

- Lücken-Skript -

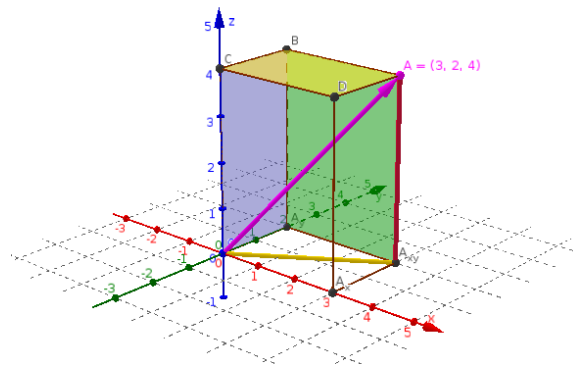
Mathematik 1

TWIE24

Studiengang Wirtschaftsingenieurwesen

Prof. Dr.-Ing. Stephan Sauter

Q1 2025



Inhaltsverzeichnis

Vorwort	1
1 Vektorrechnung	2
1.1 Einführung	2
1.2 Definitionen	2
1.3 Darstellung von Vektoren im Koordinatensystem	3
1.3.1 Basisvektoren	4
1.3.2 Betrag	4
1.3.3 Gleichheit	4
1.3.4 Addition	4
1.3.5 Multiplikation mit einem Skalar	4
1.3.6 Verbindungsvektoren	5
1.4 Lineare Abhängigkeit	5
1.4.1 Linearkombination:	5
1.4.2 Lineare Abhängigkeit:	6
1.4.3 Lineare Unabhängigkeit:	6
1.5 Skalarprodukt	7
1.5.1 Motivation	7
1.5.2 Definition	7
1.5.3 Eigenschaften	7
1.5.4 Berechnung	8
1.6 Vektorprodukt	8
1.6.1 Motivation	9
1.6.2 Definition	9
1.6.3 Regeln	10
1.6.4 Berechnung	10
1.7 Spatprodukt	11
1.7.1 Definition	11
1.7.2 geometrische Herleitung	11
1.7.3 Regeln	13
1.8 Anwendungen in der Geometrie	14
1.8.1 Geradengleichungen	14
1.8.2 Ebenengleichungen	15
1.8.3 Abstandsbestimmungen	16
2 Matrizen und Determinanten	21
2.1 Matrizen	21
2.1.1 Einführung	21
2.1.2 Definitionen	21
2.1.3 Rechenoperationen	23

2.1.4	Rechenregeln	25
2.2	Determinanten	26
2.2.1	Definition	26
2.2.2	Eigenschaften	28
2.3	Spezielle Matrizen	30
2.3.1	Reguläre Matrizen	30
2.3.2	Inverse Matrizen	30
2.3.3	Orthogonale Matrizen	32
2.4	Rang einer Matrix	32
2.4.1	Definition	32
2.4.2	Regeln	32
2.4.3	Rangbestimmung	33
3	Lineare Gleichungssysteme	34
3.1	Einführung	34
3.2	Definitionen	34
3.3	Lösungsverhalten	35
3.3.1	Lösbarkeit	35
3.3.2	Lösungsmenge	36
3.3.3	Lösungsberechnung - Cramersche Regel	36
3.3.4	Lösungsberechnung - Gaußscher Algorithmus	37
3.4	Spezialfall: quadratische lineare Gleichungssysteme	39
3.4.1	homogenes quadratisches lineares GLS	39
3.4.2	inhomogenes quadratisches lineares GLS	39
3.5	Rundungsfehler	40
3.6	Geschwindigkeit der Verfahren	41
4	Lineare Abbildungen	42
4.1	Definitionen	42
4.1.1	n-dimensionaler reeller Koordinatenraum	42
4.1.2	Lineare Abbildung	42
4.2	Beispiele	43
4.2.1	$\mathbb{R} \rightarrow \mathbb{R}$	43
4.2.2	$\mathbb{R}^2 \rightarrow \mathbb{R}$	43
4.2.3	$\mathbb{R}^2 \rightarrow \mathbb{R}^2$	43
4.2.4	$\mathbb{R}^3 \rightarrow \mathbb{R}^3$	44
4.3	Eigenwerte, Eigenvektoren quadratischer Matrizen	46
4.3.1	Definitionen	47
4.3.2	Berechnung	47
4.3.3	EW und EV einer Dreiecksmatrix	48
4.3.4	EW und EV einer symmetrischen Matrix	48
5	Folgen	49
5.1	Definition	49
5.1.1	Beispiele	49
5.1.2	Definition monotone und beschränkte Folgen	50
5.2	Konvergenz	50
5.2.1	Definitionen	50
5.2.2	Rechenregeln für konvergente Folgen	51
5.2.3	Folgen der Form $p(n)/q(n)$	51
5.2.4	Eulersche Zahl (ohne Beweis)	52
5.2.5	Konvergenzaussagen für monotone Folgen (ohne Beweis)	52

6	Funktionen	53
6.1	Darstellungsformen von Funktionen	53
6.2	Eigenschaften	54
6.2.1	Nullstellen	54
6.2.2	(un-)gerade Funktionen	54
6.2.3	monotone Funktionen	54
6.2.4	periodische Funktionen	54
6.2.5	umkehrbare Funktion	55
6.2.6	Grenzwert einer Funktion	55
6.2.7	Stetige Funktionen	56
6.3	Polynomfunktionen	57
6.3.1	Definition	57
6.3.2	Spezialfall: Polynom 1. Grades	57
6.3.3	Spezialfall: Polynom 2. Grades	59
6.3.4	Nullstellen	59
6.4	Gebrochen rationale Funktionen	60
6.4.1	Definition	60
6.4.2	Definitionsbereich, Nullstellen, Pole	61
6.4.3	Asymptoten	61
6.5	Potenzfunktionen	62
6.6	Trigonometrische Funktionen	63
6.6.1	Definition	63
6.6.2	Zusammenhänge	64
6.6.3	Allgemeine Sinus- und Kosinusfunktion	64
6.6.4	Darstellung der Sinusschwingung im Zeigerdiagramm	65
6.7	Arkusfunktionen	67
6.7.1	Arkussinus & Arkuscosinus	67
6.7.2	Arkustangens & Arkuscotangens	68
6.7.3	Trigonometrische Gleichungen	68
6.8	Exponentialfunktionen	68
6.9	Logarithmusfunktionen	69
6.10	Hyperbelfunktionen	70
6.11	Areafunktionen	72
7	Komplexe Zahlen	73
7.1	Einführung	73
7.2	Definitionen	74
7.2.1	Zahlen	74
7.2.2	Komplexe Zahlen	75
7.3	Grundrechenoperationen	78
7.3.1	Addition	78
7.3.2	Multiplikation mit einem Skalar	78
7.3.3	Multiplikation	78
7.3.4	Division	78
7.3.5	Graphische Darstellung der Grundrechenoperationen	79
7.3.6	Potenzieren und Wurzelziehen	79
7.4	Funktionen komplexer Zahlen	80
7.4.1	Potenzfunktionen	81
7.4.2	Sinus-, Cosinus- und e -Funktion	81
7.4.3	Komplexe e -Funktion	81
7.4.4	Anwendungsbeispiel aus der Elektrotechnik	82

8	Anhang	83
8.1	Einheitskreis	83

Vorwort

Das vorliegende Skript soll vorlesungsbegleitend dem Hörer das Abzeichnen bzw. Abschreiben der Inhalte ersparen. Falls eine Vorlesungsstunde versäumt wurde, kann der Hörer anhand des Skriptes ansehen, welcher Stoff z.B. mit einem Buch nachgeholt werden sollte.

Bei allen Betrachtungen steht eine anschauliche Darstellung im Vordergrund. Es soll versucht werden, dem Leser Hinweise zu geben, die ihm bei der Lösung der anstehenden Problemstellungen nützlich sind.

Insbesondere wird darauf hingewiesen, dass für die Prüfung das selbständige Lösen der Übungsaufgaben nicht nur empfohlen, sondern vorausgesetzt wird!

- Bronstein u.a.: Taschenbuch der Mathematik Edition Harri Deutsch
- Wilhelm Leupold u.a. : Mathematik - ein Studienbuch für Ingenieure. Band 1 Carl Hanser Verlag
- L. Papula: Mathematik für Ingenieure und Naturwissenschaftler. Band 1 Verlag Vieweg
- Harro Heuser: Lehrbuch der Analysis. Teil 1, Springer Verlag
- B. Neumayer, S. Kaup: Mathematik für Ingenieure 1, Shaker Verlag Aachen
- www.wolframalpha.com

Musterlösungen für die Übungsaufgaben, Formelsammlungen und Skript:

- www.Freiwilligschlauwerden.de



KAPITEL 1

Vektorrechnung

1.1 Einführung

Physikalische Größen, wie Arbeit und Temperatur können durch eine Zahl beschrieben werden. Kraft, Geschwindigkeit und das elektrische Feld in einem Punkt sind jedoch durch ihren Betrag und ihre Richtung definiert.

Um hierfür ein mathematisches Werkzeug zu besitzen, wurden Vektoren eingeführt. Vektoren sind Größen, die durch Betrag **und** Richtung festgelegt sind. Mit Hilfe der Vektorrechnung können z.B. folgende Größen mathematisch beschrieben werden.

- Addition von Kräften
- Vervielfachung einer Geschwindigkeit
- Zerlegung eines Feldes

1.2 Definitionen

- **Darstellung**

Vektoren werden als Pfeile dargestellt und repräsentiert durch: $\vec{a}, \vec{b}, \vec{c}$

Bei Angabe des Anfangspunktes A und des Endpunktes B (Pfeilspitze), werden Vektoren als $\overrightarrow{AB}$ dargestellt.

- **Betrag eines Vektors**

Der Betrag des Vektors entspricht der Länge des Pfeils. Für den Betrag des Vektors $\vec{a}$ schreibt man $|\vec{a}|$

- **Nullvektor**

Ein Vektor mit Betrag 0, heißt Nullvektor: $\vec{0}$

- **Einheitsvektor**

Ein Vektor mit Betrag 1 heißt Einheitsvektor.

Zu jedem Vektor $\vec{a} \neq \vec{0}$ gibt es einen gleichgerichteten Einheitsvektor $\vec{e}_a = \frac{\vec{a}}{|\vec{a}|}$

- **Gleichheit**

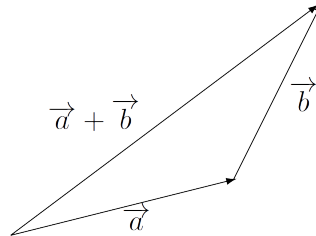
Zwei Vektoren heißen gleich, wenn sie in Betrag und Richtung übereinstimmen: $\vec{a} = \vec{b}$

- **Kolineariät**

Zwei Vektoren heißen kollinear, wenn sie parallel oder antiparallel sind.

• Vektoraddition

Zwei Vektoren $\vec{a}$ und $\vec{b}$ werden addiert, indem man den Vektor $\vec{b}$ an den Endpunkt des Vektors $\vec{a}$ ansetzt. Der dann vom Anfang von $\vec{a}$ bis zum Ende von $\vec{b}$ führende Vektor heißt der Summenvektor $\vec{a} + \vec{b}$

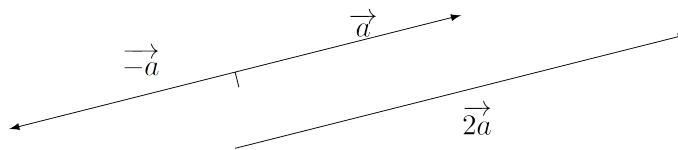


Für die Vektoraddition gelten:

Kommutativgesetz:	$\vec{a} + \vec{b} = \vec{b} + \vec{a}$
Assoziativgesetz:	$\vec{a} + (\vec{b} + \vec{c}) = (\vec{a} + \vec{b}) + \vec{c}$
Dreiecksungleichung:	$ \vec{a} + \vec{b} \leq \vec{a} + \vec{b} $

• Multiplikation mit einem Skalar

Ist $\lambda \in \mathbb{R}$, dann ist $\lambda\vec{a}$ ein Vektor, der parallel zu $\vec{a}$ ist und dessen Betrag das λ -fache des Betrag von $\vec{a}$ ist.



Für die Multiplikation eines Vektors mit einem Skalar gelten:

Kommutativgesetz:	$\lambda\vec{a} = \vec{a}\lambda$
Assoziativgesetz:	$\lambda(\mu\vec{a}) = (\lambda\mu)\vec{a}$
Distributivgesetze:	$\lambda(\vec{a} + \vec{b}) = \lambda\vec{a} + \lambda\vec{b}$
	$(\lambda + \mu)\vec{a} = \lambda\vec{a} + \mu\vec{a}$

1.3 Darstellung von Vektoren im Koordinatensystem

Zur rechnerischen Behandlung von Vektoren wird ein (zweiachsiges) dreiachsiges kartesisches Koordinatensystem zugrunde gelegt.

Die Einheitsvektoren, deren Richtungen mit der positiven Richtung der x -, y - bzw. z -Achse übereinstimmen, werden mit $\vec{e}_x$, $\vec{e}_y$, bzw. $\vec{e}_z$ bezeichnet. Sie werden **Basisvektoren** genannt.

Jeder Vektor $\vec{a} \neq \vec{0}$ kann eindeutig als Summe von Vielfachen der Basisvektoren dargestellt werden:

$$\vec{a} = a_x \vec{e}_x + a_y \vec{e}_y + a_z \vec{e}_z$$

Da ein Vektor durch ein Zahlentripel a_x, a_y, a_z gegeben ist, kann auch umgekehrt ein Zahlentripel als Vektor gedeutet werden. Es wird daher als Schreibweise für Vektoren meist eine einspaltige Matrix verwendet:

$$\vec{a} = \begin{pmatrix} a_x \\ a_y \\ a_z \end{pmatrix}$$

1.3.1 Basisvektoren

Für die Basisvektoren gilt:

$$\vec{e}_x = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \vec{e}_y = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \vec{e}_z = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$$

1.3.2 Betrag

Für den Betrag eines Vektors folgt nach dem Satz des Pythagoras:

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

1.3.3 Gleichheit

Aus der Definition der Gleichheit von Vektoren folgt für zwei Vektoren

$$\vec{a} = \vec{b} \iff a_x = b_x, \quad a_y = b_y, \quad a_z = b_z$$

1.3.4 Addition

Die Summe zweier Vektoren $\vec{a}$ und $\vec{b}$ berechnet sich als

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

1.3.5 Multiplikation mit einem Skalar

Für die Multiplikation eines Vektors mit einem Skalar erhält man:

$$\lambda \vec{a} = \begin{pmatrix} \lambda a_x \\ \lambda a_y \\ \lambda a_z \end{pmatrix}$$

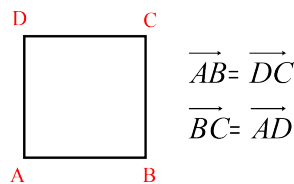
1.3.6 Verbindungsvektoren

Ein Vektor welcher zwei Punkte verbindet:

$$\vec{v} = \overrightarrow{AB} = \vec{b} - \vec{a} = \begin{pmatrix} b_x - a_x \\ b_y - a_y \\ b_z - a_z \end{pmatrix}$$

Vektoren sind im Gegensatz zu den Punkten im Raum nicht absolut. D.h. eine vektorielle Größe ist nur in Bezug auf Länge und Richtung, jedoch **nicht** bezüglich seiner Position im reellen Raum definiert.

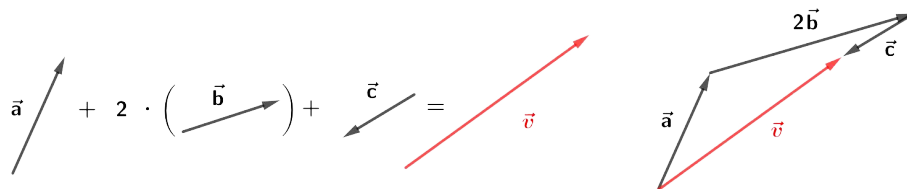
Sei $ABCD$ ein Quadrat mit den vier Punkten A, B, C und D , so gilt:



1.4 Lineare Abhängigkeit

1.4.1 Linearkombination:

Eine Linearkombination von endlich vielen Vektoren ist die Summe von beliebigen Vielfachen dieser Vektoren. Die Vielfachen heißen Koeffizienten.



Z.B. ist $\begin{pmatrix} 3 \\ 2 \end{pmatrix}$ eine Linearkombination der Vektoren $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\begin{pmatrix} 1 \\ 0 \end{pmatrix}$, denn:

$$\begin{pmatrix} 3 \\ 2 \end{pmatrix} =$$

Jeder 2-dimensionale, bzw. 3-dimensionale Vektor ist durch eine Linearkombination der Einheitsvektoren $\vec{e}_x, \vec{e}_y$ bzw. $\vec{e}_x, \vec{e}_y, \vec{e}_z$ darstellbar.

Beispiel:

Der Vektor $\begin{pmatrix} 3 \\ 4 \\ 5 \end{pmatrix}$ soll als Linearkombination der Vektoren $\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$, $\begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$ und $\begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$ geschrieben werden:

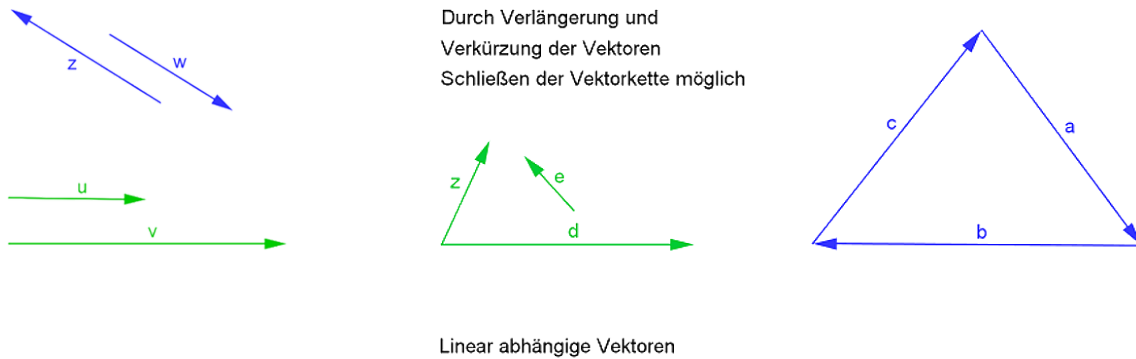
$$\begin{pmatrix} 3 \\ 4 \\ 5 \end{pmatrix} =$$

1.4.2 Lineare Abhängigkeit:

Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ heißen **linear abhängig**, wenn es Skalare $\alpha_1, \alpha_2, \dots, \alpha_n$ gibt, die nicht alle Null sind, so dass die Bedingung

$$\alpha_1 \vec{a}_1 + \alpha_2 \vec{a}_2 + \dots + \alpha_n \vec{a}_n = 0$$

erfüllt ist, d.h. einer der Vektoren ist als Linearkombination der anderen darstellbar.

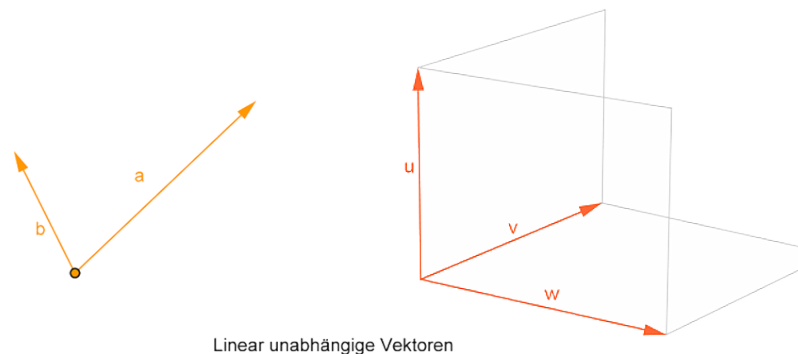


1.4.3 Lineare Unabhängigkeit:

Vektoren $\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n$ heißen **linear unabhängig**, wenn die Bedingung

$$\alpha_1 \vec{a}_1 + \alpha_2 \vec{a}_2 + \dots + \alpha_n \vec{a}_n = 0$$

nur für $\alpha_1 = \alpha_2 = \dots = \alpha_n = 0$ erfüllt ist, d.h. keiner der Vektoren ist als Linearkombination der anderen darstellbar.



Beispiele

Sind folgende Vektoren linear unabhängig?

1. $\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$

2. $\begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \begin{pmatrix} 3 \\ 1 \\ 2 \end{pmatrix}$

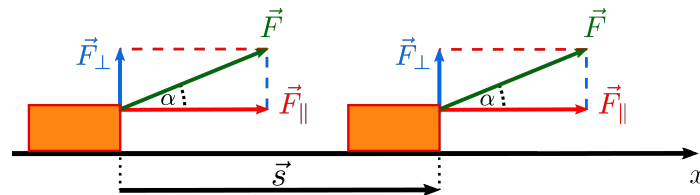
3. $\begin{pmatrix} 3 \\ 2 \\ 1 \end{pmatrix}, \begin{pmatrix} 13 \\ 2 \\ 7 \end{pmatrix}, \begin{pmatrix} 12 \\ 8 \\ 4 \end{pmatrix}$

1.5 Skalarprodukt

1.5.1 Motivation

Berechnung der Arbeit in der Physik:

Ein Körper wird entlang der x -Achse bewegt, vom Ursprung bis nach $s = 100$ m. Hierfür wird eine Kraft $F = 5$ N eingesetzt, die im Winkel von $\alpha = 30^\circ$ zur x -Achse am Körper zieht. Welche Arbeit wird hier verrichtet?



Die Arbeit ergibt sich zu:

$$W = |\vec{F}| \cdot \cos(\alpha) \cdot |\vec{s}|$$

Das Skalarprodukt zwischen zwei Vektoren ist nun genau so definiert, dass man hier kurz

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

schreiben kann.

Also ergibt sich für die Arbeit im Beispiel oben:

$$W = 5 \text{ N} \cdot \cos(30^\circ) \cdot 100 \text{ m} \approx 5 \text{ N} \cdot 0.866 \cdot 100 \text{ m} = 433 \text{ Nm}$$

1.5.2 Definition

Das Skalarprodukt zweier Vektoren ist definiert durch:

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

mit $0 \leq \varphi \leq 180^\circ$.

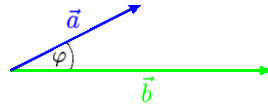
Das Ergebnis des Skalarprodukts ist, wie der Name sagt, ein Skalar, d.h. eine reelle Zahl.

