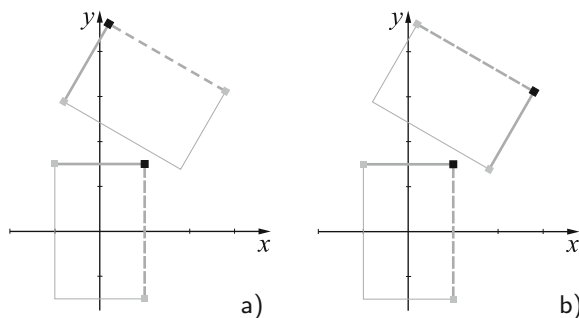


Abb. 7.16:
Kongruenzabbildungen in $\mathbb{R}^2$
a) gleichsinnig
b) ungleichsinnig

- Abbildungsvorschriften *ungleichsinniger Kongruenzabbildungen* haben die Form

$$\phi \begin{pmatrix} x \\ y \end{pmatrix} = \mathbf{A}_2^f \circ \vec{x} + \vec{v} = \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix}.$$

Abb.7.16 b) zeigt hierfür ein Beispiel mit $\alpha = 60^\circ$ und $\vec{v} = \begin{pmatrix} 1 \\ 3 \end{pmatrix}$; man beachte die vertauschte Orientierung der markierten Strecken.

Für die Beschreibung von Kongruenzabbildungen im Raum setzt man die auf S. 291 hergeleiteten Darstellungen von Isometrien in $\mathbb{R}^3$ in die allgemeine Abbildungsgleichung (7.2) affiner Abbildungen $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ (siehe S. 281) ein und erhält für *gleichsinnige Kongruenzabbildungen*

$$\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \vec{v}$$

sowie für *ungleichsinnige Kongruenzabbildungen*

$$\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x \circ \mathbf{S}_z \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \vec{v}$$

mit

$$\mathbf{D}_z = \begin{pmatrix} \cos \alpha_z & -\sin \alpha_z & 0 \\ \sin \alpha_z & \cos \alpha_z & 0 \\ 0 & 0 & 1 \end{pmatrix},$$

$$\mathbf{D}_y = \begin{pmatrix} \cos \alpha_y & 0 & \sin \alpha_y \\ 0 & 1 & 0 \\ -\sin \alpha_y & 0 & \cos \alpha_y \end{pmatrix},$$

$$\mathbf{D}_x = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \alpha_x & -\sin \alpha_x \\ 0 & \sin \alpha_x & \cos \alpha_x \end{pmatrix} \text{ sowie}$$

$$\mathbf{S}_z = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}.$$

In der Abb. 7.17 sind eine gleichsinnige und eine ungleichsinnige Kongruenzabbildung dargestellt, die Eulerschen Winkel sind in beiden Fällen $\alpha_z = 100^\circ$, $\alpha_y = -90^\circ$, $\alpha_x = 45^\circ$, der Verschiebungsvektor ist $\vec{v} = \begin{pmatrix} 3 \\ 2 \\ 2 \end{pmatrix}$.

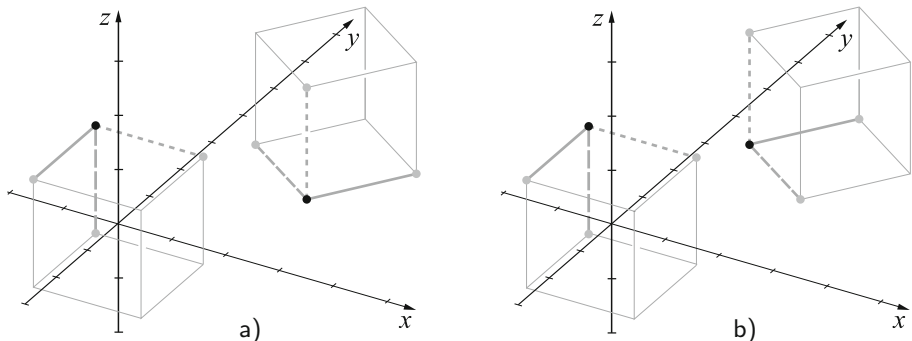


Abb. 7.17: Kongruenzabbildungen in $\mathbb{R}^3$: a) gleichsinnig, b) ungleichsinnig

7.4.3 Ähnlichkeitsabbildungen

Bereits in der Schule werden Ähnlichkeitsabbildungen als Nacheinanderausführungen von zentrischen Streckungen und Kongruenzabbildungen eingeführt. Zentrische Streckungen mit dem Streckfaktor k und dem Streckzentrum im Koordinatenursprung lassen sich durch Matrizen $\mathbf{Z} = \begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix}$ bzw. $\mathbf{Z} = \begin{pmatrix} k & 0 & 0 \\ 0 & k & 0 \\ 0 & 0 & k \end{pmatrix}$ (in $\mathbb{R}^2$ bzw. $\mathbb{R}^3$, jeweils mit $k \neq 0$) beschreiben (siehe dazu auch das Beispiel 7.5 auf S. 261).

