

Herbstsemester 2025

Mathematik I für Naturwissenschaften

Christine Zehrt-Liebendörfer

Inhaltsverzeichnis

1	Grundlagen	1
1.1	Zahlen, Mengen und Symbole	1
1.2	Funktionen und Abbildungen	4
1.3	Elementare algebraische Funktionen	10
1.4	Trigonometrische Funktionen	16
1.5	Exponential- und Logarithmusfunktionen	22
2	Komplexe Zahlen	25
2.1	Definitionen und Rechenregeln	25
2.2	Algebraische Gleichungen	27
2.3	Polarkoordinaten und exponentielle Darstellung	29
2.4	Potenzen und Wurzeln	32
3	Grenzwerte und stetige Funktionen	35
3.1	Folgen	35
3.2	Reihen	41
3.3	Grenzwerte bei Funktionen	44
3.4	Stetige Funktionen	47
4	Differentiation	50
4.1	Die Ableitung einer Funktion	50
4.2	Extremalstellen	56
4.3	Die Regeln von Bernoulli-de l'Hôpital	62
4.4	Lineare Approximation	63
4.5	Taylorpolynome und Taylorreihen	65
4.6	Näherungsverfahren zur Lösung von Gleichungen	69
5	Integration	77
5.1	Das unbestimmte Integral	77
5.2	Das bestimmte Integral	78
5.3	Integrationstechniken	83
5.4	Uneigentliche Integrale	90
6	Differentialgleichungen	93
6.1	Separierbare Differentialgleichungen	93
6.2	Lineare Differentialgleichungen erster Ordnung	101
6.3	Lineare Differentialgleichungen zweiter Ordnung	105
6.4	Systeme von linearen Differentialgleichungen	110
7	Lineare Gleichungssysteme	112
7.1	Vektoren in der Ebene und im Raum	112
7.2	Der n -dimensionale Raum	116
7.3	Lineare Gleichungssysteme und Matrizen	117
7.4	Der Gauß-Algorithmus	120
7.5	Lösbarkeitskriterien	126
7.6	Struktur der Lösungsmenge	127

8	Rechnen mit Matrizen	129
8.1	Matrixoperationen und ihre Eigenschaften	129
8.2	Invertierbare Matrizen	133
8.3	Potenzen einer Matrix	137
8.4	Determinanten	139
8.5	Zwei weitere Lösungsmethoden für lineare Gleichungssysteme	143
9	Vektorräume	145
9.1	Definition und Beispiele	145
9.2	Linearkombinationen	147
9.3	Basis und Dimension	150
9.4	Orthogonale Vektoren	156
9.5	Näherungslösungen für unlösbare lineare Gleichungssysteme	159
10	Fourierreihen	162
10.1	Fourierpolynome	162
10.2	Fourierreihen	165
10.3	Eine Orthonormalbasis bestehend aus Funktionen	166
11	Boolesche Algebra	169
11.1	Grundlegende Operationen und Gesetze	169
11.2	Boolesche Funktionen und ihre Normalformen	170
11.3	Logische Schaltungen	173

16.09.2025

Christine Zehrt-Liebendörfer

Departement Mathematik und Informatik

Spiegelgasse 1, 4051 Basel

dmi.unibas.ch/de/personen/christine-zehrt

1 Grundlagen

1.1 Zahlen, Mengen und Symbole

In diesem ersten Abschnitt werden die gebräuchlichsten Bezeichnungen und Symbole definiert.