1.5.3 Eigenschaften

1. $\vec{a} \cdot \vec{b} = \vec{b} \cdot \vec{a}$ Kommutativgesetz
2. $\vec{a} \cdot (\vec{b} + \vec{c}) = \vec{a} \cdot \vec{b} + \vec{a} \cdot \vec{c}$ Distributivgesetz
3. $\vec{a} \cdot \vec{a} = |\vec{a}|^2$
4. $\vec{a} \cdot \vec{b} = 0$ $\vec{a}$ ist senkrecht zu $\vec{b}$ (alternativ: $\vec{a} \perp \vec{b}$)
5. $\vec{e}_x \cdot \vec{e}_x = 1$, $\vec{e}_y \cdot \vec{e}_y = 1$, $\vec{e}_z \cdot \vec{e}_z = 1$
6. $\vec{e}_x \cdot \vec{e}_y = 0$, $\vec{e}_y \cdot \vec{e}_z = 0$, $\vec{e}_z \cdot \vec{e}_x = 0$
7. $\vec{a} \cdot \vec{e}_x = a_x$, $\vec{a} \cdot \vec{e}_y = a_y$, $\vec{a} \cdot \vec{e}_z = a_z$

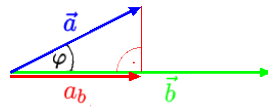
8. Der Winkel φ zwischen zwei Vektoren $\vec{a}$ und $\vec{b}$ berechnet sich aus

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



9. Die Projektionslänge a_b eines Vektors $\vec{a}$ auf einen Vektor $\vec{b}$ hat folgende Größe:

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



1.5.4 Berechnung

Das Skalarprodukt zweier n -dimensionaler Vektoren berechnet sich zu

$$\vec{a} \circ \vec{b} = a_1 b_1 + a_2 b_2 + \dots + a_n b_n$$

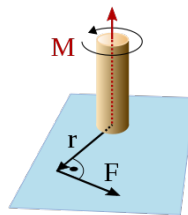
denn:

$$\begin{aligned} \vec{a} \circ \vec{b} &= (a_1 \vec{e}_1 + a_2 \vec{e}_2 + \dots + a_n \vec{e}_n) \circ (b_1 \vec{e}_1 + b_2 \vec{e}_2 + \dots + b_n \vec{e}_n) \\ &= a_1 \vec{e}_1 \circ b_1 \vec{e}_1 + a_1 \vec{e}_1 \circ b_2 \vec{e}_2 + \dots + a_1 \vec{e}_1 \circ b_n \vec{e}_n + \\ &\quad a_2 \vec{e}_2 \circ b_1 \vec{e}_1 + a_2 \vec{e}_2 \circ b_2 \vec{e}_2 + \dots + a_2 \vec{e}_2 \circ b_n \vec{e}_n + \\ &\quad \dots \\ &\quad a_n \vec{e}_n \circ b_1 \vec{e}_1 + a_n \vec{e}_n \circ b_2 \vec{e}_2 + \dots + a_n \vec{e}_n \circ b_n \vec{e}_n \\ &= a_1 b_1 + a_2 b_2 + \dots + a_n b_n \end{aligned}$$

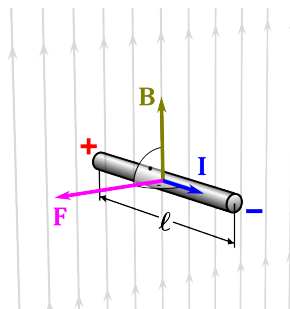
1.6 Vektorprodukt

Das Vektorprodukt $\vec{a} \times \vec{b}$ zweier Vektoren $\vec{a}$ und $\vec{b}$ ist nur im 3-dimensionalen Raum definiert. Das Ergebnis ist wieder ein Vektor, daher der Name Vektorprodukt. Gelegentlich wird das Vektorprodukt auch als Kreuzprodukt bezeichnet, da es durch ein Kreuz symbolisiert wird. Folgende physikalische Größen sind als Vektorprodukt darstellbar:

- Drehmoment einer an einem starren Körper angreifenden Kraft
- Drehimpuls eines rotierenden Körpers



- Kraft auf einen stromdurchflossenen Leiter im Magnetfeld



1.6.1 Motivation

Berechnung des Drehmoments in der Physik: An einem an einer Achse befestigter Hebel der Länge $s = 100$ m greift eine Kraft $F = 5$ N an, die im Winkel von $\alpha = 30^\circ$ zum Hebel nach unten zieht. Welches Drehmoment tritt auf?



Das Drehmoment ist ein Vektor, der in das Blatt hinein zeigt und eine Länge hat, die dem Flächeninhalt des Parallelogramms zwischen $\vec{s}$ und $\vec{F}$ entspricht.

$$|\vec{M}| = |\vec{F}| \cdot \sin(\alpha) \cdot |\vec{s}|$$

$$\Rightarrow M = 5 \text{ N} \cdot \sin(30^\circ) \cdot 100 \text{ m} = 5 \text{ N} \cdot 0.5 \cdot 100 \text{ m} = 250 \text{ Nm}$$

Das Vektorprodukt zwischen zwei Vektoren ist nun genau so definiert, dass man hier kurz

$$\vec{M} = \vec{s} \times \vec{F} \quad \text{schreiben kann.}$$

1.6.2 Definition

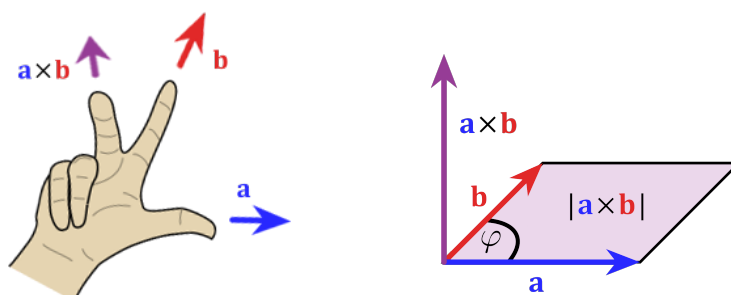
Das Vektorprodukt $\vec{c} = \vec{a} \times \vec{b}$ ist ein Vektor $\vec{c}$ mit folgenden Eigenschaften:

- $|\vec{c}|$ ist gleich dem Flächeninhalt des von $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms:
 $|\vec{c}| = |\vec{a}||\vec{b}|\sin\varphi \quad (0^\circ \leq \varphi \leq 180^\circ)$
- $\vec{c}$ steht sowohl auf $\vec{a}$ als auch auf $\vec{b}$ senkrecht, d.h.
 $\vec{c} \cdot \vec{a} = \vec{c} \cdot \vec{b} = 0$
- Die Richtung von $\vec{c}$ ist so festgelegt, dass $\vec{a}, \vec{b}$ und $\vec{c}$ ein Rechtssystem bilden.

Rechtssystem:

Als Rechtssystem wird ein System aus drei Vektoren im 3-dimensionalen Raum bezeichnet, wenn diese der Rechten-Hand-Regel entsprechen, d.h.

- Daumen in Richtung des ersten Vektors
- Zeigefinger in Richtung des zweiten Vektors
- Mittelfinger (rechtwinklig zum Daumen und Zeigefinger abgespreizt) zeigt bei einem Rechtssystem in Richtung des dritten Vektors



1.6.3 Regeln

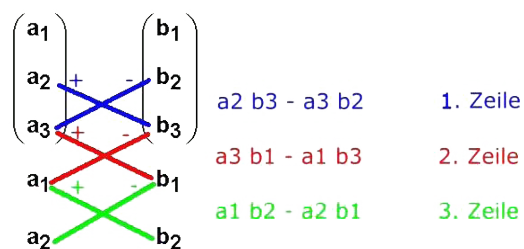
- $\vec{a} \times \vec{0} = \vec{0}$
- $\vec{a} \times \vec{a} = \vec{0}$
- $\vec{a} \times \vec{b} = \vec{0}$
Kolinearität ($\vec{a}$, $\vec{b}$ parallel oder antiparallel, $\vec{a} \parallel \vec{b}$)
- $\vec{a} \times (\vec{b} + \vec{c}) = \vec{a} \times \vec{b} + \vec{a} \times \vec{c}$
 $(\vec{a} + \vec{b}) \times \vec{c} = \vec{a} \times \vec{c} + \vec{b} \times \vec{c}$
Distributivgesetze
- $(\lambda \vec{a}) \times \vec{b} = \lambda(\vec{a} \times \vec{b}) = \vec{a} \times (\lambda \vec{b})$ für alle $\lambda \in \mathbb{R}$
Assoziativgesetz
- $\vec{a} \times \vec{b} = -(\vec{b} \times \vec{a})$
antikommutativ
- $\vec{e}_x \times \vec{e}_y = \vec{e}_z$
- Lagrange-Identität:
 $(\vec{a} \times \vec{b}) \circ (\vec{c} \times \vec{d}) = (\vec{a} \circ \vec{c}) \circ (\vec{b} \circ \vec{d}) - (\vec{b} \circ \vec{c}) \circ (\vec{a} \circ \vec{d})$

1.6.4 Berechnung

Das Vektorprodukt zweier 3-dimensionalen Vektoren berechnet sich zu

$$\vec{a} \times \vec{b} = \begin{pmatrix} a_2 b_3 - a_3 b_2 \\ a_3 b_1 - a_1 b_3 \\ a_1 b_2 - a_2 b_1 \end{pmatrix}$$

Merkregel:



Beispiele:

$$1. \begin{pmatrix} 2 \\ 4 \\ 1 \end{pmatrix} \times \begin{pmatrix} 3 \\ -2 \\ 1 \end{pmatrix} =$$

$$2. \begin{pmatrix} -2 \\ -1 \\ 2 \end{pmatrix} \times \begin{pmatrix} 2 \\ 2 \\ 0 \end{pmatrix} =$$

Beweis:

$$\begin{aligned}
 \vec{a} \times \vec{b} &= (a_1\vec{e}_1 + a_2\vec{e}_2 + a_3\vec{e}_3) \times (b_1\vec{e}_1 + b_2\vec{e}_2 + b_3\vec{e}_3) \\
 &= a_1\vec{e}_1 \times b_1\vec{e}_1 + a_1\vec{e}_1 \times b_2\vec{e}_2 + a_1\vec{e}_1 \times b_3\vec{e}_3 + \\
 &\quad a_2\vec{e}_2 \times b_1\vec{e}_1 + a_2\vec{e}_2 \times b_2\vec{e}_2 + a_2\vec{e}_2 \times b_3\vec{e}_3 + \\
 &\quad a_3\vec{e}_3 \times b_1\vec{e}_1 + a_3\vec{e}_3 \times b_2\vec{e}_2 + a_3\vec{e}_3 \times b_3\vec{e}_3 \\
 &= 0 + a_1b_2\vec{e}_3 - a_1b_3\vec{e}_2 \\
 &\quad - a_2b_1\vec{e}_3 + 0 + a_2b_3\vec{e}_1 \\
 &\quad + a_3b_1\vec{e}_2 - a_3b_2\vec{e}_1 + 0 \\
 &= (a_2b_3 - a_3b_2)\vec{e}_1 + (a_3b_1 - a_1b_3)\vec{e}_2 + (a_1b_2 - a_2b_1)\vec{e}_3 \\
 &= \vec{a} \times \vec{b} = \begin{pmatrix} a_2b_3 - a_3b_2 \\ a_3b_1 - a_1b_3 \\ a_1b_2 - a_2b_1 \end{pmatrix}
 \end{aligned}$$

NR. $\vec{e}_n \times \vec{e}_n = 0$
 $\vec{e}_3 = \vec{e}_1 \times \vec{e}_2$
 $-\vec{e}_2 = \vec{e}_1 \times \vec{e}_3$
 $\vec{e}_1 = \vec{e}_2 \times \vec{e}_3$
 $-\vec{e}_3 = \vec{e}_2 \times \vec{e}_1$
 $\vec{e}_2 = \vec{e}_3 \times \vec{e}_1$
 $-\vec{e}_1 = \vec{e}_3 \times \vec{e}_2$

1.7 Spatprodukt

Vektor- und Skalarprodukt sind über das Spatprodukt miteinander verknüpft. Das Spatprodukt ist das Ergebnis aus dem Vektorprodukt zweier Vektoren und dem Skalarprodukt mit einem dritten Vektor. Es beschreibt die Größe des **orientierten** Volumens des Parallelepipeds (Spat), das durch die drei Vektoren aufgespannt wird.

Unter orientiertem Volumen versteht man dabei das Volumen multipliziert mit dem Faktor +1, falls die Vektoren ein Rechtssystem bilden, und multipliziert mit -1, falls sie ein Linkssystem bilden.

1.7.1 Definition

Das Spatprodukt dreier 3-dimensionaler Vektoren ist ein Skalar mit

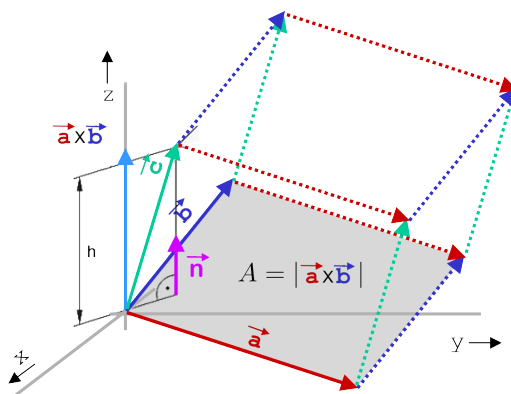
$$\vec{a}, \vec{b}, \vec{c} \mapsto (\vec{a} \times \vec{b}) \circ \vec{c} \in \mathbb{R}$$

1.7.2 geometrische Herleitung

Das Volumen eines Spats errechnet sich aus dem Produkt seiner Grundfläche und seiner Höhe.

$$V = Ah$$

Bekanntlich ist das Vektorprodukt $\vec{a} \times \vec{b}$ genau der Normalenvektor auf der durch $\vec{a}$ und $\vec{b}$ aufgespannten Grundfläche und dessen Betrag gleich dem Flächeninhalt dieser Fläche, also $A = |\vec{a} \times \vec{b}|$.



Die Höhe des Spats ist die Projektion des Vektors $\vec{c}$ auf die Richtung des Normalenvektors $\vec{n}$.
Es folgt:

$$V = Ah = |\vec{a} \times \vec{b}| \underbrace{\left(\frac{\vec{a} \times \vec{b}}{|\vec{a} \times \vec{b}|} \circ \vec{c} \right)}_{\vec{n}} = (\vec{a} \times \vec{b}) \circ \vec{c} = [\vec{a} \ \vec{b} \ \vec{c}]$$

Das Spatprodukt gibt das orientierte Volumen des Spats an.

Einfache Berechnung:

Spatprodukt = Wert der Determinante

$$(\vec{a} \times \vec{b}) \circ \vec{c} = \begin{vmatrix} a_1 & b_1 & c_1 \\ a_2 & b_2 & c_2 \\ a_3 & b_3 & c_3 \end{vmatrix} = a_1 \cdot b_2 \cdot c_3 + b_1 \cdot c_2 \cdot a_3 + c_1 \cdot a_2 \cdot b_3 - c_1 \cdot b_2 \cdot a_3 - a_1 \cdot c_2 \cdot b_3 - b_1 \cdot a_2 \cdot c_3$$

→ siehe Determinanten im nächsten Kapitel.

Beispiel:

$$\vec{a} = \begin{pmatrix} 3 \\ -3 \\ 4 \end{pmatrix} \quad \vec{b} = \begin{pmatrix} -4 \\ -7 \\ 2 \end{pmatrix} \quad \vec{c} = \begin{pmatrix} 7 \\ 2 \\ 2 \end{pmatrix}$$

Spatprodukt: $(\vec{a} \times \vec{b}) \circ \vec{c}$

Determinante D

Es gilt:

$$(\vec{a} \times \vec{b}) \circ \vec{c} > 0 \iff \vec{a}, \vec{b}, \vec{c} \text{ bilden ein Rechtssystem}$$

$$(\vec{a} \times \vec{b}) \circ \vec{c} < 0 \iff \vec{a}, \vec{b}, \vec{c} \text{ bilden ein Linkssystem}$$

$$(\vec{a} \times \vec{b}) \circ \vec{c} = 0 \iff \vec{a}, \vec{b}, \vec{c} \text{ sind linear abhängig} \Rightarrow \underline{\text{einfacher Test!}}$$

Beispiel:

Bilden die folgenden 3 Vektoren ein Rechtssystem und sind sie linear abhängig?

$$\vec{a} = \begin{pmatrix} 3 \\ -3 \\ 4 \end{pmatrix} \quad \vec{b} = \begin{pmatrix} -4 \\ -7 \\ 2 \end{pmatrix} \quad \vec{c} = \begin{pmatrix} 7 \\ 2 \\ 2 \end{pmatrix}$$

1.7.3 Regeln

- $(\vec{a} \times \vec{b}) \circ \vec{c} = (\vec{b} \times \vec{c}) \circ \vec{a} = (\vec{c} \times \vec{a}) \circ \vec{b} \Rightarrow \underline{\text{zyklische Vertauschung der Vektoren}}$
- $(\vec{a} \times \vec{b}) \circ \vec{c} = -(\vec{b} \times \vec{a}) \circ \vec{c}$
- $(\vec{a} \times \vec{a}) \circ \vec{b} = 0$
- Assoziativgesetz: $(\lambda \vec{a} \times \vec{b}) \circ \vec{c} = \lambda (\vec{a} \times \vec{b}) \circ \vec{c}$
- Distributivgesetz: $(\vec{a} \times \vec{b}) \circ (\vec{c} + \vec{d}) = (\vec{a} \times \vec{b}) \circ \vec{c} + (\vec{a} \times \vec{b}) \circ \vec{d}$

1.8 Anwendungen in der Geometrie

Mit Hilfe der Vektorrechnung lassen sich im 3-dimensionalen Raum einfach Geraden und Ebenen sowie ihre relativen Lagen vektoriell beschreiben.

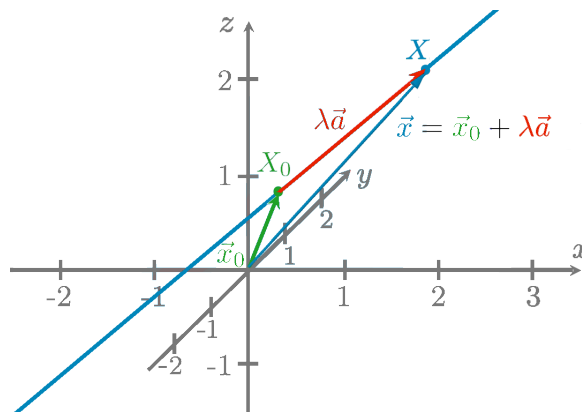
1.8.1 Geradengleichungen

Punkt-Richtungs-Form

Eine Gerade im $\mathbb{R}^3$ durch einen Punkt X_0 und mit einem Richtungsvektor $\vec{a}$ ist gegeben durch:

$$\vec{x} = \vec{x}_0 + \lambda \vec{a} \quad \text{mit } \lambda \in \mathbb{R}$$

Hierbei ist $\vec{x}$ der Ortsvektor zu den Punkten X der Geraden.



Beispiel:

Bestimmen Sie die Gleichung der Geraden durch den Punkt $(3|-2|1)$ in Richtung $\begin{pmatrix} 5 \\ 2 \\ 3 \end{pmatrix}$.

2 Punkteform

Eine Gerade im $\mathbb{R}^3$ durch zwei Punkte X_1 und X_2 ist gegeben durch:

$$\vec{x} = \vec{x}_1 + \lambda(\vec{x}_2 - \vec{x}_1) \quad \text{mit } \lambda \in \mathbb{R}$$

Hierbei ist $\vec{x}$ der Ortsvektor zu den Punkten X der Geraden.

Beispiel:

Bestimmen Sie die Gleichung der Geraden durch die Punkte $(1|1|1)$ und $(2|0|4)$.

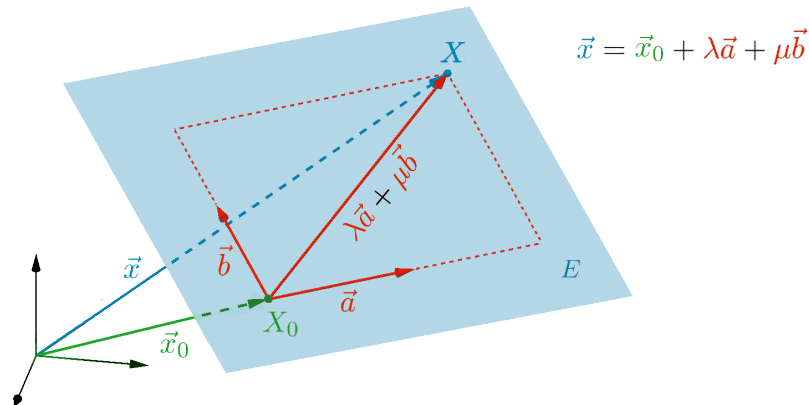
1.8.2 Ebenengleichungen

Parameterform

Eine Ebene im $\mathbb{R}^3$ durch einen Punkt X_0 und mit zwei linear unabhängigen Richtungsvektoren $\vec{a}$ und $\vec{b}$ ist gegeben durch:

$$\vec{x} = \vec{x}_0 + \lambda \vec{a} + \mu \vec{b} \quad \text{mit } \lambda, \mu \in \mathbb{R}$$

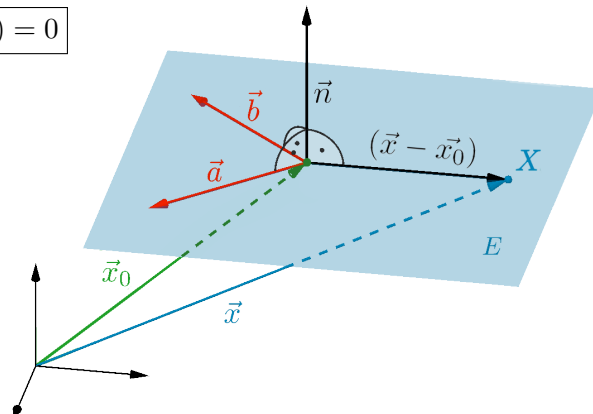
Hierbei ist $\vec{x}$ der Ortsvektor zu den Punkten X der Ebene.



Normalform

Eine Ebene im $\mathbb{R}^3$ durch einen Punkt X_0 und mit einem Vektor $\vec{n}$, senkrecht zur Ebene, ist gegeben durch:

$$\vec{n} \circ (\vec{x} - \vec{x}_0) = 0$$



Koordinatenform

Durch

$$ax + by + cz = d$$

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

Beispiel:

Bestimmen Sie einen Normalenvektor zur Ebene, gegeben durch $x_1 + 2x_2 - 3x_3 = 5$.

1.8.3 Abstandsbestimmungen

Abstand zweier Punkte:

Der Abstand zweier Punkte X_1 und X_2 im $\mathbb{R}^3$ oder im $\mathbb{R}^2$ ist:

$$d = |\vec{x}_2 - \vec{x}_1|$$

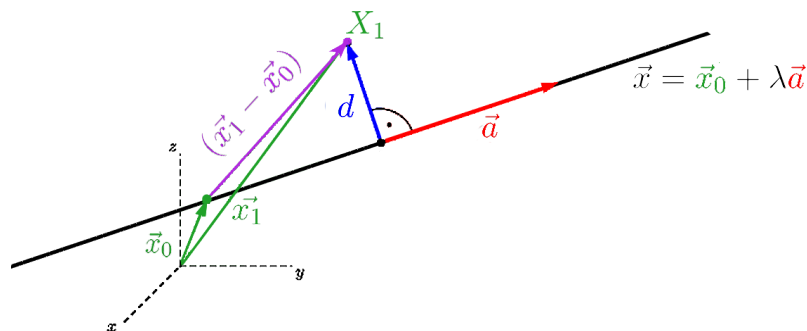
Beispiel:

Bestimmen Sie den Abstand zwischen den Punkten $(2|-5|3)$ und $(2|-1|0)$.

Abstand zwischen Punkt und Gerade:

Der Abstand zwischen einem Punkt X_1 und einer Gerade $\vec{x} = \vec{x}_0 + \lambda \vec{a}$ im $\mathbb{R}^3$ ist:

$$d = \frac{|(\vec{x}_1 - \vec{x}_0) \times \vec{a}|}{|\vec{a}|}$$



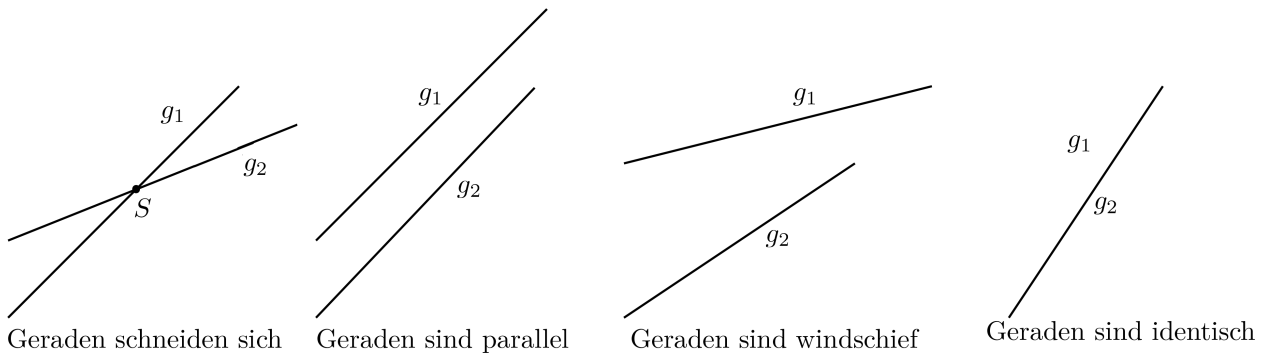
Beispiel:

Die Gleichung einer Geraden lautet

$\vec{x} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ 5 \\ 2 \end{pmatrix}$. Bestimmen Sie den Abstand des Punktes $(5|3|-2)$ von dieser Geraden.