Interessant ist ein Blick auf zentrische Streckungen mit negativen Streckfaktoren k hinsichtlich ihrer *Orientierungstreue*. Zentrische Streckungen mit $k < 0$ sind in $\mathbb{R}^2$ gleichsinnige Abbildungen (siehe Abb. 7.18 a), sie lassen sich als Streckungen mit dem Streckfaktor $|k|$ und anschließende Drehungen um 180° bzw. Punktspiegelungen auffassen. Hingegen sind zentrische Streckungen mit $k < 0$ in $\mathbb{R}^3$ ungleichsinnige Abbildungen. Dies lässt sich anhand von Abb. 7.18 b) unter Heranziehung der Rechte-Hand-Regel (siehe S. 133) anschaulich nachvollziehen, aber auch aus der Tatsache ableiten, dass Determinanten von Matrizen orientierungserhaltender Abbildungen positiv, bei ungleichsinnigen Abbildungen hingegen negativ sind. Für $k < 0$ ist

$$\det \begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix} = k^2 > 0, \quad \det \begin{pmatrix} k & 0 & 0 \\ 0 & k & 0 \\ 0 & 0 & k \end{pmatrix} = k^3 < 0.$$

Abbildungsvorschriften für Ähnlichkeitsabbildungen $\phi: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ lassen sich aus denjenigen für Kongruenzabbildungen durch Kombination mit einer zentrischen Streckung ableiten. Man erhält

$$\blacksquare \quad \phi \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \circ \begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix}$$

für gleichsinnige Ähnlichkeitsabbildungen sowie

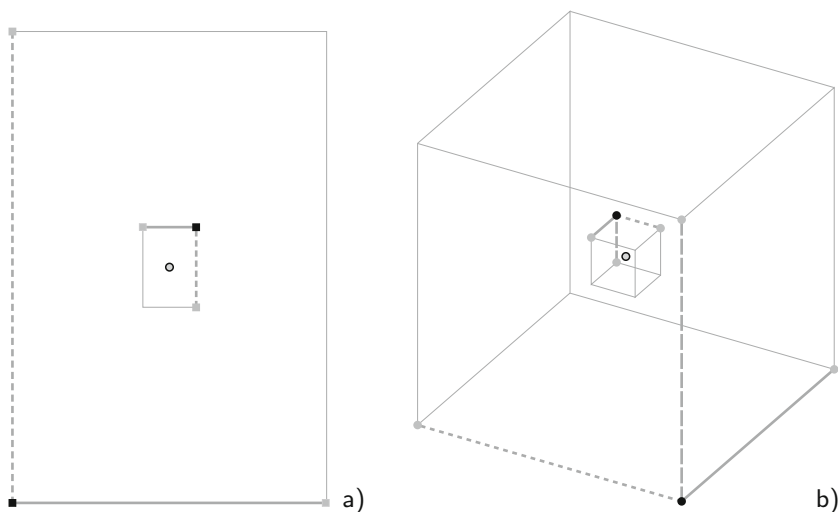


Abb. 7.18: Zentrische Streckungen mit dem Streckungsfaktor $k = -6$ in $\mathbb{R}^2$ und $\mathbb{R}^3$

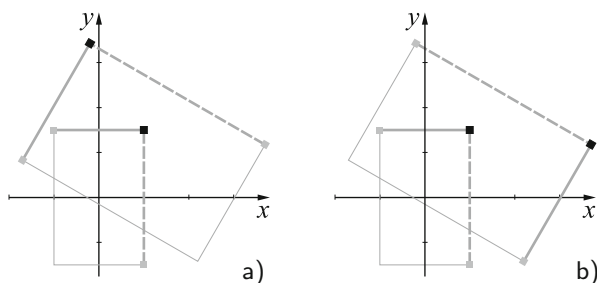


Abb. 7.19:
Ähnlichkeitsabbildungen
in $\mathbb{R}^2$

- a) gleichsinnig
b) ungleichsinnig

$$\blacksquare \quad \phi \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} \circ \begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix}$$

für ungleichsinnige Ähnlichkeitsabbildungen.

In Abb. 7.19 sind eine gleich- und eine ungleichsinnige Ähnlichkeitsabbildung in $\mathbb{R}^2$, jeweils mit $\alpha = 60^\circ$, $\vec{v} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $k = 1,5$, dargestellt.