Zahlenmengen

Die Menge $\mathbb{N}$ der *natürlichen* Zahlen ist gegeben durch

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

die Menge $\mathbb{Z}$ der *ganzen* Zahlen ist gegeben durch

$$\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, 3, \dots\},$$

und die Menge $\mathbb{Q}$ der *rationalen* Zahlen (Brüche) ist gegeben durch

$$\mathbb{Q} = \left\{ \frac{p}{q} \mid p \text{ in } \mathbb{Z}, q \text{ in } \mathbb{N} \right\}.$$

In $\mathbb{Q}$ sind alle Operationen (Addition, Subtraktion, Multiplikation, Division) in eindeutiger Weise durchführbar, ausser der Division durch 0. Jede rationale Zahl lässt sich als (endlicher oder periodisch unendlicher) Dezimalbruch darstellen, zum Beispiel

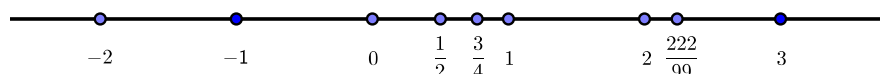
$$\begin{aligned} \frac{1}{4} &= 1 : 4 = 0,25 && \text{als endlicher Dezimalbruch und} \\ \frac{1}{3} &= 1 : 3 = 0, \overline{3} \approx 0,33333333 && \text{als periodischer Dezimalbruch.} \end{aligned}$$

Die Darstellung einer rationalen Zahl als Bruch ist nicht eindeutig, zum Beispiel ist

$$0,5 = \frac{1}{2} = \frac{2}{4} = \frac{3}{6} = \dots,$$

die Darstellung als Dezimalbruch ist hingegen eindeutig, wenn man “Neunerenden” nicht zulässt, wie zum Beispiel $0,4\overline{9} = 0,4999\dots = 0,5$. Es gilt $0,\overline{9} = 1$, wie wir später noch in den Übungen zeigen werden.

Die rationalen Zahlen kann man als Punkte auf der Zahlengeraden veranschaulichen:



Es bleiben jedoch Lücken auf der Zahlengeraden. Durch Hinzunahme aller Zahlen in diesen Lücken erhält man die Menge $\mathbb{R}$ der *reellen* Zahlen. Das heisst, die Menge der reellen Zahlen ist die Gesamtheit aller Zahlen auf der Zahlengeraden. Man sagt, dass die reellen Zahlen die rationalen Zahlen vervollständigen (mathematisch kann dies durch Grenzwertbetrachtungen exakt beschrieben werden). Eine reelle Zahl, die nicht rational ist, nennt man *irrational*. Zum Beispiel sind die Zahlen $\sqrt{2}$, $\sqrt{2025}$, $\sqrt[3]{1893}$, $\log(99)$, π , e , usw. irrational.

Wir werden später noch eine weitere Zahlenmenge kennenlernen, nämlich die Menge $\mathbb{C}$ der komplexen Zahlen.

Runden von Dezimalzahlen

Rechenmaschinen rechnen aufgrund ihres endlichen Speichers mit gerundeten Zahlen und auch bei Messungen entstehen Ungenauigkeiten.

Je nach Situation runden wir eine Zahl auf k Stellen nach dem Komma (d.h. auf k *Nachkommastellen*). Dabei werden Ziffern ≥ 5 aufgerundet und Ziffern < 5 abgerundet.

Beispiel

Die Zahl $\pi = 3,1415926\dots$ soll auf $k = 2$, $k = 3$ und auf $k = 4$ Nachkommastellen gerundet werden.

Ist umgekehrt eine Dezimalzahl gegeben, so muss man davon ausgehen, dass sie gerundet wurde. Ist eine Zahl a mit k Nachkommastellen gegeben, so weicht der exakte Wert von a höchstens um $0,5 \cdot 10^{-k}$ vom angegebenen Wert ab (gemäss den Rundungsregeln).

Beispiele

1. Die Angabe $a = 7,24$ bedeutet, dass $7,235 \leq a < 7,245$.
2. Die Angabe $a = 7,240$ bedeutet, dass $7,2395 \leq a < 7,2405$.

Bei Rechenoperationen muss man sich überlegen, wieviele Nachkommastellen bei der Angabe des Endergebnisses sinnvoll sind.