Abstand zwischen zwei Geraden

Zwei Geraden im $\mathbb{R}^3$ können zueinander folgendermaßen liegen:



Abhängig von dieser Lage wird der Abstand bestimmt:

Für zwei nicht parallele Geraden, die sich **schneiden**, berechnet sich der Schnittpunkt

$$\vec{x}_1 + \lambda_1 \vec{a}_1 = \vec{x}_2 + \lambda_2 \vec{a}_2$$

und der Schnittwinkel:

$$\vec{a}_1 \circ \vec{a}_2 = |\vec{a}_1| |\vec{a}_2| \cos \varphi$$

Der Abstand zweier **paralleler** Geraden mit den Gleichungen

$$\vec{x} = \vec{x}_1 + \lambda_1 \vec{a} \text{ und } \vec{x} = \vec{x}_2 + \lambda_2 \vec{a}$$

$$d = \frac{|(\vec{x}_2 - \vec{x}_1) \times \vec{a}|}{|\vec{a}|}$$

Der Abstand zweier **windschiefer** Geraden mit den Gleichungen

$$\vec{x} = \vec{x}_1 + \lambda_1 \vec{a}_1 \text{ und } \vec{x} = \vec{x}_2 + \lambda_2 \vec{a}_2 \text{ ist:}$$

$$d = \frac{|(\vec{x}_2 - \vec{x}_1) \circ (\vec{a}_1 \times \vec{a}_2)|}{|(\vec{a}_1 \times \vec{a}_2)|}$$

Beispiel:

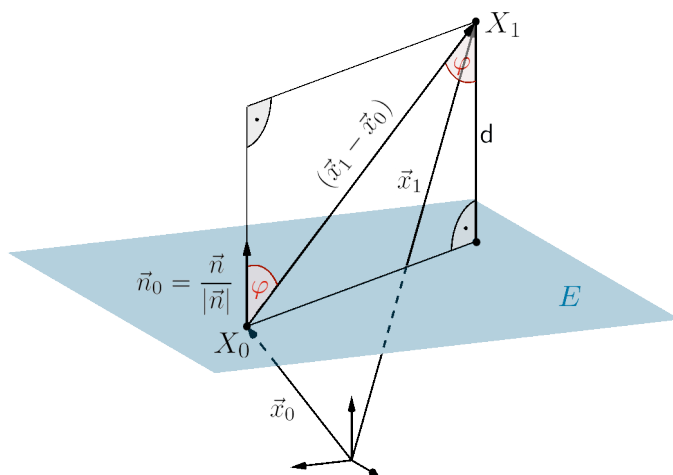
Gegeben sei $g_1 := \vec{x} = (-1|3|-1) + \lambda_1(-2|3|1)$ und $g_2 : \vec{x} = (5|-2|-3) + \lambda_2(-8|4|2)$.
Untersuchen Sie, ob g_1 und g_2 gemeinsame Punkte haben und bestimmen Sie ggf. den Schnittpunkt.

Abstand zwischen Punkt und Ebene

Der Abstand zwischen einem Punkt X_1 und einer Ebene in Normalform

$\vec{n} \circ (\vec{x} - \vec{x}_0) = 0$ ist:

$$d = \frac{|\vec{n} \circ (\vec{x}_1 - \vec{x}_0)|}{|\vec{n}|} = |\vec{n}_0 \circ (\vec{x}_1 - \vec{x}_0)| \quad \text{Hessesche Normalform (mit } \vec{n}_0 = \frac{\vec{n}}{|\vec{n}|})$$



Beispiel:

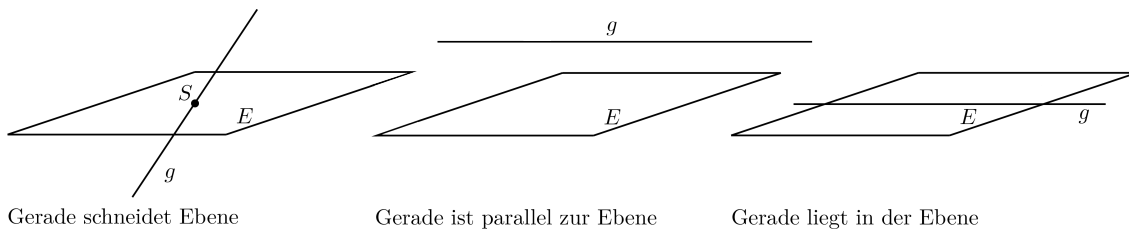
Die Ebene E enthält den Punkt $(1|0|9)$ und ihr Normalenvektor ist $\vec{n} = \begin{pmatrix} 1 \\ 3 \\ 5 \end{pmatrix}$

Bestimmen Sie den Abstand des Punktes $(-2|1|3)$ von dieser Ebene.

Abstand zwischen einer Geraden und einer Ebene

Eine Gerade und eine Ebene können

1. sich in einem Punkt schneiden, d.h. der Richtungsvektor der Geraden ist nicht senkrecht zum Normalenvektor der Ebene
2. parallel liegen, d.h. der Richtungsvektor der Geraden und der Normalenvektor der Ebene sind senkrecht
3. in einer Ebene liegen



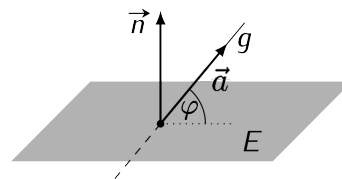
Abhängig von dieser Lage wird der Abstand bestimmt:

Für eine Gerade $\vec{x} = \vec{x}_1 + \lambda \vec{a}$ und eine Ebene $\vec{n} \circ (\vec{x} - \vec{x}_0) = 0$, die sich **schneiden**, berechnet sich der **Schnittpunkt** aus:

$$S = \vec{x}_1 + \left(\frac{\vec{n} \circ (\vec{x}_0 - \vec{x}_1)}{\vec{n} \circ \vec{a}} \right) \cdot \vec{a}$$

und der Schnittwinkel zu:

$$\varphi = \arcsin \frac{|\vec{n} \circ \vec{a}|}{|\vec{n}| |\vec{a}|}$$



Der Abstand zwischen einer Geraden mit der Gleichung

$$\vec{x} = \vec{x}_1 + \lambda \vec{a}$$

und einer zu ihr **parallel** liegenden Ebene mit der Gleichung

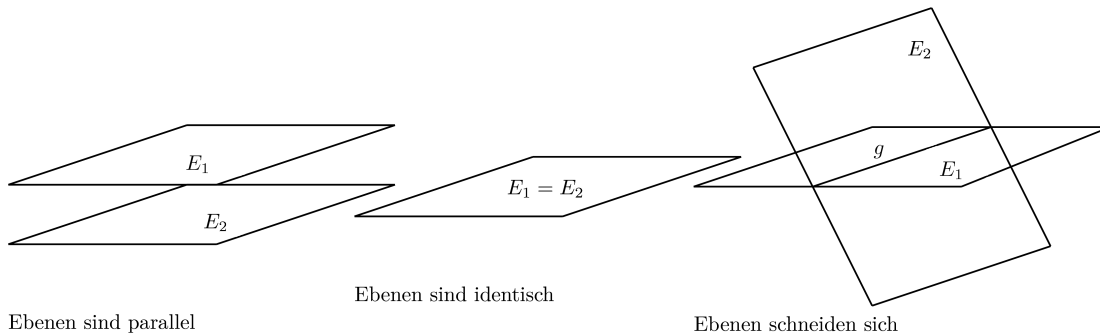
$$\vec{n} \circ (\vec{x} - \vec{x}_0) = 0 \text{ ist:}$$

$$d = \frac{|\vec{n} \circ (\vec{x}_1 - \vec{x}_0)|}{|\vec{n}|}$$

Abstand zwischen zwei Ebenen

Zwei Ebenen können

1. parallel liegen, d.h. die Normalenvektoren sind linear abhängig
2. identisch sein
3. sich in einer Geraden schneiden, d.h. die Normalenvektoren sind nicht linear abhängig



Abhängig von dieser Lage wird der Abstand bestimmt:

Der Abstand zweier paralleler Ebenen mit den Gleichungen

$\vec{n}_1 \circ (\vec{x} - \vec{x}_1) = 0$ und $\vec{n}_2 \circ (\vec{x} - \vec{x}_2) = 0$ ist:

$$d = \frac{|\vec{n}_1 \circ (\vec{x}_2 - \vec{x}_1)|}{|\vec{n}_1|}$$

Zwei Ebenen, deren Normalenvektoren nicht linear abhängig sind, **schneiden** sich in einer Geraden. Für die beiden Ebenen mit den Gleichungen

$\vec{n}_1 \circ (\vec{x} - \vec{x}_1) = 0$ und $\vec{n}_2 \circ (\vec{x} - \vec{x}_2) = 0$ lässt sich die Schnittgerade und der Schnittwinkel berechnen:

Die **Schnittgerade** ergibt sich aus:

$$\vec{x} = \vec{x}_0 + \lambda \vec{a} \quad \text{mit}$$

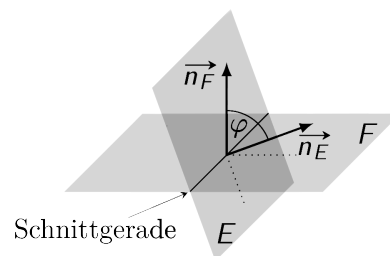
$$\vec{a} = \vec{n}_1 \times \vec{n}_2$$

und

$$\vec{n}_1 \circ (\vec{x}_0 - \vec{x}_1) = 0 \quad \text{und} \quad \vec{n}_2 \circ (\vec{x}_0 - \vec{x}_2) = 0$$

Der **Schnittwinkel** ergibt sich zu:

$$\cos \varphi = \frac{|\vec{n}_1 \circ \vec{n}_2|}{|\vec{n}_1| |\vec{n}_2|}$$



KAPITEL 2

Matrizen und Determinanten

2.1 Matrizen

2.1.1 Einführung

Eine Matrix ist ein Werkzeug, mit dessen Hilfe lineare Zusammenhänge zwischen vielen Variablen übersichtlich geschrieben und umgeformt werden können. Eine Matrix ist eine Tabelle von Zahlen oder anderen Größen.

Zum Beispiel können die Koeffizienten des Gleichungssystems:

$$\begin{aligned}2u + 5v - 3w - x + 3y + 7z &= 0 \\5u - v + x + 2y - 5z &= 0 \\4v - 2w + 3x - y + 2z &= 0\end{aligned}$$

in folgender Gestalt (Matrix) zusammengefasst werden:

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

2.1.2 Definitionen

Matrix

Unter einer Matrix vom Typ $m \times n$ versteht man ein geordnetes Schema von $m \cdot n$ Zahlen, die in m Zeilen und n Spalten dargestellt sind. Die Zahlen nennt man die Elemente der Matrix. Matrizen werden meist durch Großbuchstaben abgekürzt.

$$A = (a_{mn}) = \left(\begin{array}{cccccc} a_{11} & a_{12} & \dots & a_{1k} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2k} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ a_{j1} & a_{j2} & \dots & a_{jk} & \dots & a_{jn} \\ \dots & \dots & \ddots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mk} & \dots & a_{mn} \end{array} \right)$$

heißt $m \times n$ -Matrix.

Das Element a_{jk} steht in der j -ten Zeile und der k -ten Spalte.

Die Elemente der Matrix können reell oder komplex sein. Im Folgenden betrachten wir jedoch nur Matrizen mit reellen Elementen.

Beispiele:

1. $A = \begin{pmatrix} 2 & 5 & -3 & 1 & 3 & 7 \end{pmatrix}$ ist eine 1 $\times$ 6 -Matrix.

2. $A = \begin{pmatrix} 1 \end{pmatrix}$ ist eine 1 $\times$ 1 -Matrix.

3. $A = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$ ist eine 3 $\times$ 1 -Matrix.

4. $A = \begin{pmatrix} 1 & 2 & 4 & 0 \\ 2 & 0 & 1 & 5 \end{pmatrix}$ ist eine 2 $\times$ 4 -Matrix.

Quadratische Matrix:

Eine Matrix mit $m = n$ heißt **quadratische Matrix**.

Transponierte Matrix:

Die zu einer gegebenen Matrix A **transponierte Matrix** A^T entsteht aus A durch Vertauschen der Zeilen und Spalten.

Beispiel:

$$A = \begin{pmatrix} 1 & 4 & 7 \\ 2 & 5 & 8 \\ 3 & 6 & 9 \end{pmatrix} \quad A^T =$$

Diagonalmatrix

Eine quadratische $n \times n$ -Matrix A heißt **Diagonalmatrix**, wenn nur die Diagonalelemente a_{kk} für $k = 1, \dots, n$ ungleich 0 sind.

Beispiel:

$$A =$$

Einheitsmatrix

Eine Diagonalmatrix mit $a_{ik} = 1$ für $k = i$ und $a_{ik} = 0$ für $k \neq i$ heißt Einheitsmatrix E_n .

$$E_n =$$

Dreiecksmatrix

Eine quadratische Matrix heißt untere Dreiecksmatrix, wenn alle Elemente oberhalb der Diagonalen 0 sind: $a_{ik} = 0$ für $i < k$.

Eine quadratische Matrix heißt obere Dreiecksmatrix, wenn alle Elemente unterhalb der Diagonalen 0 sind: $a_{ik} = 0$ für $i > k$.

Beispiele: Je eine untere und eine obere Dreiecksmatrix:

$$A_u =$$

$$A_o =$$

Symmetrische Matrix

Eine quadratische $n \times n$ -Matrix heißt symmetrisch, wenn

$$a_{ik} = a_{ki} \text{ für alle } i, k = 1, \dots, n$$

gilt.

Beispiel:

$$A_s =$$

2.1.3 Rechenoperationen

Gleichheit

Zwei $m \times n$ -Matrizen sind gleich, d.h. $A = B$, falls

$$a_{ik} = b_{ik} \text{ für alle } i = 1..m \text{ und } k = 1..n \text{ gilt.}$$

Addition

Zwei $m \times n$ -Matrizen A und B werden addiert, indem man die entsprechenden Matrixelemente addiert

$$A + B = C = (c_{ik}) \text{ mit } c_{ik} = a_{ik} + b_{ik} \text{ für alle } i, k = 1, \dots, n$$

Beispiel:

$$A = \begin{pmatrix} 1 & 5 & 3 \\ 4 & 0 & 8 \end{pmatrix}, \quad B = \begin{pmatrix} 5 & 1 & 3 \\ 1 & 4 & 7 \end{pmatrix} \implies A + B =$$

Multiplikation mit einem Skalar

Eine $m \times n$ -Matrix A wird mit einem Skalar $\lambda \in \mathbb{R}$ multipliziert, indem man jedes Matrixelemente mit dem Skalar multipliziert.

$$\lambda A = (\lambda a_{ik})$$

Beispiel:

$$A = \begin{pmatrix} 1 & 5 & 3 \\ 4 & 0 & 8 \end{pmatrix} \implies 4 \cdot A =$$

Multiplikation zweier Matrizen

Eine $m \times n$ -Matrix A wird mit einer $n \times p$ -Matrix B multipliziert, indem man jede Zeile der Matrix A mit jeder Spalte der Matrix B multipliziert.

Das Ergebnis $C = A \cdot B$ ist eine $m \times p$ -Matrix mit

$$c_{ik} = \sum_{j=1}^n a_{ij} b_{jk} = a_{i1} b_{1k} + a_{i2} b_{2k} + \dots + a_{in} b_{nk}$$

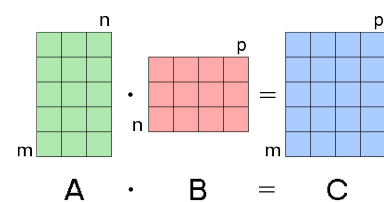
	A	$\cdot$	B	$=$	C
Typen:	$m \times n$		$n \times p$		$m \times p$

Typverträglichkeit: Spaltenzahl links = Zeilenzahl rechts

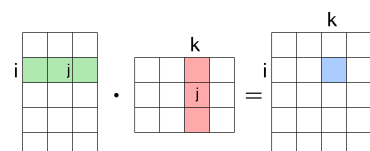
Falksches Schema:

$$\begin{pmatrix} b_{11} & \dots & b_{1k} & \dots & b_{1p} \\ a_{21} & \dots & b_{2k} & \dots & b_{2p} \\ \dots & \dots & \dots & \dots & \dots \\ b_{n1} & \dots & b_{nk} & \dots & b_{np} \end{pmatrix}$$

$$\begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ \dots & \dots & \dots & \dots \\ a_{i1} & a_{i2} & \dots & a_{in} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} c_{11} & c_{12} & \dots & c_{1p} \\ \dots & \dots & \dots & \dots \\ c_{i1} & c_{ik} & \dots & c_{ip} \\ \dots & \dots & \dots & \dots \\ c_{m1} & c_{m2} & \dots & c_{mp} \end{pmatrix}$$



Zeile mal Spalte:



$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \times \begin{bmatrix} b_{11} & b_{12} & b_{13} \\ b_{21} & b_{22} & b_{23} \\ b_{31} & b_{32} & b_{33} \end{bmatrix} \\
= \begin{bmatrix} a_{11}b_{11} + a_{12}b_{21} + a_{13}b_{31} & a_{11}b_{12} + a_{12}b_{22} + a_{13}b_{32} & a_{11}b_{13} + a_{12}b_{23} + a_{13}b_{33} \\ a_{21}b_{11} + a_{22}b_{21} + a_{23}b_{31} & a_{21}b_{12} + a_{22}b_{22} + a_{23}b_{32} & a_{21}b_{13} + a_{22}b_{23} + a_{23}b_{33} \\ a_{31}b_{11} + a_{32}b_{21} + a_{33}b_{31} & a_{31}b_{12} + a_{32}b_{22} + a_{33}b_{32} & a_{31}b_{13} + a_{32}b_{23} + a_{33}b_{33} \end{bmatrix}$$

Beispiel:

$$A = \begin{pmatrix} 1 & -3 & 2 \\ 0 & 2 & 1 \end{pmatrix}, \quad B = \begin{pmatrix} 1 \\ 5 \\ 4 \end{pmatrix}$$

$$AB =$$

$$BA =$$

2.1.4 Rechenregeln

1. Assoziativgesetz: $(A + B) + C = A + (B + C)$

$$A(BC) = (AB)C$$

2. Kommutativgesetz der Addition: $A + B = B + A$

Kommutativgesetz der Multiplikation gilt nicht,
im allgemeinen ist $AB \neq BA$. Z.B.:

$$A = \begin{pmatrix} 1 & 4 & -2 \\ 0 & 1 & 1 \\ -3 & 2 & 5 \end{pmatrix}, \quad B = \begin{pmatrix} 3 & 0 & 1 \\ -2 & 1 & 5 \\ 2 & 3 & 8 \end{pmatrix}$$

3. Distributivgesetz: $\lambda(A + B) = \lambda A + \lambda B$

$$A(B + C) = AB + AC$$

$$(A + B)C = AC + BC$$

4. Transponieren: $(A + B)^T = A^T + B^T$

$$(\lambda A)^T = \lambda A^T$$

$$(AB)^T = B^T A^T$$

5. $AE = A$

6. $EA = A$

7. $AB = 0 \nRightarrow A = 0$ oder $B = 0$, denn z.B.:

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 6 \end{pmatrix}, \quad B = \begin{pmatrix} 10 & 4 \\ -5 & -2 \end{pmatrix} \implies A \cdot B =$$

8. $AB = AC$ und $A \neq 0 \nRightarrow B = C$, denn z.B.:

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 6 \end{pmatrix}, \quad B = \begin{pmatrix} 10 & 4 \\ -5 & -2 \end{pmatrix}, \quad C = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$$

2.2 Determinanten

Eine Determinante ist eine spezielle Funktion, die einer quadratischen Matrix eine Zahl zuordnet.

- Mit Hilfe von Determinanten kann man feststellen, ob ein lineares Gleichungssystem eindeutig lösbar ist.
- Das Vorzeichen der Determinante einer Vektorbasis, gibt die Orientierung der Vektoren an
- Determinanten werden zur Berechnung von Volumina in der Vektorrechnung verwendet

2.2.1 Definition

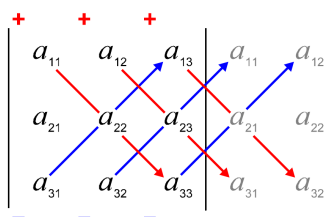
Eine Determinante ist eine Funktion auf quadratischen $n \times n$ -Matrizen A :

$\det: A \rightarrow \det(A) = |A| \in \mathbb{R}$ die folgendermaßen berechnet wird:

Für $n=1$: $A = (a_{11})$	$ A = a_{11}$
Für $n=2$: $A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$	$ A = a_{11}a_{22} - a_{12}a_{21}$
Für $n=3$: $A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$	$ 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_{11}a_{23}a_{32} - a_{12}a_{21}a_{33}$

Regel von Sarrus

Mit der Regel von Sarrus kann man sich die Gleichungen bis $n \leq 3$ sehr einfach herleiten. Dabei schreibt man die ersten beiden Spalten der Matrix rechts neben die Matrix und bildet Produkte von je 3 Zahlen, die durch die schrägen Linien verbunden sind. Dann werden die von links oben nach rechts unten verlaufenden Produkte addiert und davon die von links unten nach rechts oben verlaufenden Produkte subtrahiert.



Beispiel:

$$\begin{vmatrix} 4 & -5 & 1 \\ 0 & 4 & 2 \\ 1 & -2 & 3 \end{vmatrix} =$$

Für $n > 3$:

$$\begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{pmatrix}$$

Man entwickelt die Determinante "nach einer Zeile **oder** einer Spalte":

$$|A| = \sum_{i=1}^n (-1)^{i+j} a_{ij} |A_{ij}| \quad (\text{Entwicklung nach der } j\text{-ten Spalte})$$

$$|A| = \sum_{j=1}^n (-1)^{i+j} a_{ij} |A_{ij}| \quad (\text{Entwicklung nach der } i\text{-ten Zeile})$$

wobei A_{ij} die $(n-1) \times (n-1)$ -Untermatrix von A ist, die durch Streichen der i -ten Zeile und j -ten Spalte entsteht.

$$\begin{vmatrix} + & - & + & - & \dots \\ - & + & - & + & \dots \\ + & - & + & - & \dots \\ - & + & - & + & \dots \\ \vdots & \vdots & \vdots & \vdots & \ddots \end{vmatrix}$$

Vorzeichenschema

Beispiel:

$$\begin{vmatrix} 4 & -5 & 1 \\ 0 & 4 & 2 \\ 1 & -2 & 3 \end{vmatrix} =$$

2.2.2 Eigenschaften

Für quadratische Matrizen A und B gilt:

1. $|A| = |A^T|$, d.h. der Wert einer Determinante ändert sich nicht, wenn Zeilen und Spalten vertauscht werden.

Beispiel: $\begin{vmatrix} 8 & 5 \\ 3 & 2 \end{vmatrix} = \begin{vmatrix} 8 & 3 \\ 5 & 2 \end{vmatrix} =$

2. Bei Vertauschen zweier Zeilen (oder Spalten) ändert die Determinante ihr Vorzeichen

Beispiel: $\begin{vmatrix} 7 & 3 \\ 4 & -1 \end{vmatrix} = \begin{vmatrix} 3 & 7 \\ -1 & 4 \end{vmatrix} =$

3. Werden alle Elemente einer beliebigen Zeile (oder Spalte) einer Determinante mit einem Skalar λ multipliziert, so multipliziert sich die Determinante mit λ .

Beispiel: $\begin{vmatrix} \lambda & 1 & 1 \\ \lambda & 0 & 1 \\ \lambda & 1 & 0 \end{vmatrix} =$

4. $|\lambda A| = \lambda^n |A|$
 d.h. wird eine $n \times n$ -Matrix A mit λ multipliziert, so multipliziert sich ihre Determinante mit λ^n .

Beispiel: $\lambda A = \begin{pmatrix} \lambda & \lambda & \lambda \\ \lambda & 0 & \lambda \\ \lambda & \lambda & 0 \end{pmatrix}$

5. Der Wert einer Determinante ändert sich nicht, wenn man zu einer Zeile (oder Spalte) ein beliebig Vielfaches einer anderen Zeile (oder Spalte) elementweise addiert.