Um Ähnlichkeitsabbildungen $\phi: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ zu beschreiben, müssen nicht zwei verschiedene Abbildungsvorschriften angegeben werden. Durch

$$\phi \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x \circ \begin{pmatrix} k & 0 & 0 \\ 0 & k & 0 \\ 0 & 0 & k \end{pmatrix} \circ \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \vec{v}$$

mit den auf S. 291 sowie auf S. 293 angegebenen Drehmatrizen $\mathbf{D}_z$, $\mathbf{D}_y$, $\mathbf{D}_x$ lassen sich sowohl gleich- als auch ungleichsinnige Ähnlichkeitsabbildungen beschreiben. Mit $k > 0$ ist ϕ gleichsinnig, mit $k < 0$ hingegen ungleichsinnig. In der Abb. 7.20 sind hierfür Beispiele mit jeweils $\alpha_z = 30^\circ$, $\alpha_y = -45^\circ$, $\alpha_x = 120^\circ$ und $\vec{v} = \begin{pmatrix} 4 \\ 3 \\ 2 \end{pmatrix}$ sowie $k = 2$ in Abb. 7.20 a) und $k = -2$ in Abb. 7.20 b) dargestellt.

Mithilfe einer Datei für das CAS Maxima, die auf der Internetseite zu diesem Buch zur Verfügung steht, lassen sich diese Graphiken interaktiv aus verschiedenen Perspektiven betrachten, vor allem aber lassen sich Werte verändern und dadurch andere Ähnlichkeitsabbildungen untersuchen, siehe dazu auch die Aufgabe 11 auf S. 296.

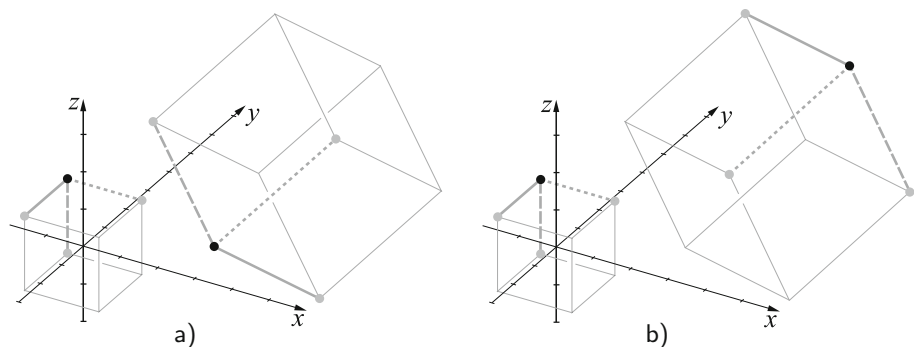


Abb. 7.20: Ähnlichkeitsabbildungen in $\mathbb{R}^3$: a) gleichsinnig, b) ungleichsinnig

7.4.4 Aufgaben zu Abschnitt 7.4

- Es seien V ein euklidischer Vektorraum, $B = \{\vec{b}_1; \dots; \vec{b}_n\}$ eine Orthonormalbasis von V sowie $f: V \rightarrow V$ eine lineare Abbildung. Zeigen Sie dass f genau dann eine Isometrie ist, wenn das Bild $B' = \{f(\vec{b}_1); \dots; f(\vec{b}_n)\}$ der Basis B ebenfalls eine Orthonormalbasis von V ist.
- Für eine lineare Abbildung $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ ist die Matrix $\mathbf{A}^f = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$ bezüglich der Standardbasis $B_0 = \{\vec{e}_1; \vec{e}_2; \vec{e}_3\}$ von $\mathbb{R}^3$ gegeben. Zeigen Sie, dass f genau dann eine Isometrie ist, wenn $(\mathbf{A}^f)^T \circ \mathbf{A}^f = \mathbf{E}_3$ gilt, das Produkt der Abbildungsmatrix mit ihrer Transponierten also die Einheitsmatrix ist.
Hinweis: Orientieren Sie sich an der Herleitung dieses Zusammenhangs für Abbildungen in $\mathbb{R}^2$, siehe S. 288.
- Beweisen Sie den Satz 7.15 auf S. 289.
- Weisen Sie nach, dass jede Isometrie ein Isomorphismus ist.
- Zeigen Sie, dass die Nacheinanderausführung der durch
$$\mathbf{A}_\alpha = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \quad \text{und} \quad \mathbf{A}_\beta = \begin{pmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{pmatrix}$$
 beschriebenen Isometrien unabhängig von der Reihenfolge der Ausführung eine Drehung mit dem Drehwinkel $\alpha + \beta$ ist.
- Untersuchen Sie die beiden auf S. 289f. betrachteten Arten von Isometrien in $\mathbb{R}^2$ für einen Drehwinkel $\alpha = 30^\circ$ auf Fixvektoren, d. h. auf Vektoren $\begin{pmatrix} x \\ y \end{pmatrix}$ mit $\begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}$ bzw. $\begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} \circ \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix}$.
- Auf S. 291 wurden Isometrien in $\mathbb{R}^3$ durch Matrizen $\mathbf{A}_1^f = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x$ bzw. $\mathbf{A}_2^f = \mathbf{D}_z \circ \mathbf{D}_y \circ \mathbf{D}_x \circ \mathbf{S}_z$ beschrieben. Zeigen Sie, dass für beliebige Eulersche Winkel $\alpha_x, \alpha_y, \alpha_z$ gilt: $\det \mathbf{A}_1^f = 1$ und $\det \mathbf{A}_2^f = -1$.
- Begründen Sie, dass jede Kongruenzabbildung bijektiv ist.
- Bestimmen Sie (falls vorhanden) den Fixpunkt bzw. die Fixpunkte der in der Abbildung 7.16 a) auf S. 292 dargestellten gleichsinnigen Kongruenzabbildung mit $\alpha = 60^\circ$, $\vec{v} = \begin{pmatrix} 1 \\ 3 \end{pmatrix}$; interpretieren Sie das Ergebnis geometrisch.
- Bestimmen Sie jeweils (falls vorhanden) den Fixpunkt bzw. die Fixpunkte der beiden in der Abbildung 7.19 auf S. 295 dargestellten Ähnlichkeitsabbildungen mit $\alpha = 60^\circ$, $\vec{v} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $k = 1, 5$.
- Stellen Sie gleichsinnige und ungleichsinnige Kongruenz- und Ähnlichkeitsabbildungen der Ebene und des Raumes für verschiedene Drehwinkel, Verschiebungsvektoren sowie (bei Ähnlichkeitsabbildungen) Streckfaktoren graphisch dar. Sie können dazu vorbereitete Dateien für das CAS Maxima nutzen, die auf der Internetseite dieses Buches zur Verfügung stehen.