Beispiel: Addieren Sie zur 1. Zeile von $\begin{vmatrix} -6 & 5 \\ 1 & 4 \end{vmatrix}$ das 6-fache der 2. Zeile

6. Für zwei Matrizen A, B gleicher Größe gilt:

$$|AB| = |A||B|$$

$$|AB| = |BA|$$

Beispiel: Nachweis von $|AB| = |A||B|$ für $n = 2$

7. $|A^n| = |A|^n$ für $n \in \mathbb{N}$

8. Die Determinante einer $n \times n$ -Dreiecksmatrix A besitzt den Wert

$$|A| = a_{11}a_{22}a_{33}\dots a_{nn}$$

9. Eine Determinante besitzt den Wert 0, wenn sie eine der folgenden Bedingungen erfüllt:

- alle Elemente einer Zeile (oder Spalte) sind 0

Beispiel: $\begin{vmatrix} 1 & 1 & 3 \\ 0 & 0 & 0 \\ 4 & 5 & 3 \end{vmatrix} = 0$

- zwei Zeilen (oder Spalten) sind gleich

Beispiel: $\begin{vmatrix} 4 & 0 & 4 \\ 1 & 3 & 1 \\ 5 & 1 & 5 \end{vmatrix} = 0$

- zwei Zeilen (oder Spalten) sind zueinander proportional

Beispiel: $\begin{vmatrix} 1 & 4 & 1 & 0 \\ 5 & 20 & 5 & 0 \\ 1 & 1 & 2 & 3 \\ 1 & 0 & 1 & 1 \end{vmatrix} = 0$

- eine Zeile (oder Spalte) ist als Linearkombination der übrigen Zeilen (oder Spalten) darstellbar.

Beispiel: $\begin{vmatrix} 1 & 1 & 5 \\ 1 & 0 & 2 \\ 1 & 2 & 8 \end{vmatrix} = 0 \quad 2 \cdot \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + 3 \cdot \begin{pmatrix} 1 \\ 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 5 \\ 2 \\ 8 \end{pmatrix}$

2.3 Spezielle Matrizen

2.3.1 Reguläre Matrizen

Eine quadratische Matrix A heißt regulär, wenn $|A| \neq 0$, andernfalls heißt sie singulär.

Beispiel:

Die Matrix $A = \begin{pmatrix} 4 & -5 & 1 \\ 0 & 4 & 2 \\ 1 & -2 & 3 \end{pmatrix}$ ist regulär.

2.3.2 Inverse Matrizen

Gibt es zu der quadratischen Matrix A eine Matrix X mit

$$AX = XA = E,$$

so heißt X die zu A inverse Matrix, gekennzeichnet durch A^{-1} .

Anmerkungen

- Eine quadratische Matrix besitzt - wenn überhaupt - genau eine Inverse.
- Besitzt eine Matrix A eine inverse Matrix A^{-1} , so heißt A invertierbar.
- Die Matrix A ist genau dann invertierbar, wenn A regulär ist.

Berechnung der Inversen

Für eine reguläre $n \times n$ -Matrix A lässt sich die inverse Matrix folgendermaßen berechnen:

$$A^{-1} = \frac{1}{|A|} \begin{pmatrix} A_{11} & A_{21} & \dots & A_{n1} \\ A_{12} & A_{22} & \dots & A_{n2} \\ \dots & \dots & \dots & \dots \\ A_{1n} & A_{2n} & \dots & A_{nn} \end{pmatrix}$$

Dabei bedeuten:

$$A_{ik} = (-1)^{i+k} D_{ik}$$

D_{ik} : $(n-1)$ -reihige Unterdeterminante von A , d.h. in A wird die i -te Zeile und die k -te Spalte gestrichen.

Wichtig: Beachten der Zeilen- und Spaltenindizes (transponiert zur Ausgangsmatrix $a_{i',k'}$)!

Beispiel:

Die 3-reihige Matrix $\mathbf{A} = \begin{pmatrix} 1 & 0 & -1 \\ -8 & 4 & 1 \\ -2 & 1 & 0 \end{pmatrix}$ ist wegen $\det \mathbf{A} = -1 \neq 0$ *regulär*

$$\det \mathbf{A} = 1 \cdot \begin{vmatrix} 4 & 1 \\ 1 & 0 \end{vmatrix} - 0 + (-1) \cdot \begin{vmatrix} -8 & 4 \\ -2 & 1 \end{vmatrix} = -1 - (-8 + 8) = -1$$

$$D_{11} = \begin{vmatrix} 4 & 1 \\ 1 & 0 \end{vmatrix} = -1, \quad D_{12} = \begin{vmatrix} -8 & 1 \\ -2 & 0 \end{vmatrix} = 2, \quad D_{13} = \begin{vmatrix} -8 & 4 \\ -2 & 1 \end{vmatrix} = 0,$$

$$D_{21} = \begin{vmatrix} 0 & -1 \\ 1 & 0 \end{vmatrix} = 1, \quad D_{22} = \begin{vmatrix} 1 & -1 \\ -2 & 0 \end{vmatrix} = -2, \quad D_{23} = \begin{vmatrix} 1 & 0 \\ -2 & 1 \end{vmatrix} = 1,$$

$$D_{31} = \begin{vmatrix} 0 & -1 \\ 4 & 1 \end{vmatrix} = 4, \quad D_{32} = \begin{vmatrix} 1 & -1 \\ -8 & 1 \end{vmatrix} = -7, \quad D_{33} = \begin{vmatrix} 1 & 0 \\ -8 & 4 \end{vmatrix} = 4$$

$$A_{11} = +D_{11} = -1, \quad A_{12} = -D_{12} = -2, \quad A_{13} = +D_{13} = 0,$$

$$A_{21} = -D_{21} = -1, \quad A_{22} = +D_{22} = -2, \quad A_{23} = -D_{23} = -1,$$

$$A_{31} = +D_{31} = 4, \quad A_{32} = -D_{32} = 7, \quad A_{33} = +D_{33} = 4$$

Die zu $\mathbf{A}$ inverse Matrix $\mathbf{A}^{-1}$ lautet somit:

$$\begin{aligned} \mathbf{A}^{-1} &= \frac{1}{\det \mathbf{A}} \cdot \begin{pmatrix} A_{11} & A_{21} & A_{31} \\ A_{12} & A_{22} & A_{32} \\ A_{13} & A_{23} & A_{33} \end{pmatrix} = \frac{1}{-1} \cdot \begin{pmatrix} -1 & -1 & 4 \\ -2 & -2 & 7 \\ 0 & -1 & 4 \end{pmatrix} = \\ &= -1 \cdot \begin{pmatrix} -1 & -1 & 4 \\ -2 & -2 & 7 \\ 0 & -1 & 4 \end{pmatrix} = \begin{pmatrix} 1 & 1 & -4 \\ 2 & 2 & -7 \\ 0 & 1 & -4 \end{pmatrix} \end{aligned}$$

2.3.3 Orthogonale Matrizen

Eine quadratische Matrix A heißt orthogonal, wenn

$$\boxed{A^{-1} = A^T} \quad \text{bzw. mit} \quad A \cdot A^{-1} = A \cdot A^T \quad \implies \quad \boxed{E = A \cdot A^T}$$

Beispiele:

1. $A = E$
2. $A = \begin{pmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{pmatrix}$

Eigenschaften

1. Betrachtet man die Zeilen (bzw. Spalten) einer orthogonalen Matrix als Vektoren, so sind diese orthonormiert.
2. Für orthogonale Matrizen ist $|A| = 1$ oder $|A| = -1$ denn: $1 = |E| = |AA^{-1}|$.
3. Die Umkehrung obiger Aussage gilt nicht, d.h. es gibt Matrizen A mit $|A| = -1$ oder $|A| = 1$, die nicht orthogonal sind.

2.4 Rang einer Matrix

Eine beliebige $n \times m$ -Matrix A kann man als Anordnung von Spaltenvektoren bzw. Zeilenvektoren sehen. Diese Vektoren können linear abhängig oder linear unabhängig sein.

2.4.1 Definition

Die größte Anzahl r linear unabhängiger Spaltenvektoren einer $n \times m$ -Matrix A bezeichnet man als

Rang von A , kurz: $\text{rg}(A)$ (engl. rank)

2.4.2 Regeln

1. Die größte Anzahl linear unabhängiger Zeilenvektoren ist ebenfalls r .

Beispiel: $A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}$

2. Die $n \times n$ -Matrix A ist genau dann invertierbar, wenn $\text{rg}(A) = n$ gilt.
3. Der Rang einer Matrix ändert sich nicht, wenn:
 - zwei Zeilen (oder Spalten) miteinander vertauscht werden
 - die Elemente einer Zeile (oder Spalte) mit einem Skalar $\neq 0$ multipliziert werden
 - zu einer Zeile (oder Spalte) ein Vielfaches einer anderen Zeile (oder Spalte) addiert wird

2.4.3 Rangbestimmung

Die $m \times n$ -Matrix A wird durch Umformungen (s.o.), die den Rang nicht ändern, auf Trapezform gebracht:

$$\begin{array}{cccccc}
 b_{11} & b_{12} & \dots & b_{1r} & b_{1,r+1} & \dots & b_{1n} \\
 0 & b_{22} & \dots & b_{2r} & b_{2,r+1} & \dots & b_{2n} \\
 \dots & \dots & \dots & \dots & \dots & \dots & \dots \\
 0 & 0 & \dots & b_{rr} & b_{r+1,r+1} & \dots & b_{rn} \\
 0 & 0 & \dots & 0 & 0 & \dots & 0 \\
 \dots & \dots & \dots & \dots & \dots & \dots & \dots \\
 0 & 0 & \dots & 0 & 0 & \dots & 0
 \end{array}$$

mit $m - r$ Nullzeilen.

Der Rang von A ist gleich der Anzahl r der nicht verschwindenden Zeilen, $\text{rg}(A) = r$

Beispiele:

Bestimmen Sie den Rang nachfolgender Matrizen:

$$1. \ A = \begin{pmatrix} 1 & 1 & 1 & 0 \\ 2 & 1 & 1 & 3 \\ 1 & 2 & 0 & 3 \end{pmatrix} \implies \text{rg}(A) =$$

$$2. \ A = \begin{pmatrix} 1 & 3 & -5 & 0 \\ 2 & 7 & -8 & 7 \\ -1 & 0 & 11 & 21 \end{pmatrix} \implies \text{rg}(A) =$$

KAPITEL 3

Lineare Gleichungssysteme

3.1 Einführung

Ein lineares Gleichungssystem mit m Gleichungen und n Unbekannten hat folgende Gestalt:

$$\begin{array}{ccccccccc} a_{11}x_1 & + & a_{12}x_2 & + & \dots & + & a_{1n}x_n & = & c_1 \\ a_{21}x_1 & + & a_{22}x_2 & + & \dots & + & a_{2n}x_n & = & c_2 \\ \dots & & & & & & & & \\ a_{m1}x_1 & + & a_{m2}x_2 & + & \dots & + & a_{mn}x_n & = & c_n \end{array}$$

Folgende Fragen werden in diesem Kapitel beantwortet:

- Existenz einer Lösung, d.h. gibt es eine Lösung?
- Dimension der Lösungsmenge
- Lösungsalgorithmus

Die Antworten werden mit Hilfe der Matrizeneigenschaften gegeben, denn obiges Gleichungssystem lässt sich durch Matrizen beschreiben

$$Ax = c$$

mit $A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & & & \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}$ $x = \begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{pmatrix}$ und $c = \begin{pmatrix} c_1 \\ c_2 \\ \dots \\ c_n \end{pmatrix}$

3.2 Definitionen

- **homogenes Gleichungssystem**

$$Ax = 0$$

- **inhomogenes Gleichungssystem**

$$Ax = c \text{ mit } c \neq 0$$

- **quadratisches Gleichungssystem**

A ist eine quadratische Matrix

Beispiele:

$$1. \begin{aligned} x_1 - 2x_2 + x_3 &= 1 \\ x_1 + x_2 - 4x_3 &= 8 \end{aligned}$$

$$A =$$

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

$$A =$$

$$3. \begin{aligned} x_1 - 2x_2 + x_3 &= 1 \\ x_1 + x_2 - 4x_3 &= 8 \\ x_1 + 3x_2 + x_3 &= 2 \\ x_1 + x_2 - 4x_3 &= 0 \end{aligned}$$

$$A =$$

3.3 Lösungsverhalten**3.3.1 Lösbarkeit**

- Ein homogenes lineares Gleichungssystem ist immer lösbar. Eine Lösung ist

$$\vec{x} = \vec{0}$$

- Ein beliebiges lineares Gleichungssystem ist genau dann lösbar, wenn

$$\text{rg}(A) = \text{rg}(A|\vec{c})$$

$$\text{mit } \text{rg}(A|\vec{c}) = \text{rg} \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} & c_1 \\ a_{21} & a_{22} & \dots & a_{2n} & c_2 \\ \dots & & & & \\ a_{m1} & a_{m2} & \dots & a_{mn} & c_m \end{pmatrix}$$

3.3.2 Lösungsmenge

Falls das lineare Gleichungssystem, mit n Unbekannten, lösbar ist, lässt sich eine Aussage über die Dimension der Lösungsmenge machen:

- $\text{rg}(A) = \text{rg}(A|\vec{c}) = n$

dann besitzt das lineare Gleichungssystem genau eine Lösung

- $\text{rg}(A) = \text{rg}(A|\vec{c}) < n$

dann besitzt das lineare Gleichungssystem unendliche viele Lösungen

$\text{rg}(A) = \text{rg}(A \vec{c}) = n - 1$	Lösungsmenge ist eindimensional
$\text{rg}(A) = \text{rg}(A \vec{c}) = n - 2$	Lösungsmenge ist zweidimensional
...	...

- $\text{rg}(A) \neq \text{rg}(A|\vec{c})$

dann besitzt das lineare Gleichungssystem keine Lösung

3.3.3 Lösungsberechnung - Cramersche Regel

Die Cramersche Regel oder Determinantenmethode ist eine mathematische Formel für die Lösung eines linearen Gleichungssystems. Sie ist bei der theoretischen Betrachtung linearer Gleichungssysteme hilfreich.

Die Cramersche Regel ist nach Gabriel Cramer benannt, der sie im Jahr 1750 veröffentlichte, jedoch wurde sie bereits vorher von Leibniz gefunden.

Gegeben sei ein lineares Gleichungssystem der Dimension $n \times n$ in Matrixschreibweise

$$Ax = b.$$

Ist die quadratische Koeffizientenmatrix A regulär, also $\det(A) \neq 0$, dann ist das Gleichungssystem eindeutig lösbar und die Komponenten x_i des eindeutig bestimmten Lösungsvektors x sind gegeben durch:

$$x_i = \frac{\det(A_i)}{\det(A)} \text{ für alle } i.$$

Hierbei ist A_i die Matrix, die gebildet wird, indem die i -te Spalte der Koeffizientenmatrix A durch die rechte Seite des Gleichungssystems b ersetzt wird:

$$A_i = (a_{mn}) = \begin{pmatrix} a_{1,1} & \dots & a_{1,i-1} & \mathbf{b_1} & a_{1,i+1} & \dots & a_{1,n} \\ a_{2,1} & \dots & a_{2,i-1} & \mathbf{b_2} & a_{2,i+1} & \dots & a_{2,n} \\ \vdots & & \vdots & \vdots & \vdots & & \vdots \\ a_{n,1} & \dots & a_{n,i-1} & \mathbf{b_n} & a_{n,i+1} & \dots & a_{n,n} \end{pmatrix}$$

Nachteil der Cramerschen Regel:

Für die Berechnung einer Lösung ist der Rechenaufwand jedoch in der Regel zu hoch.

Bei der Berechnung einer $n \times n$ -Matrix auf einem Rechner mit 10^8 Gleitkommaoperationen pro Sekunde (100 Mflops) würden sich die folgenden Rechenzeiten ergeben:

n	10	12	14	16	18	20
Rechenzeit	0.4 s	1 min	3,6 h	41 Tage	38 Jahre	16000 Jahre

Beispiel:

Cramersche Regel an einem Beispiel von 2 Gleichungen mit 2 Unbekannten: Wir betrachten das Gleichungssystem

$$\begin{aligned} 5x_1 + 2x_2 &= 1 \\ 2x_1 + 4x_2 &= 6 \end{aligned}$$

mit

$$A = \begin{pmatrix} 5 & 2 \\ 2 & 4 \end{pmatrix} \quad \text{und} \quad b = \begin{pmatrix} 1 \\ 6 \end{pmatrix}.$$

Wollen wir die Cramersche Regel zur Lösung dieses Gleichungssystems benutzen, so benötigen wir die drei Determinanten $\det A$, $\det A_1 \begin{pmatrix} 1 \\ 6 \end{pmatrix}$ und $\det A_2 \begin{pmatrix} 1 \\ 6 \end{pmatrix}$, wie sie zu Beginn dieses Abschnitts eingeführt wurden

$$\begin{aligned} x_1 &= \frac{\det A_1 \begin{pmatrix} 1 \\ 6 \end{pmatrix}}{\det A} = \frac{\begin{vmatrix} 1 & 2 \\ 6 & 4 \end{vmatrix}}{\begin{vmatrix} 5 & 2 \\ 2 & 4 \end{vmatrix}} = \frac{1 \cdot 4 - 2 \cdot 6}{5 \cdot 4 - 2 \cdot 2} = \frac{4 - 12}{20 - 4} = \frac{-8}{16} = -\frac{1}{2} \\ x_2 &= \frac{\det A_2 \begin{pmatrix} 1 \\ 6 \end{pmatrix}}{\det A} = \frac{\begin{vmatrix} 5 & 1 \\ 2 & 6 \end{vmatrix}}{\begin{vmatrix} 5 & 2 \\ 2 & 4 \end{vmatrix}} = \frac{5 \cdot 6 - 1 \cdot 2}{5 \cdot 4 - 2 \cdot 2} = \frac{30 - 2}{20 - 4} = \frac{28}{16} = \frac{7}{4}. \end{aligned}$$

3.3.4 Lösungsberechnung - Gaußscher Algorithmus

Das gaußsche Eliminationsverfahren oder einfach Gauß-Verfahren (nach Carl Friedrich Gauß) ist ein Algorithmus zum Lösen linearer Gleichungssysteme.

Durch schrittweises Eliminieren von Unbekannten aus einem gegebenen System wird ein System in gestaffelter Form erzeugt, aus dem rückwärts rechnend die Unbekannten bestimmt werden können. Erlaubt sind dabei folgende Umformungen:

- Zwei Gleichungen dürfen miteinander vertauscht werden
- Jede Gleichung darf mit einem beliebigen Skalar $\neq 0$ multipliziert werden
- Zu jeder Gleichung darf ein beliebig Vielfaches einer anderen Gleichung addiert werden

Es gibt verschiedene Varianten des Gauß-Algorithmus, die hier vorgestellte ist die Sukzessive Elimination und Substitution. Das bedeutet, dass zunächst in der Eliminationsphase im Tableau eine Dreiecksform hergestellt wird, sodass eine Variable abgelesen werden kann. Die Dreiecksform kann implizit oder explizit hergestellt werden (hier explizit).

x_1	x_2	x_3	x_4	RS
a_{11}	a_{12}	a_{13}	a_{14}	b_1
0	a_{22}	a_{23}	a_{24}	b_2
0	0	a_{33}	a_{34}	b_3
0	0	0	1	b_4

→

x_1	x_2	x_3	x_4	RS
1	0	0	0	b_1
0	1	0	0	b_2
0	0	1	0	b_3
0	0	0	1	b_4

Gauß
Gauß-Jordan

Gauß-Jordan Verfahren:

Dies ist eine Erweiterung des gaußschen Eliminationsverfahrens, bei dem in einem zusätzlichen Schritt das Gleichungssystem bzw. dessen erweiterte Koeffizientenmatrix auf die reduzierte Stufenform gebracht wird. Daraus lässt sich dann die Lösung direkt ablesen.

Nachteile des Gaußschen Algorithmus:

Das Gaußsche Eliminationsverfahren ist ein direktes Verfahren, das i.a. verwendet wird, falls A eine vollbesetzte Matrix ist. Bei vollbesetzten Matrizen wächst der Rechenaufwand des Gaußschen Eliminationsverfahrens mit bis zu der dritten Potenz der Anzahl der Unbekannten (Schreibweise: $W = O(N^2)$ bis (N^3)).

Ablauf anhand eines Beispiels

Aus folgendem Gleichungssystem sollen die Unbekannten x_1, x_2, x_3, x_4 bestimmt werden

$$\begin{aligned} 2x_1 + 4x_2 + 6x_3 + 2x_4 &= 14 \\ x_1 - x_2 + x_3 - x_4 &= 10 \\ 4x_1 + 2x_2 + 14x_3 + 2x_4 &= -4 \\ 2x_1 + 7x_2 + 10x_3 - x_4 &= 4 \end{aligned}$$

1. Tabellenschema erstellen

$$\begin{array}{rrrrr} 2 & 4 & 6 & 2 & 14 \\ 1 & -1 & 1 & -1 & 10 \\ 4 & 2 & 14 & 2 & -4 \\ 2 & 7 & 10 & -1 & 4 \end{array}$$

2. Vereinfachung durch zeilenweises Kürzen

$$\begin{array}{rrrrr} 1 & 2 & 3 & 1 & 7 \\ 1 & -1 & 1 & -1 & 10 \\ 2 & 1 & 7 & 1 & -2 \\ 2 & 7 & 10 & -1 & 4 \end{array}$$

3. Eliminieren der 1. Unbekannten mit Hilfe der 1. Gleichung

Die letzte Spalte beschreibt den Rechenvorgang

$$\begin{array}{rrrrr} 1 & 2 & 3 & 1 & 7 & \text{I} \\ 0 & -3 & -2 & -2 & 3 & \text{II-I} \\ 0 & -3 & 1 & -1 & -16 & \text{III-2·I} \\ 0 & 3 & 4 & -3 & -10 & \text{IV-2·I} \end{array}$$

4. Eliminieren der 2. Unbekannten mit Hilfe der 2. Gleichung

$$\begin{array}{rrrrr} 1 & 2 & 3 & 1 & 7 & \text{I} \\ 0 & -3 & -2 & -2 & 3 & \text{II} \\ 0 & 0 & 3 & 1 & -19 & \text{III-II} \\ 0 & 0 & 2 & -5 & -7 & \text{IV+II} \end{array}$$

5. Eliminieren der 3. Unbekannten mit Hilfe der 3. Gleichung

$$\begin{array}{rrrrr} 1 & 2 & 3 & 1 & 7 & \text{I} \\ 0 & -3 & -2 & -2 & 3 & \text{II} \\ 0 & 0 & 3 & 1 & -19 & \text{III} \\ 0 & 0 & 0 & -17 & 17 & 3\cdot\text{IV}-2\cdot\text{III} \end{array}$$

6. Schrittweises Einsetzen von unten nach oben $-17x_4 = 17 \Rightarrow x_4 = -1$

$$3x_3 - 1 = -19 \Rightarrow x_3 = -6$$

$$-3x_2 + 12 + 2 = 3 \Rightarrow x_2 = \frac{11}{3}$$

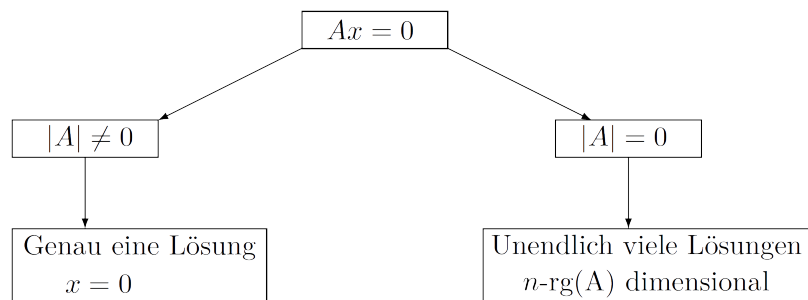
$$x_1 + \frac{22}{3} - 18 - 1 = 7 \Rightarrow x_1 = \frac{56}{3}$$

3.4 Spezialfall: quadratische lineare Gleichungssysteme

Von einem quadratischen Gleichungssystem ist die Rede, wenn die Zahl der Unbekannten gleich der Zahl der Gleichungen ist. Ein Gleichungssystem dieser Form kann, wenn die Zeilen oder Spalten linear unabhängig sind, eindeutig gelöst werden.

Die Besonderheit der quadratischen linearen Gleichungssysteme liegt in der Möglichkeit, hierfür die Determinante berechnen zu können, und diese zur Untersuchung des Lösungsverhaltens heranziehen zu können.

3.4.1 homogenes quadratisches lineares GLS



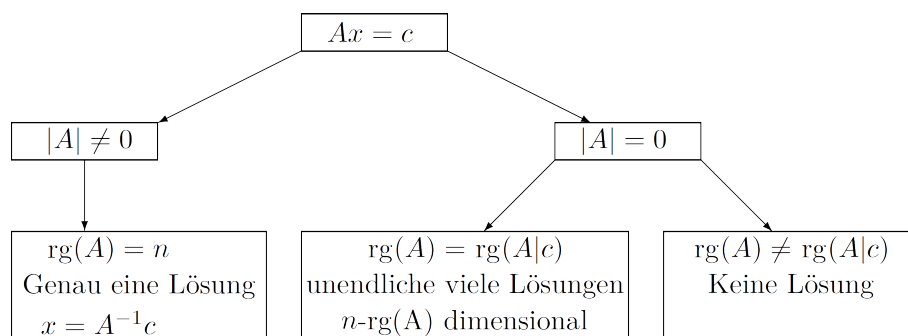
Beispiel:

$$2x_1 + 5x_2 - 3x_3 = 0$$

$$4x_1 - 4x_2 + x_3 = 0$$

$$4x_1 - 2x_2 = 0$$

3.4.2 inhomogenes quadratisches lineares GLS



Beispiel:

Prüfen Sie nachfolgende linearen Gleichungssysteme auf Lösbarkeit und berechnen Sie die Lösung mit Hilfe des Gaußschen Algorithmus.

$$2x_1 + 3x_2 + 2x_3 = 2$$

$$-x_1 - x_2 - 3x_3 = -5$$

$$3x_1 + 5x_2 + 5x_3 = 3$$

3.5 Rundungsfehler

Rundungsfehler können bei linearen Gleichungssystemen, wenn man unbedacht vorgeht, manchmal einen katastrophalen Einfluss haben. Man kann das bereits an sehr kleinen Gleichungssystemen mit wenigen Unbekannten beobachten

Beispiel:

Wir betrachten das Gleichungssystem $Ax = b$ mit

$$\left(\begin{array}{ccc|c} 2 & 1.01 & 2.52 & 9.57 \\ 0.4 & 0.203 & -1.8 & -0.385 \\ 0.6 & -1.05 & 0.8 & -3.85 \end{array} \right)$$

Das Gaußsche Eliminationsverfahren liefert bei **4-stelliger** Genauigkeit, d.h. mit dezimalen Gleitkommazahlen, deren Mantisse vierstellig ist, die exakte Lösung:

$$x_1 = 1, \quad x_2 = 5, \quad x_3 = 1.$$

Wird die gleiche Rechnung jedoch mit **3-stelligen** Gleitkommazahlen durchgeführt, so kommt folgende, völlig unsinnige Lösung heraus:

$$\hat{x}_1 = 3.53, \quad \hat{x}_2 = 0, \quad \hat{x}_3 = 1.4$$

Mit geeigneten Varianten des Gaußschen Eliminationsverfahrens kann man derartige "Katastrophen" verhindern, und daran ist man in der Praxis natürlich interessiert.

Wir werden hier keine derartigen problematischen Gleichungssysteme betrachten. Trotzdem ist es interessant zu wissen, wann dieses Problem auftreten kann.

Bei zwei Gleichungen mit zwei Unbekannten beschreibt jede Gleichung eine Gerade. Sind diese Geraden **nahezu parallel**, dann können winzige Änderungen in den Koeffizienten der Geradengleichung (oder entsprechend im gegebenen Gleichungssystem) dazu führen, dass sich der Schnittpunkt der Geraden (d.h. die Lösung des Gleichungssystems) erheblich verschiebt.

Auch bei der zeichnerischen Lösung wird in derartigen Fällen die genaue Bestimmung des Schnittpunktes zweier Geraden problematisch.

Analog ist es im dreidimensionalen Fall kritisch, wenn die Ebenen, die durch die Gleichungen beschrieben werden, fast parallel sind, oder wenn die Schnittgeraden von je zwei Ebenen nahezu parallel sind.

3.6 Geschwindigkeit der Verfahren

Iterative Lösungsverfahren:

In der Praxis hat A häufig eine spezielle Struktur und/oder ist schwach besetzt (engl. sparse), d.h. die meisten Elemente von A sind 0. N ist eventuell sehr groß, z.B. 10^6 .

Dies bedeutet, dass man Gleichungssysteme mit einer Million von Unbekannten (oder mehr) zu lösen hat.

In solchen Fällen sind iterative Lösungsverfahren eine interessante Alternative. Bei ihnen startet man mit irgendeiner Startnäherung für die Unbekannten (im Zweifelsfall nimmt man an, dass alle Unbekannten 0 sind) und reduziert den Fehler sukzessiv durch eine geeignete Iterationsvorschrift.

Als Beispiele sind das **Jacobi-Verfahren**, **Gauss-Seidel**- sowie **SOR** (*Successive OverRelaxation*) Verfahren genannt.

Aktuelle **Mehrgitterverfahren** bilden in der numerischen Mathematik eine Klasse von effizienten Algorithmen zur näherungsweisen Lösung von Gleichungssystemen, die aus der Diskretisierung partieller Differentialgleichungen stammen.

Es wird dazu auf die weiterführende Literatur verwiesen.

Geschwindigkeitsvergleich:

Betrachten wir ein lineares Gleichungssystem mit $N = 10^6$ Unbekannten. Die folgende Tabelle vergleicht die für die Lösung dieses Gleichungssystems benötigte Anzahl an Rechenoperationen und die benötigte Rechenzeit für verschiedene Verfahren auf einem Standard-PC.

Sie gibt anschaulich die Leistungssteigerung wieder, die mit Hilfe von neuen Methoden der Numerischen Mathematik erzielt wurde.

Verfahren	Anzahl der Operationen	Rechenzeit
Cramersche Regel	$\sim N!$	" ∞ "
Gaußsches Eliminationsverfahren für Bandmatrizen	$\sim N^2$	14 h
Überrelaxationsverfahren, SOR (1960)	$\sim N^{1.5}$	5 min
Mehrgitter-Verfahren (1980)	$\sim N$	1 sec

KAPITEL 4

Lineare Abbildungen

Lineare Abbildungen, wie z.B. Projektionen, Drehungen, Spiegelungen, finden in zahlreichen Bereichen Anwendung (z.B. Zoomen des Bildschirminhalts). In diesem Kapitel untersuchen wir ganz allgemein lineare Abbildungen, die Vektoren aus dem n -dimensionalen Raum ($\mathbb{R}^n$) in einen m -dimensionalen Raum ($\mathbb{R}^m$) abbilden.

4.1 Definitionen

4.1.1 n -dimensionaler reeller Koordinatenraum

Der n -dimensionale reelle Koordinatenraum $\mathbb{R}^n$ ist die Menge der n -Tupel $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$

mit $x_i \in \mathbb{R}$. Die Elemente des $\mathbb{R}^n$ bezeichnet man als Punkte oder als Vektoren. Wie bei Vektoren gilt die Addition, die Multiplikation mit einem Skalar und das Skalarprodukt.

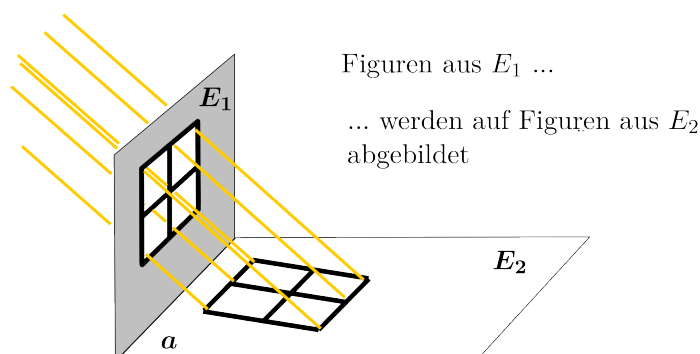
4.1.2 Lineare Abbildung

Unter einer linearen Abbildung

$$y = Ax$$

versteht man die Transformation eines Vektors $x \in \mathbb{R}^n$ in einen Vektor $y \in \mathbb{R}^m$ durch Multiplikation mit einer Matrix $A_{m,n}$.

A heit Abbildungsmatrix.



4.2 Beispiele

4.2.1 $\mathbb{R} \rightarrow \mathbb{R}$

Bekannt sind die linearen Abbildungen von $\mathbb{R}$ in $\mathbb{R}$. Letztere sind alle Funktionen f mit

$$f(x) = mx \quad \text{mit } m \in \mathbb{R}, \text{ für alle } x \in \mathbb{R}$$

Mit der 1×1 -Matrix $A = (m)$ lässt sich dies in der Form:

$$y = Ax$$

schreiben.

4.2.2 $\mathbb{R}^2 \rightarrow \mathbb{R}$

Betrachtet man lineare Abbildungen vom $\mathbb{R}^2$ in $\mathbb{R}$, so können diese folgendermaßen beschrieben werden:

$$y = a_1x_1 + a_2x_2 \quad \text{bzw. } y = Ax \quad \text{mit } A = \begin{pmatrix} a_1 & a_2 \end{pmatrix} \quad \text{und } x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

Zum Beispiel ist $y = x_1 + x_2$ die Abbildung, die jedem Punkt $(x_1|x_2)$ der reellen Ebene den reellen Wert y zuordnet.

4.2.3 $\mathbb{R}^2 \rightarrow \mathbb{R}^2$

Betrachtet man lineare Abbildungen vom $\mathbb{R}^2$ in $\mathbb{R}^2$, so können diese folgendermaßen beschrieben werden:

$$\begin{aligned} y_1 &= a_{11}x_1 + a_{12}x_2 \\ y_2 &= a_{21}x_1 + a_{22}x_2 \end{aligned}$$

bzw.

$$y = Ax \quad \text{mit } y = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}, \quad A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \quad \text{und } x = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

Die Abbildungsmatrix $A = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$

bildet jeden Punkt der Ebene auf sich selbst ab.

Die Abbildungsmatrix $A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$

projiziert jeden Punkt der Ebene senkrecht auf die x -Achse.

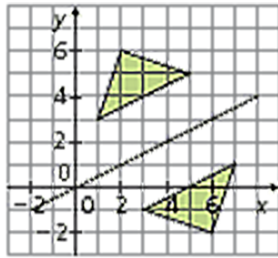
Die Abbildungsmatrix $A = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}$

spiegelt jeden Punkt der Ebene an der y -Achse.

Abbildung

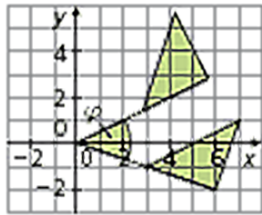
Matrix

Spiegelung an der Geraden mit der Gleichung $y = m \cdot x$ mit $m = \tan \varphi$.



$$A = \begin{pmatrix} \cos 2\varphi & \sin 2\varphi \\ \sin 2\varphi & -\cos 2\varphi \end{pmatrix}$$

Drehung um den Ursprung mit Drehwinkel φ .



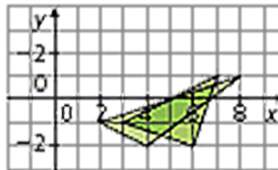
$$A = \begin{pmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{pmatrix}$$

Zentrische Streckung mit dem Ursprung als Zentrum und dem Faktor k ($k \neq 0$).



$$A = \begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix}$$

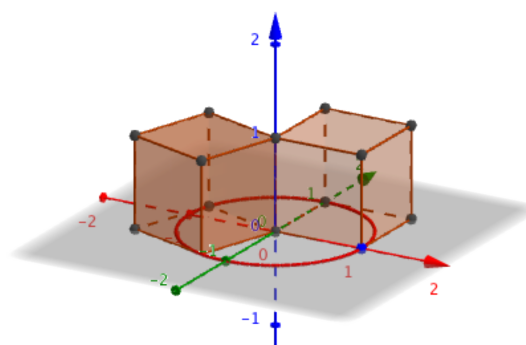
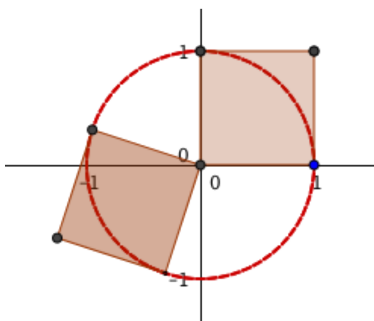
Scherung mit dem Scherungswinkel φ und der x -Achse als Scherungsachse.



$$A = \begin{pmatrix} 1 & \tan \varphi \\ 0 & 1 \end{pmatrix}$$

4.2.4 $\mathbb{R}^3 \rightarrow \mathbb{R}^3$

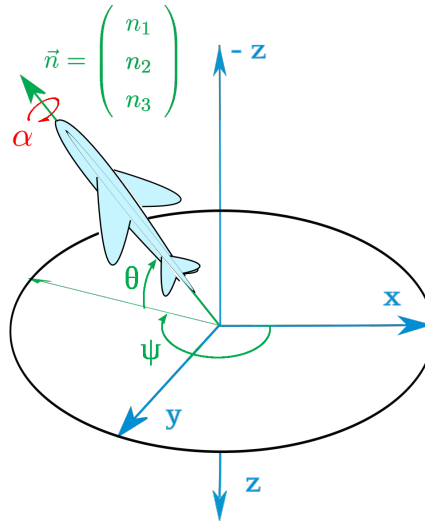
Drehungen im $\mathbb{R}^3$ werden in Drehungen um die Achsen zerlegt.



Drehung um die z -Achse:

$$\text{Drehmatrix } D_z = \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) & 0 \\ \sin(\alpha) & \cos(\alpha) & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Eine wichtige lineare Abbildung im $\mathbb{R}^3$ ist die Drehung um eine Drehachse $\vec{n} = \begin{pmatrix} n_1 \\ n_2 \\ n_3 \end{pmatrix}$ um den Winkel α :



Die Drehmatrix der zusammengesetzten Drehung $D_{\vec{n}}$ erhält man durch Matrixmultiplikation aus den Matrizen der einzelnen Drehungen um die x, y, z -Achsen.

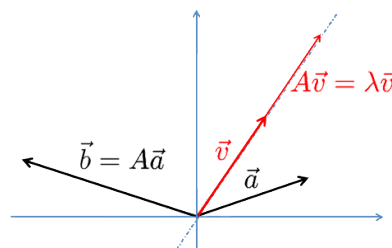
$$D_{\vec{n}} = \begin{pmatrix} n_1^2(1 - \cos\alpha) + \cos\alpha & n_1n_2(1 - \cos\alpha) - n_3\sin\alpha & n_1n_3(1 - \cos\alpha) + n_2\sin\alpha \\ n_2n_1(1 - \cos\alpha) + n_3\sin\alpha & n_2^2(1 - \cos\alpha) + \cos\alpha & \cos\alpha(1 - \cos\alpha) - n_1\sin\alpha \\ n_3n_1(1 - \cos\alpha) - n_2\sin\alpha & n_3n_2(1 - \cos\alpha) + n_1\sin\alpha & n_3^2(1 - \cos\alpha) + \cos\alpha \end{pmatrix}$$

4.3 Eigenwerte, Eigenvektoren quadratischer Matrizen

Quadratische Matrizen beschreiben Abbildungen von einem Koordinatenraum in denselben Koordinatenraum. Hierbei führt folgende Fragestellung zur Betrachtung von Eigenvektoren und Eigenwerten:

Die Darstellung ein und derselben physikalischen Größe ist in verschiedenen Koordinatensystemen verschieden. Daher die Frage nach einem dem physikalischen Vorgang, bzw. der physikalischen Größe besonders angepassten Koordinatensystem.

Betrachtet man Abbildungen von einem Koordinatenraum in sich, so ist ein **Eigenvektor** dieser Abbildung ein vom Nullvektor verschiedener Vektor, dessen Richtung durch die Abbildung nicht verändert wird. Ein Eigenvektor wird also nur gestreckt. Man bezeichnet den Streckungsfaktor als **Eigenwert** zum Eigenvektor.



Eigenwerte charakterisieren wesentliche Eigenschaften linearer Abbildungen, etwa ob ein entsprechendes lineares Gleichungssystem eindeutig lösbar ist oder nicht. In vielen Anwendungen beschreiben Eigenwerte auch physikalische Eigenschaften eines mathematischen Modells.

Anwendungsbeispiele:

Schwingungsfähige Systeme besitzen bevorzugte Frequenzen - Resonanzfrequenzen -, die durch Eigenvektoren beschrieben werden können.

Erwünschte Resonanzfrequenzen: Musikinstrumente

Unerwünschte Resonanzfrequenzen: Eigenschwingungen von Bauwerken

- Im Jahr 1850 marschierten 730 französische Soldaten im Gleichschritt über die Hängebrücke von Angers. Die Brücke geriet in heftige Schwingungen und stürzte ein. Es ist heute verboten, vgl. §27 StVO, im Gleichschritt über eine Brücke zu marschieren. (1883 → Broughton Suspension Bridge, Manchester)
- Die Hängebrücke über den Tacoma Narrows stürzte 1940 ein, nachdem sie durch den Wind zu immer stärkeren Schwingungen angeregt wurde.

Zum Schutz vor solchen Resonanzkatastrophen werden Konstruktionen auf eine Eigenschwingung ausgelegt, die typischerweise nicht im Betrieb auftritt. In Erdbebengebieten richtet man sich dabei an die lokal typischen Schwingungsfrequenzen der Erderschütterung.

Die Eigenwerttheorie liefert mathematische Lösungsmethoden für diese und folgende Themen

- Diagonalisierung symmetrischer Matrizen; vgl. Spannungstensoren
- Normalformen von Kegelschnitten, Ellipsoiden, etc.
- Lösungen linearer Differentialgleichungssysteme, zB bei Schwingungen

4.3.1 Definitionen

Ist A eine Abbildungsmatrix vom $\mathbb{R}^n$ auf sich, und ist $0 \neq x \in \mathbb{R}^n$ mit

$$Ax = \lambda x \quad \text{mit } \lambda \in \mathbb{R}$$

So heißt x Eigenvektor zum Eigenwert

4.3.2 Berechnung

Methode zur Berechnung der Eigenwerte und Eigenvektoren einer quadratischen Matrix:

Es sei A eine $n \times n$ -Matrix.

Für die Eigenvektoren und Eigenwerte der Matrix A gilt:

$$\begin{aligned} Ax &= \lambda x \\ \iff Ax - \lambda x &= 0 \\ \iff Ax - \lambda Ex &= 0 \\ \iff (A - \lambda E)x &= 0 \end{aligned}$$

Dies ist ein homogenes quadratisches lineares Gleichungssystem, welches immer $x = 0$ als Lösung hat.

Nicht triviale Lösungen hat dieses Gleichungssystem genau dann wenn gilt:

$$|(A - \lambda E)| = 0$$

Die Auflösung dieser Determinante liefert ein Polynom n -ten Grades für λ , das sogenannte **charakteristische Polynom**.

Das charakteristische Polynom, das für **quadratische Matrizen** von endlichdimensionalen Vektorräumen definiert ist, gibt Auskunft über Eigenschaften einer Matrix oder einer linearen Abbildung.

Die Lösungen dieses Polynoms sind n (reelle und/oder imaginäre) **Eigenwerte**:

$$\lambda_1, \lambda_2, \lambda_3, \dots, \lambda_n$$

Zu jedem dieser Eigenwerte kann der korrespondierende Eigenvektor berechnet werden, durch Lösen des linearen Gleichungssystems

$$Ax - \lambda_i x = 0, \quad i = 1, 2, \dots, n$$

Beispiel:

Berechnen Sie das charakteristische Polynom sowie die Eigenwerte und Eigenvektoren

folgender Matrix: $A = \begin{pmatrix} 1 & 3 \\ 0 & 2 \end{pmatrix}$

Bemerkungen

1. Ist x ein Eigenvektor zum Eigenwert λ , so ist auch kx ein Eigenvektor zum Eigenwert λ
2. Als Lösung wird üblicherweise (vgl. Punkt 1) ein Vektor exemplarisch angegeben, oder sogar der normierte Vektor.
3. Sind alle Eigenwerte voneinander verschieden, so gehört zu jedem Eigenwert genau ein linear unabhängiger Eigenvektor.
4. Tritt ein Eigenvektor k -fach auf, so gehören hierzu mindestens ein, höchstens k linear unabhängige Eigenvektoren.
5. Die zu verschiedenen Eigenwerten gehörenden Eigenvektoren sind immer linear unabhängig

4.3.3 EW und EV einer Dreiecksmatrix

Ist A eine $n \times n$ -Dreiecksmatrix, so sind die Eigenwerte identisch mit den Hauptdiagonalelementen der Matrix A :

$$\lambda_i = a_{ii} \quad \text{für } i = 1, 2, \dots, n$$

4.3.4 EW und EV einer symmetrischen Matrix

Die Eigenwerte und Eigenvektoren einer symmetrischen $n \times n$ -Matrix A besitzen folgende Eigenschaften:

- alle Eigenwerte sind reell
- es gibt genau n linear unabhängige Eigenvektoren
- zu jedem einfachen Eigenwert gehört genau ein linear unabhängiger Eigenvektor, zu jedem k -fachen Eigenwert gehören genau k linear unabhängige Eigenvektoren
- Eigenvektoren zu verschiedenen Eigenwerten sind orthogonal

KAPITEL 5

Folgen

Eine Folge ist eine geordnete und numerierte Liste von Zahlen, die entweder in aufzählender Schreibweise oder durch eine Rechenvorschrift gegeben sein kann. Folgen können gegen einen Grenzwert konvergieren.

Die Theorie der Grenzwerte von Folgen ist eine wichtige Grundlage der Analysis, denn auf ihr beruhen die Berechnung von Grenzwerten von Funktionen, die Definition der Ableitung (Differentialquotient als Grenzwert einer Folge von Differenzenquotienten) und der Riemannsche Integralbegriff. Wichtige Folgen erhält man auch als Koeffizienten von Taylorreihen analytischer Funktionen.

5.1 Definition

Eine Folge ist eine Abbildung $f : \mathbb{N} \mapsto \mathbb{R}$ oder $\mapsto \mathbb{C}$
Notation: $a = (a_n)$, wobei a_n das n -te Folgeelement ist.

5.1.1 Beispiele

Es sei $a = (a_n)$ eine Folge mit

- $a_n = n \quad \forall n \in \mathbb{N}$
- $a_n = \frac{1}{n} \quad \forall n \in \mathbb{N}$
- $a_n = (-1)^n \quad \forall n \in \mathbb{N}$

Arithmetische Folgen

Folgen deren Elemente folgendermaßen berechnet werden

$\boxed{a_{n+1} - a_n = d}$ bzw. $\boxed{a_n = a_1 + (n - 1) \cdot d}$ bei beliebigem Anfangswert a_1 und Konstante d heißen arithmetische Folgen.

Bemerkung:

Jedes Folgeelement ist das arithmetische Mittel seiner beiden Nachbarn (ohne Beweis)

$$a_n = \frac{a_{n-1} + a_{n+1}}{2}$$

Geometrische Folgen

Folgen deren Elemente folgendermaßen berechnet werden

$$\boxed{a_n = a_1 \cdot q^{n-1}} \quad \text{bei beliebigem Anfangswert } a_1 \text{ und Konstante } q$$

heißen geometrische Folgen (oft auch für $n \in \mathbb{N}_0$ als $a_n = a_0 \cdot q^n$ bezeichnet).

Bemerkung:

Jedes Folgeelement ist das geometrische Mittel seiner beiden Nachbarn

$$a_n = \sqrt{a_{n-1} \cdot a_{n+1}}.$$

5.1.2 Definition monotone und beschränkte Folgen

Eine Folge (a_n) heißt

monoton wachsend	wenn $a_k \leq a_m$ für $k < m$.
monoton fallend	wenn $a_k \geq a_m$ für $k < m$.
streng monoton wachsend	wenn $a_k < a_m$ für $k < m$.
streng monoton fallend	wenn $a_k > a_m$ für $k < m$.
beschränkt	wenn es eine Zahl M gibt mit $ a_k \leq M \quad \forall k \in \mathbb{N}$.

5.2 Konvergenz

5.2.1 Definitionen

- Eine Zahl g heißt Grenzwert oder Limes der Zahlenfolge (a_n) , symbolisch

$$\lim_{n \rightarrow \infty} a_n = g$$

wenn es zu jedem $\varepsilon > 0$ eine natürliche Zahl n_0 so gibt, dass $|a_n - g| < \varepsilon$ für alle $n \geq n_0$ gilt.