Literaturverzeichnis

- Anton, H. (1998): Lineare Algebra. Spektrum, Heidelberg.
- Arens, T., Hettlich, F., Karpfinger, C., Kockelkorn, U., Lichtenegger, K., Stachel, H. (2008): Mathematik. Spektrum, Heidelberg.
- Bär, G. (2001): Geometrie. Eine Einführung für Ingenieure und Naturwissenschaftler. Teubner, Stuttgart.
- Beutelspacher, A. (2010): Lineare Algebra. 7. Auflage. Vieweg+Teubner, Wiesbaden.
- Bosch, S. (2008): Lineare Algebra. 4. Auflage. Springer, Heidelberg.
- Filler, A. (2007): Herausarbeiten funktionaler und dynamischer Aspekte von Parameterdarstellungen durch die Erstellung von Computeranimationen. *Mathematische Semesterberichte* 54, 155–176.
- Filler, A. (2008): 3D-Computergrafik und die Mathematik dahinter: Einbeziehung von Elementen der Computergrafik in den Mathematikunterricht der Sekundarstufe II im Stoffgebiet Analytische Geometrie. VDM Verlag, Saarbrücken.
- Fischer, G. (2001): Analytische Geometrie. 7. Auflage. Vieweg, Braunschweig.
- Fischer, G. (2010): Lineare Algebra. 17. Auflage. Vieweg+Teubner, Wiesbaden.
- Griesel H., Postel H. (Hrsg.) (1986): Mathematik heute. Leistungskurs Lineare Algebra/ Analytische Geometrie. Schroedel Schöningh, Hannover.
- Jänich, K. (2008): Lineare Algebra. 11. Auflage. Springer, Heidelberg.
- Kaufmann, St.-H. (2011): Der Parameter in der Mathematik – die Geschichte einer untergeordneten Variablen? In: Folkerts, M. (Hrsg.): ALGORISMUS. Studien zur Geschichte der Mathematik und der Naturwissenschaften, Heft 75. Dr. Erwin Rauner Verlag, Augsburg, 115-123.
- Kowalsky, H.-J., Michler, G. O. (2003): Lineare Algebra. 12. Auflage. Walter de Gruyter, Berlin.
- Lehmann, E. (1983): Lineare Algebra mit dem Computer. Teubner, Stuttgart.
- Lehmann, I., Schulz, W. (2007): Mengen – Relationen – Funktionen. 3. Auflage. Teubner, Wiesbaden.
- Lind, D. (1997): Koordinaten, Vektoren, Matrizen. Spektrum, Heidelberg.
- Modler, F., Kreh, M. (2010): Tutorium Analysis 1 und Lineare Algebra 1. Spektrum, Heidelberg.
- Muthsam, H. J. (2006): Lineare Algebra und ihre Anwendungen. Spektrum, Heidelberg.
- Nyman M. (2004): 4 Farben – 1 Bild. 4. Auflage. Springer, Berlin.
- Padberg, F., Kütting, H. (1991): Lineare Algebra. Eine elementare Einführung. BI-Wissenschaftsverlag, Mannheim.
- Scheid, H., Schwarz, W. (2009): Elemente der Linearen Algebra und der Analysis. Spektrum, Heidelberg.
- Schupp H. (2000): Kegelschnitte. Franzbecker, Hildesheim.
- Tietze, U.-P., Klika, M., Wolpers, H. (2000): Mathematikunterricht in der Sekundarstufe II. Bd. 2: Didaktik der Analytischen Geometrie und Linearen Algebra. Vieweg, Braunschweig.
- Vollrath, H.-J., Weigand, H.-G. (2006): Algebra in der Sekundarstufe. 3. Auflage. Spektrum, Heidelberg.
- Wittmann, G. (2003): Schülerkonzepte zur Analytischen Geometrie. Franzbecker, Hildesheim.
- Wittmann, G. (2008): Elementare Funktionen und ihre Anwendungen. Spektrum, Heidelberg.