- Eine Folge mit Grenzwert 0 heißt Nullfolge.
- Eine Folge heißt **konvergent**, wenn sie einen endlichen Grenzwert besitzt.

Andernfalls **divergent**

5.2.2 Rechenregeln für konvergente Folgen

Seien (a_n) und (b_n) zwei konvergente Folgen mit $\lim_{n \rightarrow \infty} a_n = a$ und $\lim_{n \rightarrow \infty} b_n = b$.

Dann gilt:

- $\lim_{n \rightarrow \infty} (a_n + b_n) = a + b$
- $\lim_{n \rightarrow \infty} (a_n b_n) = ab$
- $\lim_{n \rightarrow \infty} \left(\frac{a_n}{b_n} \right) = \frac{a}{b}$, falls $b \neq 0$ und $b_n \neq 0$ für alle $n \in \mathbb{N}$

5.2.3 Folgen der Form $p(n)/q(n)$

Zur Bestimmung des Grenzwertes von Folgen der Form $a_n = p(n)/q(n)$, wobei $p(n)$ und $q(n)$ Polynome der Variablen n sind, gibt es folgendes Vorgehen:

1. Man erweitert den Bruch mit $\frac{1}{n^k}$, wobei k der höchste Exponent ist, der in den Polynomen $p(n)$ und $q(n)$ auftritt

$$\lim_{n \rightarrow \infty} \frac{p(n)}{q(n)} = \lim_{n \rightarrow \infty} \frac{p(n) \cdot \frac{1}{n^k}}{q(n) \cdot \frac{1}{n^k}}$$

2. Nun stehen in Zähler und Nenner jeweils konvergente Folgen, deren Grenzwert bestimmt wird:

$$= \frac{\lim_{n \rightarrow \infty} p(n) \cdot \frac{1}{n^k}}{\lim_{n \rightarrow \infty} q(n) \cdot \frac{1}{n^k}} = \frac{p}{q}$$

3. Der Grenzwert der Folge (a_n) ist

$$\lim_{n \rightarrow \infty} a_n = \begin{cases} \infty & \text{für } q = 0 \\ \frac{p}{q} & \text{für } q \neq 0 \end{cases}$$

Beispiel:

$$\lim_{n \rightarrow \infty} \frac{n^2 + n + 1}{n - 1}$$

Alternative: Satz von l'Hospital:

$$\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = \lim_{n \rightarrow \infty} \frac{f'(n)}{g'(n)} \quad \text{mit } g'(n) \neq 0$$

Beispiel:

$$\lim_{n \rightarrow \infty} \frac{n^2 + n + 1}{n - 1}$$

5.2.4 Eulersche Zahl (ohne Beweis)

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = e = 2,71828182\dots$$

5.2.5 Konvergenzaussagen für monotone Folgen (ohne Beweis)

1. Eine monoton wachsende beschränkte Folge ist konvergent
2. Eine monoton fallende beschränkte Folge ist konvergent

KAPITEL 6

Funktionen

Unter einer **Funktion** versteht man eine Vorschrift f , die jedem Element x aus einer Definitionsmenge $\mathbb{D}$ ein Element y aus einer Wertemenge $\mathbb{W}$ zuordnet.

Notation: $y = f(x)$

Im folgenden betrachten wir nur eindimensionale reellwertige Funktionen, also Funktionen $f : \mathbb{R} \mapsto \mathbb{R}$

6.1 Darstellungsformen von Funktionen

Es gibt unterschiedliche Darstellungsformen einer Funktion:

1. **explizite Darstellung** $y = f(x)$

Hier kann der Funktionswert y für jedes x direkt berechnet werden.

Beispiel:

$$y = 2x + 3$$

2. **implizite Darstellung** $F(x,y) = 0$

Hier kann der Funktionswert nicht direkt aus einer Zuordnungsvorschrift berechnet werden, sondern nur indirekt über den Zusammenhang $F(x,y) = 0$.

Beispiele:

$$y - 2x - 3 = 0$$

$$x^2y = 1$$

Es ist nicht immer möglich eine implizite Darstellung in eine explizite umzuformen.

3. **Parameterdarstellung** $x = x(t), y = y(t)$

Bei der mathematischen Beschreibung eines Bewegungsablaufs wird oft die Lage eines Körpers durch seine Koordinaten (x,y) , die sich mit der Zeit t verändern, beschrieben. Eine Darstellung dieser Art ist eine Parameterdarstellung.

Beispiele:

$$x(t) = \cos(t) \text{ und } y(t) = \sin(t)$$

$$x(t) = t^2 - 1 \text{ und } y(t) = t(t^2 - 1)$$

6.2 Eigenschaften

6.2.1 Nullstellen

Eine Funktion $y = f(x)$ besitzt an der Stelle x_0 eine Nullstelle, wenn $f(x_0) = 0$.

6.2.2 (un-)gerade Funktionen

Eine Funktion f mit einem symmetrischen Definitionsbereich $\mathbb{D}$ heißt

gerade, wenn $f(-x) = f(x)$ für jedes $x \in \mathbb{D}$

ungerade, wenn $f(-x) = -f(x)$ für jedes $x \in \mathbb{D}$

Beispiel:

Die Funktion $y = \sin(x)$ ist ungerade. Die Funktion $y = \cos(x)$ ist gerade.

6.2.3 monotone Funktionen

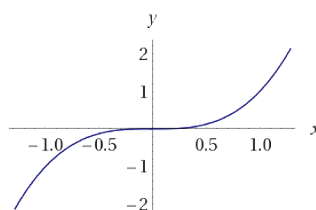
Eine Funktion f heißt

monoton wachsend	wenn für alle $a < b$ gilt:	$f(a) \leq f(b)$
streng monoton wachsend	wenn für alle $a < b$ gilt:	$f(a) < f(b)$
monoton fallend	wenn für alle $a < b$ gilt:	$f(a) \geq f(b)$
streng monoton fallend	wenn für alle $a < b$ gilt:	$f(a) > f(b)$
konstant	wenn für alle a, b gilt:	$f(a) = f(b)$

Beispiele:

$f(x) = x^3$ ist streng monoton wachsend

$f(x) = 1$ ist konstant



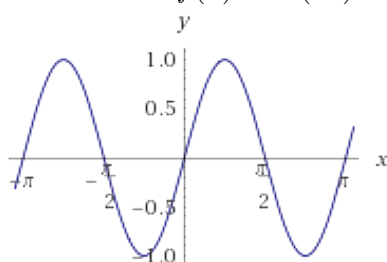
6.2.4 periodische Funktionen

Eine Funktion f heißt periodisch mit der Periode T , wenn mit jedem $x \in \mathbb{D}$ auch $x \pm T \in \mathbb{D}$ ist und es gilt:

$$f(x \pm T) = f(x)$$

Beispiel:

Die Funktion $f(x) = \sin(2x)$ ist periodisch mit Periode $T = \pi$



6.2.5 umkehrbare Funktion

Eine Funktion f heißt umkehrbar, wenn aus

$$x_1, x_2 \in \mathbb{D} \text{ mit } x_1 \neq x_2 \text{ stets } f(x_1) \neq f(x_2) \text{ folgt.}$$

Bestimmung der Umkehrfunktion

Ist die Funktion $y = f(x)$ umkehrbar, so bestimmt man die Umkehrfunktion in folgenden Schritten:

1. Vertauschen der Variablen x und y : $x = f(y)$
Dies ist die Umkehrfunktion in impliziter Schreibweise
2. Auflösen nach y (nicht immer möglich)
Dies ist die explizite Darstellung der Umkehrfunktion $f^{-1}(x)$

Beispiele:

1. $y = (x + 1)^2$ für $x \geq 0$

2. $y = x^2 - 2x + 3$ für $x \geq 2$

6.2.6 Grenzwert einer Funktion

Eine Funktion f sei in einer Umgebung der Stelle x_0 definiert. Gilt dann für jede im Definitionsbereich der Funktion liegende und gegen die Stelle x_0 konvergierende Zahlenfolge (x_n) mit $x_n \neq x_0$ stets

$$\lim_{n \rightarrow \infty} f(x_n) = g$$

so heißt g der **Grenzwert von f an der Stelle x_0** .

Notation: $\lim_{x \rightarrow x_0} f(x) = g$

Beispiele:

1. $\lim_{x \rightarrow 2} \left(\frac{x^2 - 2x}{x - 2} \right) =$

2. $\lim_{x \rightarrow 0} \left(\frac{1}{x} \right) =$

Besitzt eine Funktion f die Eigenschaft, dass die Folge ihrer Funktionswerte $(f(x_n))$ für jede wachsende Zahlenfolge $(x_n) \in \mathbb{D}$ gegen eine Zahl g strebt, so heißt g der

Grenzwert von f für $x \rightarrow \infty$.

$$\text{Notation: } \lim_{x \rightarrow \infty} f(x) = g$$

Analog gilt:

Besitzt eine Funktion f die Eigenschaft, dass die Folge ihrer Funktionswerte $(f(x_n))$ für jede fallende Zahlenfolge $(x_n) \in \mathbb{D}$ gegen eine Zahl g strebt, so heißt g der

Grenzwert von f für $x \rightarrow -\infty$.

$$\text{Notation: } \lim_{x \rightarrow -\infty} f(x) = g$$

Beispiele:

$$1. \lim_{x \rightarrow \infty} \left(\frac{2x-1}{x} \right) =$$

$$2. \lim_{x \rightarrow \pm\infty} \left(\frac{x^3}{x^2+1} \right) =$$

Rechenregeln für Grenzwerte

Falls die jeweiligen Grenzwerte existieren, gelten folgende Regeln:

$$1. \lim_{x \rightarrow x_0} (Cf(x)) = C \left(\lim_{x \rightarrow x_0} f(x) \right) \text{ für beliebige Konstante } C$$

$$2. \lim_{x \rightarrow x_0} (f(x) \pm g(x)) = \lim_{x \rightarrow x_0} f(x) \pm \lim_{x \rightarrow x_0} g(x)$$

$$3. \lim_{x \rightarrow x_0} (f(x)g(x)) = \lim_{x \rightarrow x_0} f(x) \lim_{x \rightarrow x_0} g(x)$$

$$4. \lim_{x \rightarrow x_0} \left(\frac{f(x)}{g(x)} \right) = \frac{\lim_{x \rightarrow x_0} f(x)}{\lim_{x \rightarrow x_0} g(x)} \text{ falls } \lim_{x \rightarrow x_0} g(x) \neq 0$$

Bemerkungen

- Diese Regeln gelten entsprechend für Grenzwerte vom Typ $x \rightarrow \pm\infty$
- Grenzwerte, die zu einem Ausdruck $\frac{0}{0}$ oder $\frac{\infty}{\infty}$ führen, werden in Mathematik 2 behandelt.

6.2.7 Stetige Funktionen

Eine in x_0 und in einer Umgebung von x_0 definierte Funktion f heißt an der Stelle x_0 **stetig**, wenn der Grenzwert der Funktion an dieser Stelle vorhanden ist und mit dem dortigen Funktionswert übereinstimmt:

$$\lim_{x \rightarrow x_0} f(x) = f(x_0)$$

Eine in x_0 und in einer Umgebung von x_0 definierte Funktion f heißt an der Stelle x_0 **unstetig**, wenn eine der folgenden Aussagen zutrifft:

1. Der Grenzwert von f an der Stelle x_0 ist zwar vorhanden, stimmt jedoch nicht mit dem Funktionswert überein, d.h.

$$\lim_{x \rightarrow x_0} f(x) \neq f(x_0)$$

2. Der Grenzwert von f an der Stelle x_0 existiert nicht.

Eine Funktion f heißt stetig, wenn sie für jedes $x_0 \in \mathbb{D}$ stetig ist.

6.3 Polynomfunktionen

6.3.1 Definition

Funktionen vom Typ

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$$

werden als ganzrationale Funktionen oder Polynomfunktionen bezeichnet.

Die Zahlen $a_0, a_1, \dots, a_n \in \mathbb{R}$ heißen Koeffizienten.

Der höchste Exponent n in der Funktionsgleichung mit $a_n \neq 0$ heißt Grad des Polynoms.

Bemerkungen

Polynomfunktionen besitzen viele besonders einfache und angenehme Eigenschaften:

Ein Polynom vom Grade n hat genau n (ev. komplexe) Nullstellen. Sie lassen sich problemlos differenzieren und integrieren. Aus diesem Grunde versucht man die bei technischen Problemen auftretenden Funktionen durch Polynome zu approximieren.

Beispiele:

1. $y = 4$
2. $y = 2x - 3$
3. $y = 2x^2 - 3x + 5$
4. $y = 4x^8 - x^5 + 3x$

6.3.2 Spezialfall: Polynom 1. Grades

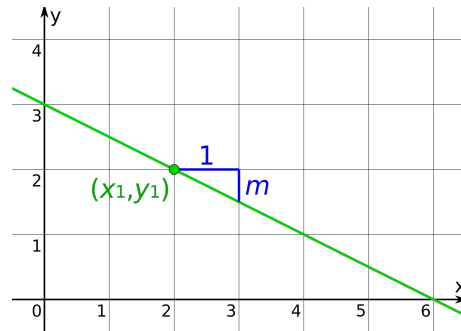
Polynome ersten Grades haben folgende Funktionsgleichung:

$$y = a_1 x + a_0 \quad \text{oder} \quad y = mx + b$$

Der Graph ist eine Gerade mit Steigung m und y -Achsenabschnitt b . Abhängig von der Problemstellung wird die Geradengleichung in folgenden Formen aufgestellt:

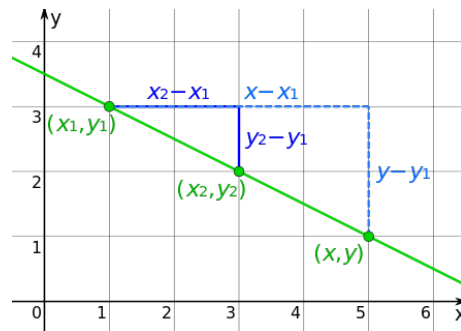
- **Punkt-Steigungs-Form:** Gegeben ein Punkt $(x_1|y_1)$ und die Steigung m der Geraden:

$$\frac{y - y_1}{x - x_1} = m$$



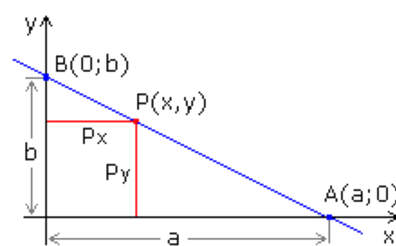
- **Zwei-Punkte-Form:** Gegeben zwei Punkte $(x_1|y_1)$ und $(x_2|y_2)$ der Geraden:

$$\frac{y - y_1}{x - x_1} = \frac{y_2 - y_1}{x_2 - x_1}$$



- **Achsenabschnitts-Form:** Gegeben die beiden Achsenabschnitte a der x -Achse und b der y -Achse der Geraden:

$$\frac{x}{a} + \frac{y}{b} = 1$$



$$\frac{x}{a} + \frac{y}{b} = 1 \quad a \neq 0 \quad \text{und} \quad b \neq 0$$

auflösen nach y

$$y = b \left(1 - \frac{x}{a} \right) \Rightarrow$$

$$y = f(x) = -\frac{b}{a}x + b$$

6.3.3 Spezialfall: Polynom 2. Grades

Polynome zweiten Grades haben folgende Funktionsgleichung:

$$y = a_2x^2 + a_1x + a_0 \quad \text{oder} \quad y = ax^2 + bx + c$$

Der Graph ist eine Parabel. Das Vorzeichen des Koeffizienten a entscheidet über die Öffnung der Parabel:

$a > 0$ Parabel nach oben geöffnet, Scheitelpunkt ist Tiefpunkt

$a < 0$ Parabel nach unten geöffnet, Scheitelpunkt ist Hochpunkt

Abhängig von der Problemstellung wird die Parabelgleichung in folgenden Formen aufgestellt:

- **Koordinatenform**

$$y = ax^2 + bx + c$$

- **Produktform**

Gegeben a und die Nullstellen x_1, x_2 der Parabel:

$$y = a(x - x_1)(x - x_2)$$

- **Scheitelpunktsform**

Gegeben a und die Koordinaten des Scheitelpunktes $S = (x_0|y_0)$ der Parabel:

$$y - y_0 = a(x - x_0)^2$$

6.3.4 Nullstellen

Anzahl der Nullstellen

Ein Polynom n -ten Grades besitzt genau n eventuell komplexe Nullstellen, also höchstens n reelle Nullstellen.

Mehrfach auftretende Nullstellen werden entsprechend oft mitgezählt.

Produktdarstellung

Sind die Nullstellen der Polynomfunktion n -ten Grades bekannt $x_1, x_2, \dots, x_n$, so lässt sich die Funktion auch in Form eines Produktes darstellen:

$$\begin{aligned} f(x) &= a_nx^n + a_{n-1}x^{n-1} + \dots + a_1x + a_0 \\ &= a_n(x - x_1)(x - x_2)\dots(x - x_n) \end{aligned}$$

Die n Faktoren $x - x_1, x - x_2, \dots, x - x_n$ werden als Linearfaktoren der Produktdarstellung bezeichnet.

Bemerkungen

1. Bei einer r -fachen Nullstelle tritt der zugehörige Linearfaktor r -fach auf.
2. Ist die Anzahl der reellen Nullstellen k kleiner als der Polynomgrad n , so besitzt die reelle Produktdarstellung folgende Form:

$$f(x) = a_n(x - x_1)(x - x_2)\dots(x - x_k)f^*(x)$$

wobei f^* ein Polynom vom Grade $n - k$ ohne reelle Nullstellen ist.

Nullstellenberechnung

Die Nullstellen einer Polynomfunktion f vom Grade $n > 2$ lassen sich schrittweise berechnen:

1. Im ersten Schritt wird durch Probieren versucht eine (reelle) Nullstelle x_1 als Faktor von a_0 zu bestimmen.
2. Hat man eine Nullstelle gefunden, so wird die Polynomfunktion durch den Linearfaktor $x - x_1$ dividiert. Das Restpolynom hat einen Grad $n - 1$.
3. Durch Wiederholung der Schritte 1 und 2 wird so lange verfahren, bis das Restpolynom vom Grad ≤ 2 ist, wofür es eine Berechnungsvorschrift für die Nullstellen gibt.

6.4 Gebrochen rationale Funktionen

6.4.1 Definition

Funktionen, die als Quotient zweier Polynomfunktionen $g(x)$ und $h(x)$ darstellbar sind, heißen gebrochen-rationale Funktionen:

$$y = \frac{g(x)}{h(x)} = \frac{a_mx^m + a_{m-1}x^{m-1} + \dots + a_1x + a_0}{b_nx^n + b_{n-1}x^{n-1} + \dots + b_1x + b_0}$$

Gebrochen rationale Funktionen sind für alle $x \in \mathbb{R}$, außer den Nennernullstellen, definiert.

Gebrochen rationale Funktionen heißen für

- $n > m$ echt gebrochen rationale Funktionen
- $n \leq m$ unecht gebrochen rationale Funktionen

Beispiele

1. $f(x) = \frac{x^3 - 1}{x + 1}$
2. $f(x) = \frac{x^2 + 2x + 3}{x^2 + x}$
3. $f(x) = \frac{x - 1}{x^2 + x + 1}$

6.4.2 Definitionsbereich, Nullstellen, Pole

1. Der Definitionsbereich einer gebrochen rationalen Funktion

$f = \frac{g}{h}$ besteht aus allen reellen Zahlen außer den Nennernullstellen, d.h.

$$\mathbb{D} = \mathbb{R} \setminus \{\text{Nennernullstellen}\}$$

2. Die **Nullstellen** einer gebrochen rationalen Funktion $f = \frac{g}{h}$ sind alle Zählernullstellen aus $\mathbb{D}$, d.h. $x_0 \in \mathbb{D}$ mit $g(x_0) = 0$.
3. Definitionslücken, in deren unmittelbarer Umgebung die Funktionswerte über alle Grenzen wachsen heißen **Pole**.

Vorgehensweise

Zur Bestimmung des Definitionsbereiches, der Nullstellen und der Pole einer gebrochen rationalen Funktion, wird folgendermassen vorgegangen:

1. Bestimmung aller Nullstellen des Nenners, ergibt den Definitionsbereich
2. Bestimmen aller Zählernullstellen im Definitionsbereich, ergibt die Nullstellen
3. Zerlegung von Zähler und Nenner in Linearfaktoren und Kürzen dieser Faktoren soweit möglich
4. Die Nennernullstellen der gekürzten gebrochen rationalen Funktion ergeben die Pole
5. Die Definitionslücken, welche keine Pole sind, sind hebbare Lücken der Funktion

6.4.3 Asymptoten

Um das Verhalten einer gebrochen rationalen Funktion f für große x -Werte, d.h. für $x \rightarrow \pm\infty$, zu bestimmen, wird die gebrochen rationale Funktion durch Polynomdivision in eine Summe aus Polynom p und echt gebrochen rationale Funktion r zerlegt.

$$f(x) = p(x) + r(x)$$

Da die echt gebrochen rationale Funktion r für große x gegen 0 konvergiert,

$$\lim_{x \rightarrow \pm\infty} r(x) = 0$$

nähert sich die gebrochen rationale Funktion für große x dem Polynom an. Man nennt $p(x)$ die Asymptote von f .

$$f(x) \approx p(x) \text{ für } x \rightarrow \pm\infty,$$

Bemerkung

An den Polstellen x_i spricht man ebenfalls von Asymptoten $x = x_i$

6.5 Potenzfunktionen

Potenzfunktionen sind vom Typ

$$f(x) = x^r \quad \text{mit } x > 0 \text{ und } r \in \mathbb{R}$$

$r = n \in \mathbb{N}$	$f(x) = x^n$ ist eine Polynomfunktion
$r = -n \in \mathbb{N}$	$f(x) = \frac{1}{x^n}$ ist eine gebrochen rationale Funktion
$r = \frac{1}{n}$ mit $n \in \mathbb{N}$	$f(x) = \sqrt[n]{x}$ mit $x \geq 0$ ist eine Wurzelfunktion
$r = \frac{m}{n}$ mit $n, m \in \mathbb{N}$	$f(x) = \sqrt[n]{x^m}$ mit $x \geq 0$ ist eine Potenzfunktion mit rationalem Exponenten
$r \in \mathbb{R}$	$f(x) = x^r = e^{\ln x^r} = e^{r \ln x}$ ist eine Potenzfunktion mit beliebigem reellen Exponenten

6.6 Trigonometrische Funktionen

6.6.1 Definition

Mit trigonometrischen Funktionen oder auch Winkelfunktionen bezeichnet man rechnerische Zusammenhänge zwischen Winkel und Seitenverhältnissen (ursprünglich in rechtwinkligen Dreiecken). Tabellen mit Verhältniswerten für bestimmte Winkel ermöglichen Berechnungen bei Vermessungsaufgaben, die Winkel und Seitenlängen in Dreiecken nutzen.

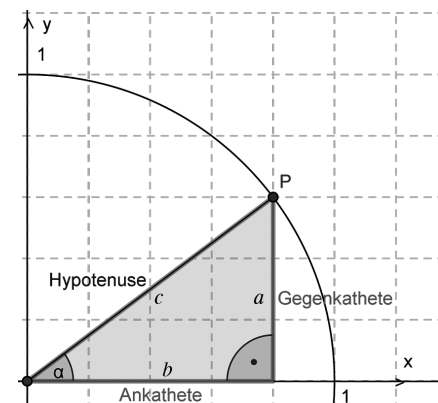
Die trigonometrischen Funktionen sind außerdem die grundlegenden Funktionen zur Beschreibung periodischer Vorgänge in den Naturwissenschaften.

Sie finden u.a. Anwendung bei

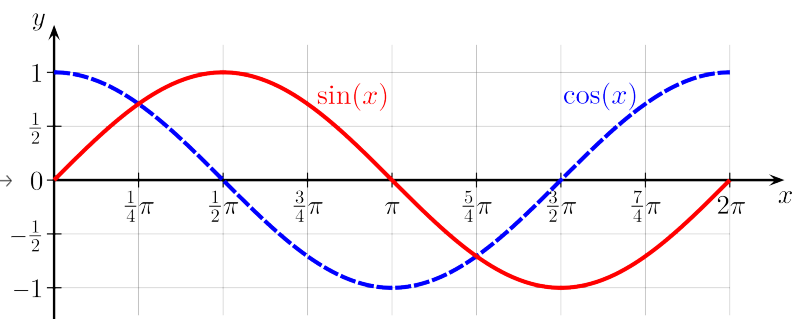
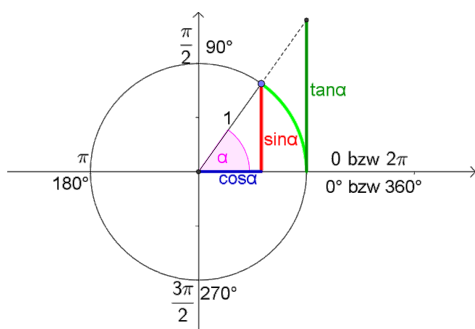
- mechanischen und elektromagnetischen Schwingungen
- gekoppelten Schwingungen
- Ausbreitung von Wellen

Die 4 trigonometrischen Funktionen **Sinus**, **Kosinus**, **Tangens**, **Kotangens** sind folgendermaßen im rechtwinkligen Dreieck definiert:

$\sin \alpha$	$=$	$\frac{\text{Gegenkathete}}{\text{Hypotenuse}}$	$=$	$\frac{a}{c}$
$\cos \alpha$	$=$	$\frac{\text{Ankathete}}{\text{Hypotenuse}}$	$=$	$\frac{b}{c}$
$\tan \alpha$	$=$	$\frac{\text{Gegenkathete}}{\text{Ankathete}}$	$=$	$\frac{a}{b} = \frac{a/c}{b/c} = \frac{\sin \alpha}{\cos \alpha}$
$\cot \alpha$	$=$	$\frac{\text{Ankathete}}{\text{Gegenkathete}}$	$=$	$\frac{b}{a} = \frac{b/c}{a/c} = \frac{\cos \alpha}{\sin \alpha} = \frac{1}{\tan \alpha}$



Für beliebige Winkel x werden die trigonometrischen Funktionen mit Hilfe des Einheitskreises definiert.



6.6.2 Zusammenhänge

Es gelten folgende nützliche Gleichungen:

$\sin(x + \pi) = -\sin(x)$ und $\cos(x + \pi) = -\cos(x)$
$\sin(x + \frac{\pi}{2}) = \cos(x)$ und $\cos(x + \frac{\pi}{2}) = -\sin(x)$
$\sin(-x) = -\sin(x)$ und $\cos(-x) = \cos(x)$
$\sin^2 x + \cos^2 x = 1$
$\sin(x \pm y) = \sin(x)\cos(y) \pm \sin(y)\cos(x)$
$\cos(x \pm y) = \cos(x)\cos(y) \mp \sin(x)\sin(y)$
$\sin(2x) = 2\sin(x)\cos(x)$
$\cos(2x) = \cos^2(x) - \sin^2(x)$
$\tan(x \pm y) = \frac{\tan(x) \pm \tan(y)}{1 \mp \tan(x)\tan(y)}$

6.6.3 Allgemeine Sinus- und Kosinusfunktion

Bei der Beschreibung von (mechanischen, elektromagnetischen) Schwingungsvorgängen benötigt man Sinus- und Kosinusfunktionen in der allgemeinen Form:

$$y = a \cdot \sin(bx + c)$$

$$y = a \cdot \cos(bx + c)$$

Die Parameter a, b, c mit $a > 0$ und $b > 0$, bewirken gegenüber den elementaren Sinus- und Kosinusfunktionen $y = \sin x$ bzw. $y = \cos x$ folgende Änderung:

$y = a \cdot \sin(bx + c)$	Periode:	$p = \frac{2\pi}{b}$
	1. Nullstelle	$x_0 = -\frac{c}{b}$
	Wertebereich	$-a \leq y \leq a$
$y = a \cdot \cos(bx + c)$	Periode:	$p = \frac{2\pi}{b}$
	1. Maximum:	$x_m = -\frac{c}{b}$
	Wertebereich:	$-a \leq y \leq a$

Anwendungsbeispiel: harmonische Schwingung eines Federpendels

Bei der Schwingung eines Federpendels kann die Auslenkung y als Sinusschwingung abhängig von der Zeit t , betrachtet werden.

$$y = A \sin(\omega t + \varphi)$$

Dabei bedeuten:

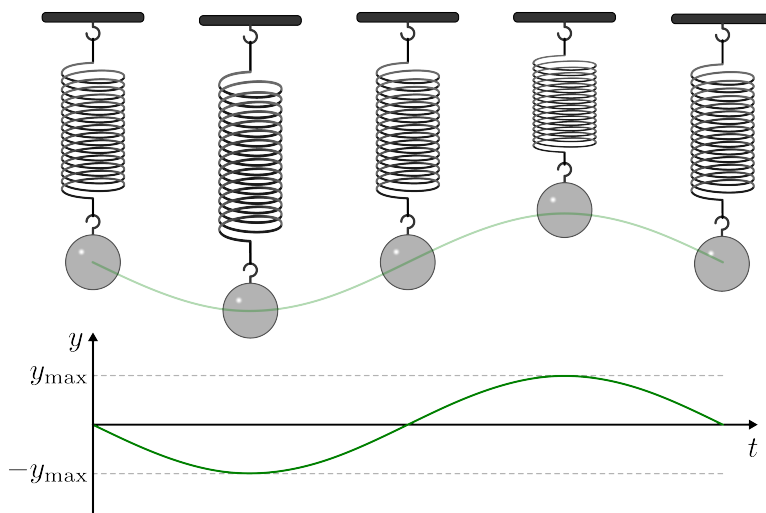
A Amplitude, d.h. maximale Auslenkung ($y_{\max}$)

ω Kreisfrequenz der Schwingung

φ Phase

$T = \frac{2\pi}{\omega}$ ist die Periodendauer (Schwingungsdauer)

$t_0 = \frac{-\varphi}{\omega}$ Phasenverschiebung



6.6.4 Darstellung der Sinusschwingung im Zeigerdiagramm

Eine Sinusschwingung vom Typ

$$y = A \sin(\omega t + \varphi)$$

lässt sich durch einen Zeiger mit

Länge A

Winkel φ

symbolisch darstellen.

Überlagerung gleichfrequenter Sinusschwingungen

Nach dem Superpositionsprinzip der Physik, entsteht bei der Überlagerung zweier gleichfrequenter Sinusschwingungen

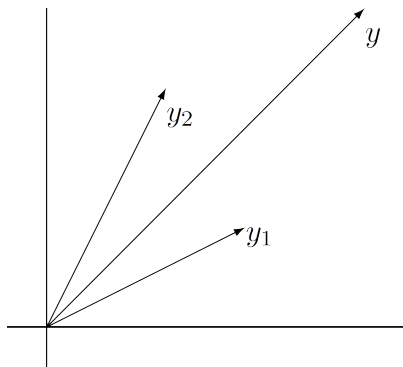
$$y_1 = A_1 \sin(\omega t + \varphi_1) \text{ und } y_2 = A_2 \sin(\omega t + \varphi_2)$$

eine resultierende Schwingung gleicher Frequenz

$$y = A \sin(\omega t + \varphi)$$

Die Amplitude A und der Phasenwinkel φ der resultierenden Schwingung lassen sich zeichnerisch im Zeigerdiagramm ermitteln.

Der Zeiger der resultierenden Schwingung ergibt sich durch vektorielle Addition der beiden anderen Zeiger



Ergebnis der Überlagerung

Die Überlagerung zweier gleichfrequenter Sinusschwingungen

$$y_1 = A_1 \sin(\omega t + \varphi_1) \text{ und } y_2 = A_2 \sin(\omega t + \varphi_2)$$

ergibt eine resultierende Schwingung gleicher Frequenz

$$y = A \sin(\omega t + \varphi)$$

mit

$$A = \sqrt{A_1^2 + A_2^2 + 2A_1A_2\cos(\varphi_2 - \varphi_1)}$$

und

$$\tan\varphi = \frac{A_1\sin\varphi_1 + A_2\sin\varphi_2}{A_1\cos\varphi_1 + A_2\cos\varphi_2}$$

6.7 Arkusfunktionen

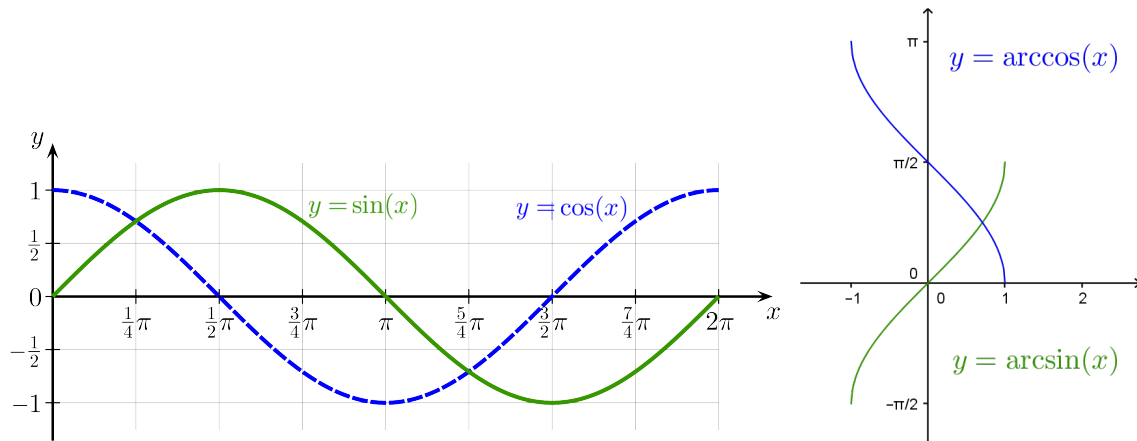
Die Arkusfunktionen sind die Umkehrfunktionen der trigonometrischen Funktionen. Grundsätzlich lassen sich die trigonometrischen Funktionen nicht umkehren, da sie periodisch sind. Beschränkt man sich jedoch auf gewisse Intervalle, in denen die Funktionen streng monoton verlaufen, so sind sie diesbezüglich umkehrbar.

Die **Umkehrfunktionen** werden als **Arkusfunktionen** bezeichnet. Ihre Funktionswerte sind im Bogenmaß dargestellte Winkel.

6.7.1 Arkussinus & Arkuscosinus

Die **Arkussinusfunktion** $y = \arcsin(x)$ ist die Umkehrfunktion der auf das Intervall $-\frac{\pi}{2} \leq x \leq \frac{\pi}{2}$ beschränkten Sinusfunktion $y = \sin(x)$.

Die **Arkuscosinusfunktion** $y = \arccos(x)$ ist die Umkehrfunktion der auf das Intervall $0 \leq x \leq \pi$ beschränkten Kosinusfunktion $y = \cos(x)$.

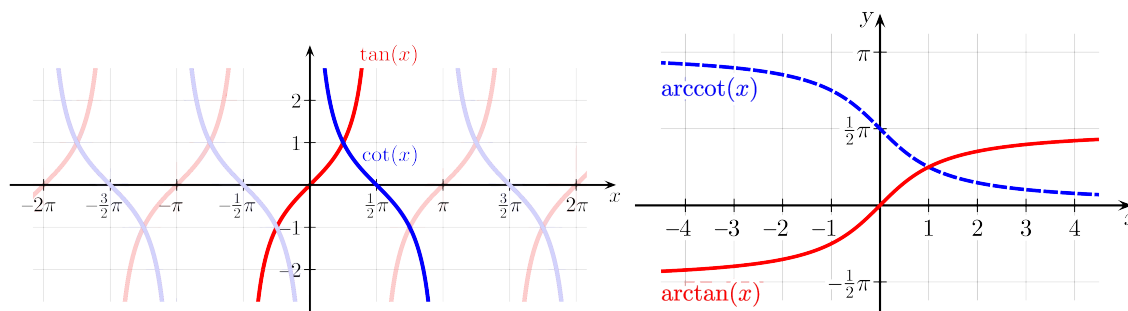


	$y = \sin(x)$	$y = \arcsin(x)$	$y = \cos(x)$	$y = \arccos(x)$
Definitionsbereich	$-\frac{\pi}{2} \leq x \leq \frac{\pi}{2}$	$-1 \leq x < 1$	$0 \leq x \leq \pi$	$-1 \leq x < 1$
Wertebereich	$-1 \leq y \leq 1$	$-\frac{\pi}{2} \leq y \leq \frac{\pi}{2}$	$-1 \leq y \leq 1$	$0 \leq y \leq \pi$
Nullstellen	$x_0 = 0$	$x_0 = 0$	$x_0 = \frac{\pi}{2}$	$x_0 = 1$
Monotonie	s. m. steigend	s. m. steigend	s. m. fallend	s. m. fallend

6.7.2 Arkustangens & Arkuscotangens

Die **Arkustangensfunktion** $y = \arctan(x)$ ist die Umkehrfunktion der auf das Intervall $-\frac{\pi}{2} < x < \frac{\pi}{2}$ beschränkten Tangensfunktion $y = \tan(x)$.

Die **Arkuscotangensfunktion** $y = \operatorname{arccot}(x)$ ist die Umkehrfunktion der auf das Intervall $0 < x < \pi$ beschränkten Kotangensfunktion $y = \cot(x)$.



	$y = \tan(x)$	$y = \arctan(x)$	$y = \cot(x)$	$y = \operatorname{arccot}(x)$
Definitionsbereich	$-\frac{\pi}{2} \leq x \leq \frac{\pi}{2}$	$-\infty \leq y \leq \infty$	$0 \leq x \leq \pi$	$-\infty \leq y \leq \infty$
Wertebereich	$-\infty \leq y \leq \infty$	$-\frac{\pi}{2} < y < \frac{\pi}{2}$	$-\infty \leq y \leq \infty$	$0 \leq y \leq \pi$
Nullstellen	$x_0 = 0$	$x_0 = 0$	$x_0 = \frac{\pi}{2}$	keine
Monotonie	s. m. steigend	s. m. steigend	s. m. fallend	s. m. fallend

6.7.3 Trigonometrische Gleichungen

Unter einer trigonometrischen Gleichung versteht man eine Gleichung, bei der die Unbekannte x in den Argumenten trigonometrischer Funktionen auftritt (z.B. $\sin(2x) = \frac{3}{2}\cos(x)$). Es gibt hierzu kein allgemeines Lösungsverfahren.

6.8 Exponentialfunktionen

Funktionen vom Typ

$$y = a^x \quad \text{mit } a > 0 \text{ und } a \neq 1$$

heißen **Exponentialfunktionen**.

Beispiele

- $y = 2^x$
- $y = \left(\frac{1}{3}\right)^x$
- $y = e^x$
- $y = e^{-x}$

Anwendungsbeispiele

- Abklingfunktion $y = a^{-\lambda t}$ mit $a > 0, \lambda > 0$
- Sättigungsfunktion $y = a(1 - e^{-\lambda t})$ mit $a > 0, \lambda > 0$
- aperiodischer Schwingungsvorgang

Der aperiodische Schwingungsvorgang tritt ein, wenn ein schwingungsfähiges System infolge zu großer Reibung zu keiner echten Schwingung mehr fähig ist, sondern sich asymptotisch der Gleichgewichtslage nähert.

$$y = 10e^{-2t} - 10e^{-4t} \text{ für } t > 0$$

- Gauß-Funktion

$$y = e^{-x^2} \text{ mit } x \in \mathbb{R}$$

Die Gauß-Funktion spielt eine wesentliche Rolle in der Wahrscheinlichkeitsrechnung.

6.9 Logarithmusfunktionen

Die **Logarithmusfunktion** $y = \log_a x$ ist die Umkehrfunktion der Exponentialfunktion

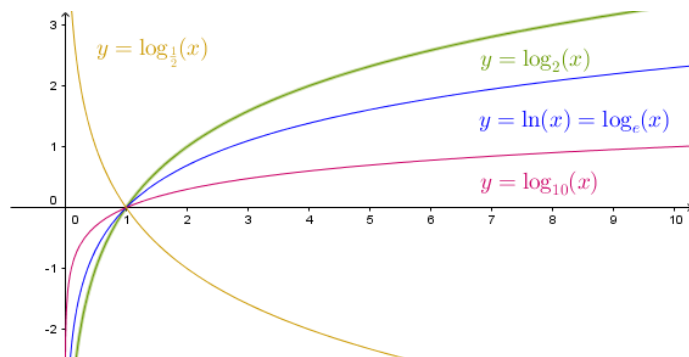
$$y = a^x \text{ mit } a > 0, a \neq 1.$$

spezielle Logarithmen

natürl. Logarithmus $\ln x = \log_e x$

Zehnerlogarithmus $\lg x = \log_{10} x$

Zweierlogarithmus $\lg x = \log_2 x$



Rechenregeln für Logarithmen

$$\log_a(uv) = \log_a(u) + \log_a(v)$$

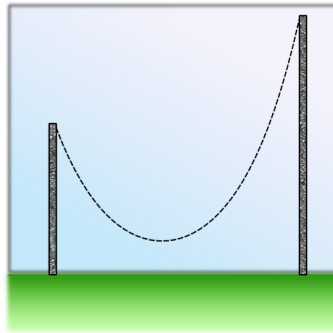
$$\log_a\left(\frac{u}{v}\right) = \log_a u - \log_a v$$

$$\log_a u^n = n \log_a u$$

$$\log_b r = \frac{\log_a r}{\log_a b} = \frac{1}{\log_a b} \cdot \log_a r$$

6.10 Hyperbelfunktionen

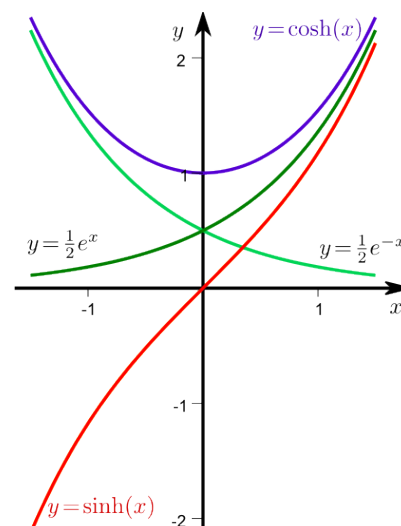
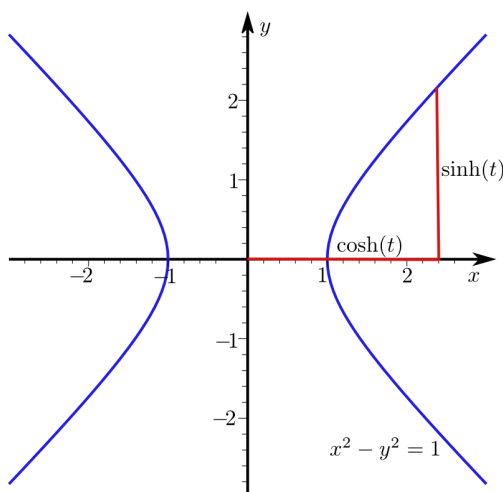
Die Hyperbelfunktionen sind spezielle Kombinationen aus den beiden e -Funktionen $y = e^x$ und $y = e^{-x}$, die in dieser Gestalt häufig in den Anwendungen vorkommen. So beispielsweise der Kosinus hyperbolicus im Brückenbau. Er beschreibt eine Kettenlinie, die immer dann entsteht, wenn eine ideale Kette in zwei Punkten aufgehängt wird und im Schwerfeld durchhängen kann.



Definition:

Sinus hyperbolicus:	$y = \sinh(x) = \frac{e^x - e^{-x}}{2}$
Kosinus hyperbolicus:	$y = \cosh(x) = \frac{e^x + e^{-x}}{2}$
Tangens hyperbolicus:	$y = \tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$
Kotangens hyperbolicus:	$y = \coth(x) = \frac{e^x + e^{-x}}{e^x - e^{-x}}$

Die Hyperbelfunktionen heißen deshalb so, da die Punkte $(\cosh(a), \sinh(a))$ auf der Hyperbel $x^2 - y^2 = 1$ liegen, ähnlich wie man auch Sinus und Cosinus Kreisfunktionen nennt, weil alle Punkte $(\cos(a), \sin(a))$ auf dem Einheitskreis liegen



Zusammenhänge zwischen den Hyperbelfunktionen

$\tanh(x)$	$= \frac{\sinh(x)}{\cosh(x)}$
$\coth(x)$	$= \frac{\cosh(x)}{\sinh(x)} = \frac{1}{\tanh(x)}$
$\sinh(x \pm y)$	$\sinh(x) \cdot \cosh(y) \pm \cosh(x) \cdot \sinh(y)$
$\cosh(x \pm y)$	$= \cosh(x) \cdot \cosh(y) \pm \sinh(x) \cdot \sinh(y)$
$\tanh(x \pm y)$	$= \frac{\tanh(x) \pm \tanh(y)}{1 \pm \tanh(x) \cdot \tanh(y)}$

Weitere Zusammenhänge

- $\cosh^2(x) - \sinh^2(x) = 1$
- $\sinh(2x) = 2\sinh(x) \cdot \cosh(x)$
- $\cosh(2x) = \sinh^2(x) + \cosh^2(x)$
- $e^x = \cosh(x) + \sinh(x)$
- $e^x = \cosh(x) \pm \sinh(x)$

6.11 Areafunktionen

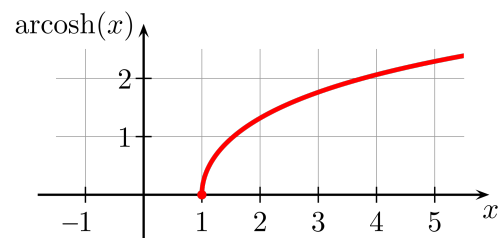
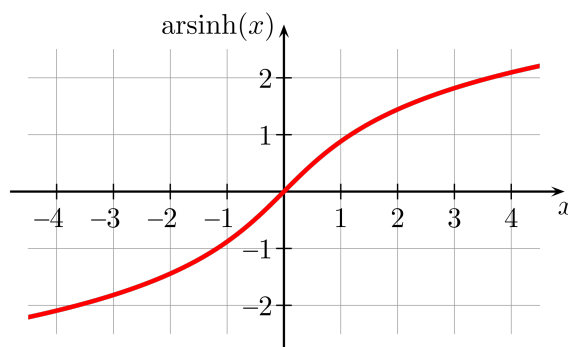
Die Umkehrfunktionen der Hyperbelfunktionen heißen Areafunktionen. Der Name Area rührt daher, da sich die Umkehrfunktion des Kosinus hyperbolicus als Fläche (Area) deuten lässt.

Die Hyperbelfunktionen $\sinh$, $\tanh$ und $\coth$ sind streng monoton und damit umkehrbar. Die Funktion $\cosh$ muss auf ein Teilintervall ($x \geq 0$) eingeschränkt werden, damit sie ebenfalls umkehrbar ist.

Definition:

Die Umkehrfunktionen der Hyperbelfunktionen $\sinh x$, $\cosh x$ eingeschränkt auf $x \geq 0$, $\tanh x$ und $\coth x$ sind:

Areasinus hyperbolicus:	$y = \operatorname{arsinh}(x)$
Areakosinus hyperbolicus:	$y = \operatorname{arcosh}(x)$
Areatangens hyperbolicus:	$y = \operatorname{artanh}(x)$
Areakotangens hyperbolicus:	$y = \operatorname{arcoth}(x)$



KAPITEL 7

Komplexe Zahlen

7.1 Einführung

Die komplexen Zahlen erweitern den Zahlenbereich der reellen Zahlen derart, dass auch Wurzeln negativer Zahlen berechnet werden können. Dies gelingt durch Einführung einer neuen Zahl i als Lösung der Gleichung $i^2 = -1$. Diese Zahl i wird imaginäre Einheit bezeichnet.