Index

- Abbildung, 257–296
 - Ähnlichkeitsabbildung, 294–295
 - affine, 279–287
 - auf, 283
 - axiale Streckung, 261, 282
 - bijektive, 258, 275, 285
 - Definition, 258
 - Drehung, 260, 289–291
 - Fixpunkt, 281
 - geometrische, 259–264, 279–284, 292–295
 - gleichsinnige, 290
 - injektive, 258, 265
 - inverse, *siehe* Umkehrabbildung
 - Isometrie, 287–292
 - Kongruenzabbildung, 292–293
 - lineare, 264–279
 - Bild, 278
 - Hauptsatz über lineare Abbildungen, 268
 - Isometrie, 287–292
 - Isomorphismus, 274–277, 285
 - Kern, 278
 - matrizielle Darstellung, 270–273
 - Nacheinanderausführung, 260, 273–274
 - orientierungserhaltende, 290
 - Parallelprojektion, 262
 - Scherung, 262, 283
 - Schubspiegelung, 259, 279
 - Spiegelung, 259, 290
 - surjektive, 258, 265, 283
 - Umkehrabbildung, 275
 - ungleichsinnige, 290
 - Verschiebung, 259, 280, 281
 - winkeltreue, 288
 - zentrische Streckung, 261, 281, 294
- Abstand
 - Gerade-Ebene, 163
 - Punkt-Ebene, 162
 - Punkt-Gerade im Raum, 163
 - Punkt-Gerade in der Ebene, 57, 161
 - von Ebenen, 163
 - von Geraden, 164
 - von Punkten, 55, 222
- Additionstheoreme, 260
- Additionsverfahren, 9
- Additivität, 219
- Ähnlichkeitsabbildung, 294–295
- Äquivalenzrelation, 89
- Äquivalenzumformungen, 11, 232, 233
- affine Abbildung, *siehe* Abbildung
- affine Hülle, 212
- affiner Raum, 203–219
 - Dimension, 210
- affiner Unterraum, *siehe* Unterraum
- Anschauungsraum, 204
- Arbeit
 - mechanische, 131
- archimedische Spirale, 148
- Austauschsatz
 - von Steinitz, 195
- Basis, 187–197
 - geordnete, 270, 271
 - kanonische Basis, 187
 - Koordinaten von Vektoren, 188
 - Orthonormalbasis, 224
 - Standardbasis, 187, 269
 - Steinitzscher Austauschsatz, 195
 - Verkürzungssatz, 192
- Basisergänzungssatz, 193
- Betrag eines Vektors, 124, 222
- Bijektivität, 258, 274, 275, 285
- Bilinearform (positiv definite, symmetrische), 219
- CAS, *siehe* Computeralgebrasystem
- Cauchy-Schwarzsche Ungleichung, 124, 222
- Computeralgebrasystem, *siehe* Maxima, 44–48, 137–138, 256
- Determinante, 251–255
 - einer Orthogonalmatrix, 289
- Diagonalform, 24
- Differentialoperator, 267
- Dimension
 - eines affinen Raumes, 210
 - eines Unterraumes, 199
 - eines Vektorraumes, 197–201
- Dimensionsformel
 - für affine Unterräume, 213
 - für lineare Unterräume, 199
- Drehung, 260, 289, 290
 - im Raum, 291
- Dreiecksform, 23
- Dreiecksungleichung, 222
- Durchschnitt linearer Unterräume, 176
- Ebene, 151–154
 - als affiner Unterraum, 210
 - geneigte, 131
 - Gleichung in Hessescher Normalform, 161
 - im vierdimensionalen Raum, 213
 - Koordinatengleichung, 151
 - Lagebeziehungen, 152–154
 - Normalengleichung, 157
 - Normalenvektor, 157
 - Parameterdarstellung, 151
 - Richtungsvektoren, 151
 - Spannvektoren, 151

- Stützvektor, 151
- Einheitskreis, 147
- Einheitsmatrix, 240, 245
- Einheitsvektor, 161
- Einselement, 240
- Einsetzungsverfahren, 9
- Eliminationsverfahren, *siehe*
 - Gauß-Algorithmus
- Ellipse, 70–84
 - Brennpunkte, 72
 - Gleichung in achsenparalleler Lage, 80
 - Gleichung in Hauptachsenlage, 76
 - Hauptachse, 73
 - Mittelpunktsgleichung, 76
 - Nebenachse, 73
 - Ortsdefinition, 72
 - Tangenten, 81
- Erzeugendensystem, 179–181, 190–192
- Erzeugnis, *siehe* lineare Hülle
- Euklidischer Punktraum, 219–226
- Euklidischer Vektorraum, 219–226
- Eulersche Winkel, 291, 293