Beispiel:

1. $x^2 - x - 1 = 0$

pq -Formel: $x^2 + p \cdot x + q = 0$ mit Lösungen: $x_{1/2} = \frac{-p}{2} \pm \sqrt{\left(\frac{p}{2}\right)^2 - q}$

$$\Rightarrow x_{1/2} = \frac{1}{2} \pm \sqrt{\left(\frac{1}{2}\right)^2 + 1} = \frac{1}{2} \pm \frac{1}{2}\sqrt{5}$$

2. $x^2 + 2x + 3 = 0$

$$\Rightarrow x_{1/2} = -1 \pm (-1)^2 - 3 = -1 \pm \sqrt{-2}$$

Es existiert keine reelle Lösung, denn es gibt keine reelle Zahl $w \in \mathbb{R}$ mit $w^2 = -2$.

$$\text{Setzt man allerdings } w = \sqrt{2} \cdot i \Rightarrow w^2 = (\sqrt{2} \cdot i)^2 = 2 \cdot i^2 = -2$$

Es gibt also zwei komplexe Lösungen $x = -1 \pm w$

$$x_1 = -1 + \sqrt{2} \cdot i \text{ und } x_2 = -1 - \sqrt{2} \cdot i$$

Komplexe Zahlen werden meist in der Form

$$z = a + ib$$

dargestellt, wobei a und b reelle Zahlen sind und i die imaginäre Einheit ist. Auf die so dargestellten komplexen Zahlen lassen sich die üblichen Rechenregeln für reelle Zahlen anwenden, wobei stets i^2 durch -1 ersetzt werden kann.

In der Elektrotechnik wird als Symbol statt i ein j benutzt, um Verwechslungen mit der Stromstärke zu vermeiden.

Für die Menge der komplexen Zahlen wird das Symbol $\mathbb{C}$ verwendet.

Der so konstruierte Zahlenbereich der komplexen Zahlen hat eine Reihe vorteilhafter Eigenschaften, die sich in vielen Bereichen der Natur- und Ingenieurwissenschaften als äußerst nützlich erwiesen haben:

- Einer der Gründe für diese positiven Eigenschaften ist die algebraische Abgeschlossenheit der komplexen Zahlen. Dies bedeutet, dass jede algebraische Gleichung n -ter Ordnung

$$a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0 = 0$$

über den komplexen Zahlen genau n Lösungen besitzt, was für reelle Zahlen nicht gilt. Diese Eigenschaft ist Inhalt des Fundamentalsatzes der Algebra.

- Ferner ist jede einmal komplex differenzierbare Funktion von selbst beliebig oft differenzierbar, anders als in der Mathematik der reellen Zahlen.
- Ein weiterer Grund ist ein Zusammenhang zwischen den trigonometrischen Funktionen $\sin$ und $\cos$ mit der Exponentialfunktion, der über die komplexen Zahlen hergestellt werden kann.
- Die Integraltransformationen Fourier-Transformation, Laplace-Transformation und z -Transformation, die z.B. in der Regelungstechnik Anwendung finden sind Transformationen im komplexen Raum
- Schliesslich ermöglichen die komplexen Zahlen eine vereinfachte Beschreibung von Phasenverschiebungen in der Elektrotechnik

7.2 Definitionen

7.2.1 Zahlen

natürliche Zahlen $\mathbb{N} = \{1, 2, 3, \dots\}$

natürliche Zahlen mit Null $\mathbb{N}_0 = \{0, 1, 2, 3, \dots\}$

ganze Zahlen $\mathbb{Z} = \{\dots - 2, -1, 0, 1, 2, \dots\}$

rationale Zahlen $\mathbb{Q} = \left\{ \frac{n}{m} \mid n, m \in \mathbb{Z} \right\}$
 $= \{ \text{endliche und periodische Dezimalbrüche} \}$

reelle Zahlen $\mathbb{R} = \{ \text{endliche und unendliche Dezimalbrüche} \}$

komplexe Zahlen $\mathbb{C} = \{a + ib \mid a, b \in \mathbb{R}\}$

Komplexe Zahlen werden typischerweise durch z dargestellt.

7.2.2 Komplexe Zahlen

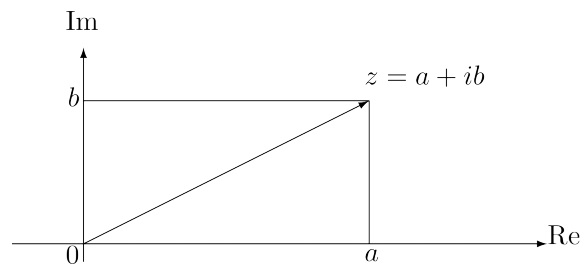
Komplexe Zahlenebene:

Die Menge der reellen Zahlen lässt sich durch Punkte auf einer Zahlengeraden veranschaulichen. Die Menge der komplexen Zahlen lässt sich als Punkte in einer Ebene darstellen. Diese Ebene wird durch 2 Achsen aufgespannt:

Die reelle Achse Re und die imaginäre Achse Im.

Die Teilmenge der reellen Zahlen liegt auf der waagrechte Achse Re, die Teilmenge der imaginären Zahlen, d.h. Zahlen ohne realen Anteil liegen auf der senkrechten Achse Im.

Eine komplexe Zahl besitzt dann die horizontale Koordinate a und die vertikale Koordinate b .



Imaginäre Einheit:

Die spezielle komplexe Zahl mit Abstand 1 vom Nullpunkt auf der imaginären Achse wird imaginäre Einheit i genannt: und es wird festgelegt:

$$i^2 = -1$$

Mit Hilfe der imaginären Einheit lässt sich jede komplexe Zahl z darstellen durch:

$$z = a + ib \text{ mit } a, b \in \mathbb{R}$$

Real- und Imaginärteil:

Ist $z = a + ib \in \mathbb{C}$, so heißt:

$a = \operatorname{Re}(z)$ Realteil von z

$b = \operatorname{Im}(z)$ Imaginärteil von z

Konjugiert komplexe Zahl

Für $z = a + ib \in \mathbb{C}$ ist

$$\bar{z} = a - ib$$

die konjugiert komplexe Zahl.

Betrag von z

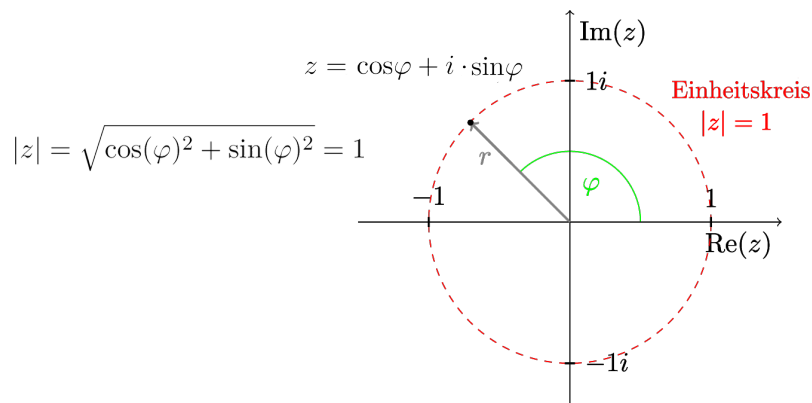
Der Betrag einer komplexen Zahl z ist definiert durch

$$|z| = \sqrt{a^2 + b^2} = \sqrt{z\bar{z}}$$

und entspricht ihrem Abstand in der komplexen Zahlenebene vom Nullpunkt. Ist die Zahl z eine reelle Zahl, also ist $b = 0$, so ist wie gewohnt $|z| = \sqrt{a^2} = |a|$.

Beispiel:

$\{z \in \mathbb{C} \mid |z| = 1\}$ entspricht dem Einheitskreis um den Ursprung in der komplexen Zahlenebene.



Polarform einer komplexen Zahl

Statt komplexe Zahlen in kartesischen Koordinaten zu beschreiben, kann man auch polare Koordinaten verwenden.

In der Mathematik versteht man unter einem **Polarkoordinatensystem** ein zweidimensionales Koordinatensystem, in dem jeder Punkt auf einer Ebene durch einen Winkel und einen Abstand definiert werden kann.

Das Polarkoordinatensystem ist hilfreich, wenn sich das Verhältnis zwischen 2 Punkten leichter durch Winkel und Abstände beschreiben lässt, als durch kartesische Koordinaten.

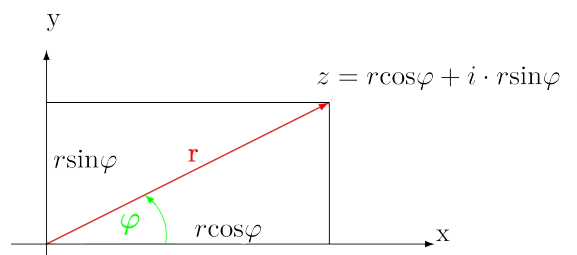
Die komplexe Zahl $z = x + iy$ wird in Polarform durch

r , den Abstand zum Ursprung in der komplexen Ebene und

φ , den Winkel zur reellen Achse

angegeben. Es ist dann:

$$z = r\cos\varphi + i \cdot r\sin\varphi$$



Üblicherweise nennt man r hier den Betrag von z und den Winkel φ das Argument (oder auch die Phase) von z .

Umrechnung kartesische Koordinaten - Polarkoordinaten

	gegeben:	
1.	$z = x + iy$	$r = \sqrt{x^2 + y^2}$ und $\varphi = \begin{cases} \arccos(\frac{x}{r}) & \text{für } y \geq 0 \\ -\arccos(\frac{x}{r}) & \text{für } y < 0 \end{cases}$
2.	r und φ	$z = r(\cos \varphi + i \sin \varphi)$

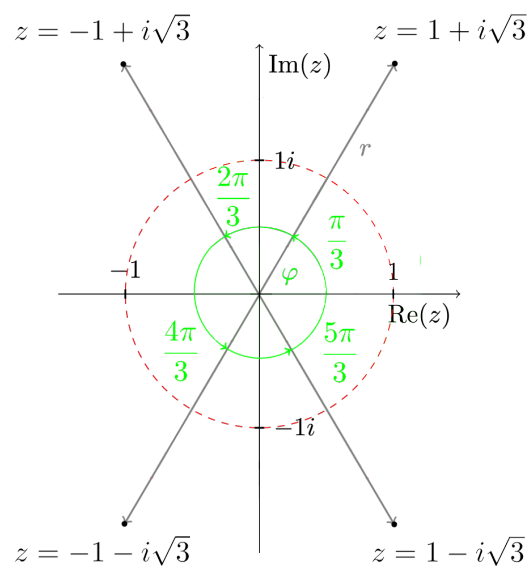
Beispiele:

$$1. \ z = 1 + i\sqrt{3} \Rightarrow r = \sqrt{1^2 + (\sqrt{3})^2} = 2 \text{ und } \varphi = \arccos(\frac{1}{2}) =$$

$$2. \ z = -1 + i\sqrt{3} \Rightarrow r = \sqrt{1 + 3} = 2 \text{ und } \varphi = \arccos(-\frac{1}{2}) =$$

$$3. \ z = -1 - i\sqrt{3} \Rightarrow r = \sqrt{1 + 3} = 2 \text{ und } \varphi = -\arccos(-\frac{1}{2}) =$$

$$4. \ z = 1 - i\sqrt{3} \Rightarrow r = \sqrt{1 + 3} = 2 \text{ und } \varphi = -\arccos(\frac{1}{2}) =$$



7.3 Grundrechenoperationen

Für die komplexen Zahlen sind die Grundrechenarten definiert.

Es seien im Folgenden $z_1 = x_1 + i \cdot y_1$ und $z_2 = x_2 + i \cdot y_2 \in \mathbb{C}$ und $k \in \mathbb{R}$

7.3.1 Addition

$$\begin{aligned} z_1 + z_2 &= x_1 + iy_1 + x_2 + iy_2 \\ &= x_1 + x_2 + i(y_1 + y_2) \end{aligned}$$

Beispiel:

$$(2 - i) + (-1 + 2i) =$$

7.3.2 Multiplikation mit einem Skalar

$$\begin{aligned} kz_1 &= k(x_1 + iy_1) \\ &= kx_1 + i \cdot k \cdot y_1 \end{aligned}$$

Beispiel:

$$-2(1 + i) =$$

7.3.3 Multiplikation

$$\begin{aligned} z_1 z_2 &= (x_1 + iy_1)(x_2 + iy_2) \\ &= x_1 x_2 + iy_1 x_2 + iy_2 x_1 + i^2 y_1 y_2 \\ &= (x_1 x_2 - y_1 y_2) + i(y_1 x_2 + y_2 x_1) \end{aligned}$$

Beispiel:

$$(2 + i)(3 - 2i) = 6 - i - 2i^2 =$$

7.3.4 Division

Für $z_2 \neq 0$ ist:

$$\begin{aligned} \frac{z_1}{z_2} &= \frac{x_1 + iy_1}{x_2 + iy_2} \\ &= \frac{(x_1 + iy_1)(x_2 + iy_2)}{(x_2 + iy_2)(x_2 + iy_2)} \\ &= \frac{x_1 x_2 - y_1 y_2 + i(y_1 x_2 + y_2 x_1)}{x_2^2 + y_2^2} \\ &= \frac{x_1 x_2 - y_1 y_2}{|z_2|^2} + i \cdot \frac{y_1 x_2 + y_2 x_1}{|z_2|^2} \end{aligned}$$

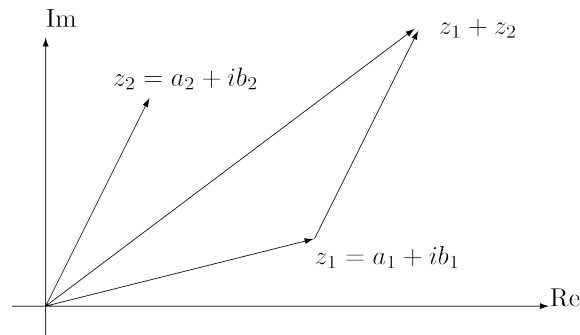
Beispiel:

$$\frac{2 + i}{1 - 2i} =$$

7.3.5 Graphische Darstellung der Grundrechenoperationen

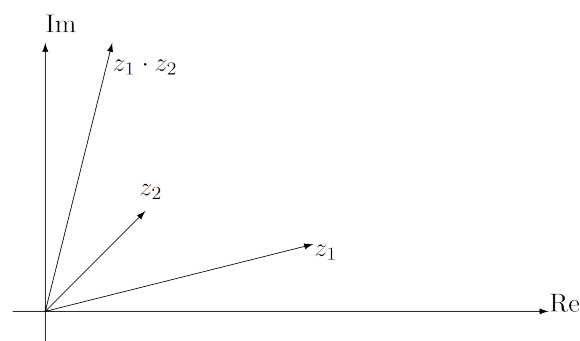
Addition

Die Addition komplexer Zahlen entspricht der Vektoraddition.



Multiplikation

Die Multiplikation zweier komplexer Zahlen in der komplexen Ebene entspricht einer **Drehstreckung**, d.h. die Winkel werden addiert und die Beträge multipliziert. Dies wird nach Einführung der e -Funktion weiter unten klarer werden.



Division

Bei der Division zweier komplexer Zahlen in der komplexen Ebene werden die Winkel subtrahiert und die Beträge dividiert.

7.3.6 Potenzieren und Wurzelziehen

Potenzieren:

Die n -te Potenz einer komplexen Zahl $z = r\cos\varphi + i\sin\varphi$ berechnet sich zu:

$$z^n = r^n(\cos n\varphi + i\sin n\varphi)$$

oder für die algebraische Form $z = a + ib$ zu

$$z^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} (ib)^k$$

Dabei ist $\binom{n}{k} = \frac{n!}{k!(n-k)!}$ (der Binomialkoeffizient).

Beispiele:

1. $i^3 =$
2. $i^4 =$
3. $i^5 =$
4. $i^{12} =$
5. $i^{33} =$

Wurzelziehen:

Als die n -ten Wurzeln einer komplexen Zahl $z \in \mathbb{C}$ bezeichnet man die Lösungen der Gleichung $z^n = a$.

Die n -ten Wurzeln $\sqrt[n]{z}$ einer komplexen Zahl $z = r\cos\varphi + i\sin\varphi$ berechnen sich wie folgt:

$$\sqrt[n]{r} \left(\cos \frac{\varphi + k2\pi}{n} + i \sin \frac{\varphi + k2\pi}{n} \right) \quad \text{für } k = 0, 1, 2, \dots, n-1$$

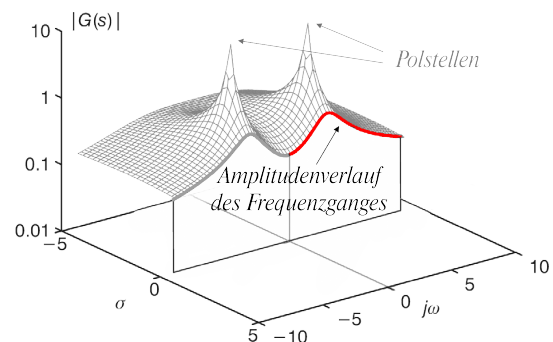
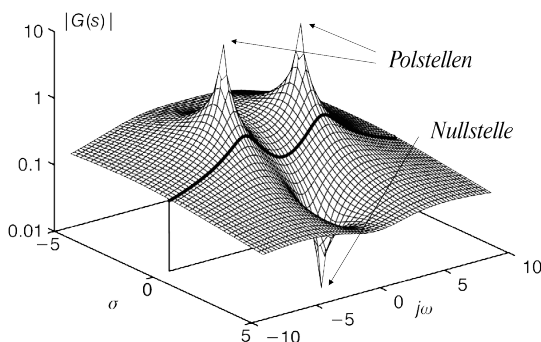
7.4 Funktionen komplexer Zahlen

Eine komplexe Funktion ordnet einer komplexen Zahl eine weitere komplexe Zahl zu, also:

$$f : z \in \mathbb{C} \rightarrow f(z) \in \mathbb{C}$$

In der Regelungstechnik sehr verbreitet der Betrag (Frequenzgang) einer komplexen Zahl (Übertragungsfunktion): $|G_s| = |\sigma + j\omega|$:

$$f : G_s \in \mathbb{C} \rightarrow f(|G_s|) \in \mathbb{R}$$



Darstellung des Betrages (Amplitude) einer komplexen Übertragungsfunktion G_s im Frequenzbereich. Technisch relevant ist nur der Abschnitt für $\sigma = 0$ sowie für Frequenzen $\omega \geq 0$ und spiegelt den Amplitudenverlauf dieses speziellen Übertragungssystems wider.

7.4.1 Potenzfunktionen

$$f(z) = z^2 = (x + iy)^2 = (x^2 - y^2) + i(2xy)$$

$$f(z) = z^3 = (x + iy)^3 = x^3 + 3x^2iy + 3xi^2y^2 + i^3y^3 = (x^3 - 3xy^2) + i(3x^2y - 3y^3)$$

$$f(z) = 2z^2 + z - 1 = 2(x^2 - y^2) + i(4xy) + x + iy - 1 = (2x^2 - 2y^2 + x - 1) + i(4xy + y)$$

$$f(z) = \frac{z^3 - 1}{z - 1} = \frac{(z^3 - z + 1)(x^2 - y^2 - 2ixy - 2)}{((x^2 - y^2 - 2) + i2xy)(x^2 - y^2 - 2 - 2ixy)}$$

7.4.2 Sinus-, Cosinus- und e-Funktion

Die folgenden Funktionen sind durch die Taylor-Reihe definiert.

- $f(z) = \sin z$ für $z \in \mathbb{C}$
- $f(z) = \cos z$ für $z \in \mathbb{C}$
- $f(z) = e^z$ für $z \in \mathbb{C}$

7.4.3 Komplexe e-Funktion

Die komplexe e-Funktion lässt sich mit Hilfe der Eulerschen Formel (Beweis folgt später mit Hilfe der Taylor-Reihen) leicht veranschaulichen.

Eulersche Formel:

$$e^{ib} = \cos(b) + i \cdot \sin(b) \quad \text{für } b \in \mathbb{R}$$

Die Zahlen e^{ib} liegen also alle wegen

$$|e^{ib}| = |\cos(b) + i \cdot \sin(b)| = \sqrt{\cos^2(b) + \sin^2(b)} = 1$$

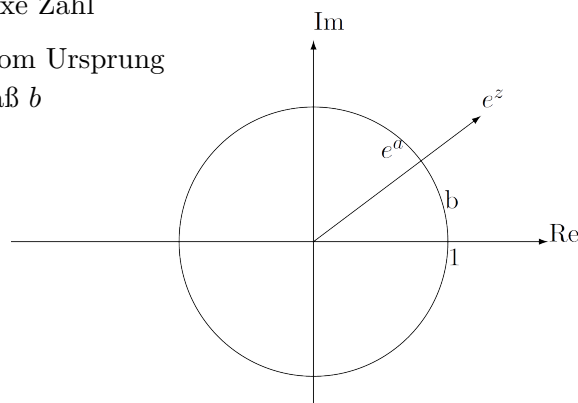
auf dem komplexen Einheitskreis. Die reelle Zahl b gibt das **Bogenmaß** des Punktes auf dem Einheitskreis an

Das Bild einer beliebigen komplexen Zahl $z = a + ib$ unter der e-Funktion ist

$$e^z = e^{a+ib} = e^a \cdot e^{ib} = e^a(\cos(b) + i\sin(b))$$

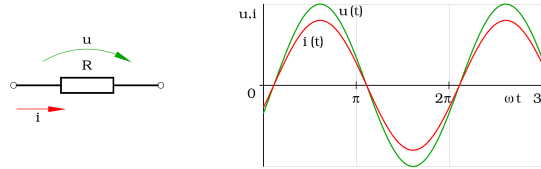
e^z ist also eine komplexe Zahl

mit Abstand e^a vom Ursprung
und mit Bogenmaß b



7.4.4 Anwendungsbeispiel aus der Elektrotechnik

An einem ohmschen Widerstand R wird die von der Zeit abhängige Wechselspannung $U = U_0 \sin(\omega t)$ angelegt (ω ist die Kreisfrequenz).

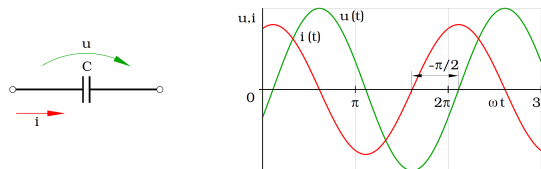


Aufgrund des Ohmschen Gesetztes ergibt sich für die Stromstärke

$$I = \frac{U}{R} = \frac{U_0}{R} \sin(\omega t) = I_0 \sin(\omega t)$$

Spannung und Stromstärke sind in Phase.

Befindet sich an der Stelle des ohmschen Widerstandes R ein Kondensator mit der Kapazität C , so wird der Stromfluss durch die entgegengesetzte Aufladung des Kondensators zur Spannungsquelle $U = U_0 \sin(\omega t)$ begrenzt.



Zum Zeitpunkt der Maximalspannung U_0 ist die Spannung am Kondensator konstant, es fließt in den Kondensator keine Ladung. Fällt danach die Spannung ab, so wird der Kondensator durch einen Strom in umgekehrter Richtung entladen. Die Stromstärke erreicht ihr Maximum, wenn sich die Spannung am schnellsten ändert, also beim Wechsel des Vorzeichens.

Dies bedeutet, dass die Stromstärke der Spannung um $\frac{\pi}{2}$ voraus eilt ($\varphi = -\frac{\pi}{2}$).

Die Stromstärke lässt sich also darstellen als:

$$I = I_0 \sin(\omega t - \frac{\pi}{2}) = I_0 \cos(\omega t)$$

Aus der Physik ist bekannt, dass die maximale Stromstärke

$$I_0 = \omega \cdot C \cdot U_0$$

beträgt. Dies lässt sich durch komplexe Zahlen beschreiben:

Man interpretiert den Kondensator als komplexen Widerstand bzw. als **Blindwiderstand**

$$R_C = \frac{1}{j\omega C} = -\frac{j}{\omega C} \quad \text{mit } C \text{ Kapazität und } \omega \text{ Kreisfrequenz } (\omega = 2\pi f)$$

Die komplexe Wechselstromrechnung wird in der Elektrotechnik angewendet, um Stromstärke und Spannung in einem Netzwerk bei sinusförmiger Wechselspannung zu bestimmen. Das Verhältnis der komplexen Spannung zur komplexen Stromstärke ist eine komplexe Konstante. Dies ist die Aussage des ohmschen Gesetzes für komplexe Größen. Die Konstante wird als komplexer Widerstand bezeichnet.

Der komplexe Widerstand einer Spule im Stromkreis ist

$$R_L = j\omega L \quad \text{mit } L \text{ Induktivität und } \omega \text{ Kreisfrequenz } (\omega = 2\pi f)$$

$\varphi = \frac{\pi}{2} \rightarrow$ **Bei Induktivitäten, Ströme sich verspäten.**

KAPITEL 8

Anhang

8.1 Einheitskreis