- Farbtripel, 101
- Farbvektoren, 101
- Farbwürfel, 108
- Fibonacci-Folge, 201, 241
- Fibonacci-Zahlen, 201, 241
- Fixpunkt, 281
- Flächeninhalt, 135
- Folge, 169, 198
- Funktion, 258
 - differenzierbare, 174, 267
 - lineare, 264, 279
 - proportionale, 264
 - reellwertige, 170
 - stetige, 174

- Gauß-Algorithmus, 22
 - Matrixschreibweise, 29
- Gauß-Jordan-Algorithmus, 23
- Gegenvektor, 110, 168
- geneigte Ebene, 131
- Gerade, 50–54
 - Achsenabschnittsform, 51
 - als affiner Unterraum, 210
 - Gleichung in Hessescher Normalform, 161
 - Koordinatengleichung, 50
 - Koordinatengleichungen, 142
 - Normalengleichung, 156
 - Normalenvektor, 156
 - Parameterdarstellung, 140
 - Richtungsvektor, 141
 - Schnittwinkel zwischen Geraden, 53
 - Stützvektor, 141
 - Steigungswinkel, 53
 - Winkel zwischen Geraden, 53, 159

- Zweipunkteform, 51
- Geschwindigkeit, 86
- Gleichsetzungsverfahren, 8
- Gleichungssystem
 - lineares, 1–48, 214–218, 248
 - Äquivalenzumformungen, 11
 - Diagonalform, 24
 - Dreiecksform, 23
 - Gauß-Algorithmus, 22
 - Gauß-Jordan-Algorithmus, 23
 - homogenes, 37, 173, 191
 - inhomogenes, 37, 204, 215
 - Koeffizientenmatrix, 30
 - Lösbarkeitskriterium, 34
 - Lösung durch Matrizeninversion, 248
 - Lösung mit CAS, 44
 - Matrixschreibweise, 29
 - Rang, 32, 35
 - Zeilenstufenform, 23
- Gruppe, 94, 104, 168, 247

- Hauptsatz über lineare Abbildungen, 268
- Hessesche Normalform, 161
- Hintereinanderausführung, *siehe*
 - Nacheinanderausführung von Abbildungen
- Homogenität, 219
- Hülle, *siehe* lineare Hülle
- Hyperbel, 70–84, 209
 - Brennpunkte, 74
 - Gleichung in achsenparalleler Lage, 80
 - Gleichung in Hauptachsenlage, 78
 - Mittelpunktsgleichung, 78
 - Ortsdefinition, 74
 - Tangenten, 81
- Hyperebene, 18

- Injektivität, 258, 265
- Inversion von Matrizen, 245–249
- Isometrie, 287–292
 - gleichsinnige, 290
 - orientierungserhaltende, 290
 - ungleichsinnige, 290
- Isomorphie, 109, 276
- Isomorphismus, 274–277, 285

- kanonische Basis, 187
- Kegel, 69
- Kegelschnitte, 70, 72–84
- Kern, 278
- Kirchhoffsches Gesetz, 42
- Knotenregel, 42
- Koeffizientenmatrix, 30
 - Rang, 32, 35
- Kollinearität, 116
- Komplanarität, 116

- Komposition, *siehe*
 - Nacheinanderausführung von Abbildungen
- Kongruenzabbildung, 292–293
- konische Spirale, 149
- Koordinaten
 - affine, 205–209
 - von Vektoren bezüglich Basen, 188
- Koordinatengeometrie, 49–84
- Koordinatensystem, 205–209
 - affines, 205
 - kartesisches, 106, 224
- Koordinatentransformation, 79, 80, 207
- Kosinussatz, 223
- Kraft, 87
- Kreis
 - Einheitskreis, 147
 - Konstruktion, 60
 - Kreisgleichung, 58
 - Parameterdarstellung, 147
 - Tangenten, 64
- Kreuzprodukt, *siehe* Vektorprodukt
- Kugel, 67
- Kurven zweiter Ordnung, *siehe* Kegelschnitte

- Lagebeziehungen
 - affiner Unterräume, 210
 - Gerade-Ebene, 152
 - Kreis-Gerade, 61
 - von Ebenen, 154
 - von Ebenen im vierdimensionalen Raum, 213
 - von Geraden, 144
- LGS, *siehe* Gleichungssystem
- lineare Abbildung, *siehe* Abbildung
- lineare Abhängigkeit, 182–185
- lineare Exzentrizität, 73
- lineare Hülle, 178
- lineare Unabhängigkeit, 182–185, 190, 194
- linearer Unterraum, *siehe* Unterraum
- Lineares Gleichungssystem, *siehe* Gleichungssystem
- Linearkombination, 113
 - Anwendungen in der Geometrie, 118
- Lösbarkeitskriterium, 34
- logarithmische Spirale, 148

- magisches Quadrat, *siehe* Quadrat
- Materialverflechtung, 235–238
- Matrix, 29–36, 171, 227–256
 - Addition, 171, 240
 - Determinante, 251–255, 289
 - Diagonalmatrix, 229, 233, 254
 - Dreiecksmatrix, 254
 - einer linearen Abbildung, 270–273
 - Einheitsmatrix, 240
 - inverse, 245–249
 - Koeffizienten, 228
 - Koeffizientenmatrix, 30
 - Rang, 32, 35
 - magisches Quadrat, 197, 216
 - Materialverflechtungs-, 235–238
 - Matrizenmultiplikation, 235–241
 - Orthogonalmatrix, 289
 - Populationsmatrix, 242–245
 - Produktmatrix, 238
 - quadratische, 229
 - Rang, 32, 228–234
 - rechtsinverse, 245
 - reguläre, 234, 247, 251
 - skalare Multiplikation, 171, 240
 - Spaltenindex, 228
 - Spaltenrang, 229–234
 - Spaltenumformungen, 232
 - Spaltenvektor, 228
 - Spaltenvertauschung, 252
 - Streckungsmatrix, 261
 - symmetrische, 229
 - transponierte, 228
 - Vektorraum der $m \times n$ -Matrizen, 171
 - Zeilenindex, 228
 - Zeilenrang, 229–234
 - Zeilenstufenform, 229
 - Zeilenumformungen, 232
 - Zeilenvektor, 228
 - Zeilenvertauschung, 252
- Matrizenmultiplikation, 235–241
- Maxima, 44–48, 137–138, 256
 - Ähnlichkeitsabbildungen, 296
 - affine Abbildungen, 286
 - Berechnung von Schnittpunkten, 63
 - Darstellung von Ellipsen, 78
 - Darstellung von Kreisen, 58
 - Darstellung von Kugeln, 68
 - Darstellung von Vektoren durch Pfeile, 137
 - determinant, 256
 - draw2d, 47, 58
 - draw3d, 48
 - invert, 256
 - Kongruenzabbildungen, 296
 - Lösen linearer Gleichungssysteme, 44–48
 - Linearkombinationen, 138
 - Matrizenrechnung, 256
 - Parameterdarstellungen von Kurven, 150
 - rank, 256
 - Skalarprodukt, 138
 - solve, 45
 - transpose, 256
 - Vektorprodukt, 138
 - Vektorrechnung, 137
- mechanische Arbeit, 131
- Mittelpunkt einer Strecke, 55

- n -Tupel, 103–105
- Nacheinanderausführung von
 - Abbildungen, 260, 273–274
- Normaleneinheitsvektor, 161
- Normalengleichung, 156–158
- Normalenvektor
 - einer Ebene, 157
 - einer Geraden, 156
- Nullvektor, 110, 168
- Orientierung, 255
- Orthogonalität, 126, 222
- Orthogonalmatrix, 289
- Orthonormalbasis, 224
- Ortsdefinition
 - Ellipse, 72
 - Hyperbel, 74
 - Parabel, 75
- Parabel, 70–84, 207
 - Gleichung in achsenparalleler Lage, 80
 - Leitlinie, 75
 - Ortsdefinition, 75
 - Scheitelpunkt, 75
 - Scheitelpunktsgleichung, 79
 - Tangenten, 82
- Parallelepiped, 115, 136
- parallelgleich, 89
- Parallelotop, 255
- Parallelprojektion, 262
- Parameterdarstellung
 - archimedische Spirale, 148
 - Ebene, 151
 - Gerade, 140
 - konische Spirale, 149
 - Kreis, 147
 - logarithmische Spirale, 148
 - schräger Wurf, 146
 - Schraubenlinie, 149
- Parametergleichung, *siehe*
 - Parameterdarstellung
- Pfeilklassen, 89–99
 - Addition, 91
 - Repräsentantenunabhängigkeit, 92
 - im Koordinatensystem, 106
 - Multiplikation mit reellen Zahlen, 95
 - Repräsentantenunabhängigkeit, 96
 - Repräsentant, 90
- Polynom, 170, 174, 181, 189, 220, 267
- Populationsmatrizen, 242–245
- positiv definite symmetrische
 - Bilinearform, 219
- Preisvektor, 122
- Produktmatrix, 238
- Punktraum, *siehe* affiner Raum
- Quadrat
 - magisches, 174, 197, 216
- Rang
 - einer Koeffizientenmatrix, 32–36
 - einer Matrix, 228–234
 - eines linearen Gleichungssystems, 32–36
- Raum
 - affiner, *siehe* affiner Raum
- Raumdiagonale im Quader, 129
- Rechenmauern, 202
- Rechte-Hand-Regel, 133, 294
- Rechtssystem, 133, 255
- Regel von Sarrus, 253
- Regularität
 - von Matrizen, 234, 247, 251
- RGB-Vektoren, 101
- RGB-Würfel, 108
- Sarrussche Regel, 253
- Satz des Pythagoras, 223
- Satz des Thales, 130
- Satz von Kronecker-Capelli, 34
- Satz von Varignon, 118
- Scherung, 262, 283
- Schnittgerade zweier Ebenen, 154
- Schnittpunkte
 - Gerade-Ebene, 152
 - von Geraden, 144
- Schnittwinkel
 - zwischen Ebenen, 160
 - zwischen Geraden im Raum, 159
 - zwischen Geraden in der Ebene, 53
 - zwischen Geraden und Ebenen, 160
- schräger Wurf, 146
- Schraubenlinie, 149
- Schubspiegelung, 259, 279
- Schwerpunkt, 119, 120
- Seitenhalbierende, 119
- Sinussatz, 130
- Skalarprodukt, 122–131, 219–221
 - Darstellung bezüglich einer Basis, 221
 - geometrische Deutung, 124
 - Rechenregeln, 123
 - verallgemeinertes, 219
- Spaltenindex, 228
- Spaltenrang, 229–234
- Spaltenumformungen, 232
- Spaltenvektor, 228
- Spat, 115, 136, 255
- Spatprodukt, 136, 255
- Spiegelung, 259, 290
- Spirale, 148
- Standardbasis, 187, 269, 288
- Steinitzcher Austauschatz, 195
- Streckung
 - axiale, 261, 282
 - zentrische, 261, 281, 294
- Streckungsmatrix, 261
- Strukturgleichheit, 109, 110, 274, 276
- Strukturtreue, 264

- Stückliste, 100
- Summe linearer Unterräume, 176
- Surjektivität, 258, 265, 283
- Tangente
 - Ellipsen-, 81
 - Hyperbel-, 81
 - Kreis-, 64
 - Parabel-, 82
- Teilraum, *siehe* Unterraum
- Teilvektorraum, *siehe* Unterraum
- Transponierte, 228
- Umkehrabbildung, 275
- Unterraum
 - affiner, 210–214
 - Dimension, 210
 - Ebene, 210
 - Gerade, 210
 - Lösungsmengen inhomogener linearer Gleichungssysteme, 215
 - magische Quadrate, 216
 - Verbindungsraum, 212
 - linearer, 172–177
 - Basis, 190
 - differenzierbarer Funktionen, 174, 267
 - Dimensionsformel, 199
 - Durchschnitt linearer Unterräume, 176
 - Fibonacci-Folge, 201
 - Lösungsmengen homogener linearer Gleichungssysteme, 173
 - magische Quadrate, 174, 197
 - Rechenmauern, 202
 - stetiger Funktionen, 174
 - Summe linearer Unterräume, 176
- Unterraumkriterium, 172
- Untervektorraum, *siehe* Unterraum
- Vektor, 110
 - Addition
 - Assoziativität, 110, 168
 - Kommutativität, 110, 168
 - Beschleunigungsvektor, 146
 - Betrag, 124, 222
 - Distributivgesetz, 110, 168
 - Farbvektor, 101
 - Gegenvektor, 110, 168
 - Geschwindigkeitsvektor, 86, 146
 - Kollinearität, 116
 - Komplanarität, 116
 - Kraftvektor, 87
 - Linearkombination, 113
 - n -Tupel, 103–105
 - Nullvektor, 110, 168
 - orthogonal, 126, 222
 - Pfeilkategorie, 89–99
 - Populationsvektor, 242
 - Preisvektor, 122
 - RGB-Vektor, 101
 - skalare Multiplikation
 - Assoziativität, 110, 168
 - Stückliste, 100
 - Verschiebungsvektor, 88, 259, 280, 293
- Vektorprodukt, 133–136
- Vektorraum, 168–177
 - Basis eines Vektorraumes, 187–197
 - der $m \times n$ -Matrizen, 171
 - der Folgen reeller Zahlen, 169, 198
 - der n -Tupel reeller Zahlen, 169
 - der Pfeilkategorie, 169
 - der Polynome höchstens n -ten Grades, 170, 174, 220, 267
 - Basis, 189
 - Erzeugendensystem, 181
 - der Verschiebungen, 169
 - Dimension, 197–201
 - Euklidischer, 219–226
 - reellwertiger Funktionen, 170, 181
 - Unterraum eines Vektorraumes, 172
- Vektorraumhomomorphismus, 264
- Verbindungsraum, 212
- Verkürzungssatz, 192
- Verkettung, *siehe* Nacheinanderausführung von Abbildungen
- Verschiebung, 88, 259, 280, 281
- Volumen, 136
- Winkel
 - Eulersche, 291, 293
 - zwischen Ebenen, 160
 - zwischen Geraden im Raum, 159
 - zwischen Geraden in der Ebene, 53
 - zwischen Geraden und Ebenen, 160
 - zwischen Vektoren, 128, 129, 223
- Winkeltreue, 288
- Zahlenfolge, 169, 198
- Zeilenindex, 228
- Zeilenrang, 229–234
- Zeilenstufenform, 23, 229
- Zeilenumformungen, 232
- Zeilenvektor, 228
- Zuordnung, 258