Beispiele

Sinnvolle Endergebnisse sind:

$$1,234 + 1,100 = 2,334 \quad \text{aber} \quad 1,234 + 1,1 = 2,3$$

Der Rundungsfehler kann sich bei Rechenoperationen vergrössern. Bei mehreren Rechenschritten ist es daher wichtig, dass erst das Endergebnis gerundet wird, und nicht schon jedes Zwischenergebnis!

Beispiel

Wir berechnen die Zahl

$$a = \frac{9}{7} \cdot \left[2 + 10\,000 \cdot \left(\frac{1}{9} - 0,111 \right) \right]$$

- exakt:
- durch Ausmultiplizieren aller Klammern und Runden auf jeweils drei Nachkommastellen nach jedem Zwischenschritt:

$$\begin{aligned} a &= \frac{18}{7} + \frac{90\,000}{7} \cdot \frac{1}{9} - \frac{9}{7} \cdot 1110 \\ &\approx 2,571 + 90\,000 \cdot 0,016 - 1,286 \cdot 1110 \approx 2,571 + 1440 - 1427,46 \approx 15,11 \end{aligned}$$

Dieses Beispiel zeigt, dass auch die Reihenfolge der Rechenschritte eine Rolle spielt. Bei einer Summe mit sehr kleinen und sehr grossen Summanden sollte man die Summanden nach aufsteigender Grösse aufsummieren.

Mengentheoretische und logische Symbole

Die wichtigsten mengentheoretischen Symbole sind die Folgenden:

$x \in \mathbb{R}$	Dies bedeutet, dass x eine reelle Zahl ist, d.h. x gehört zur Menge $\mathbb{R}$ (bzw. x ist Element der Menge $\mathbb{R}$).
$\frac{1}{2} \notin \mathbb{N}$	Dies heisst, dass $\frac{1}{2}$ keine natürliche Zahl ist, d.h. dass $\frac{1}{2}$ nicht zur Menge $\mathbb{N}$ der natürlichen Zahlen gehört.
$M \subset \mathbb{R}$	Die Menge M ist eine Teilmenge der reellen Zahlen (eventuell auch ganz $\mathbb{R}$), zum Beispiel gilt $\mathbb{Q} \subset \mathbb{R}$.
$\mathbb{R} \setminus M$	= Menge der reellen Zahlen, die nicht in M liegen = $\{x \in \mathbb{R} \mid x \notin M\}$, zum Beispiel ist $\mathbb{R} \setminus \mathbb{Q}$ die Menge der irrationalen Zahlen.
$\mathbb{R}^2$	= Menge der geordneten Zahlenpaare = $\{(x, y) \mid x, y \in \mathbb{R}\}$.
$\mathbb{N} \times \mathbb{Z}$	= Menge der geordneten Zahlenpaare = $\{(x, y) \mid x \in \mathbb{N}, y \in \mathbb{Z}\}$.
$[a, b]$	= abgeschlossenes Intervall zwischen a und $b = \{x \in \mathbb{R} \mid a \leq x \leq b\}$.
(a, b)	= offenes Intervall zwischen a und $b = \{x \in \mathbb{R} \mid a < x < b\}$.

Eine logische Folgerung beschreibt man mit einem Pfeil. Man unterscheidet:

$A \implies B$	Aus der Aussage A folgt Aussage B , d.h. wenn A wahr ist, so ist auch B wahr. (Dies ist gleichbedeutend mit: Wenn B falsch ist, dann ist auch A falsch.)
$A \iff B$	Aus A folgt B und aus B folgt A , d.h. A ist genau dann wahr, wenn B wahr ist. Die Aussagen A und B sind gleichbedeutend, man sagt auch <i>äquivalent</i> .

Beispiele

1. n ist teilbar durch 6 $\implies$ n ist teilbar durch 2.

Die Umkehrung " $\impliedby$ " gilt hier nicht. Die Aussage ist gleichbedeutend mit:

n ist nicht teilbar durch 2 $\implies$ n ist nicht teilbar durch 6.

2. $x^2 - 16 = 0 \iff (x - 4)(x + 4) = 0 \iff x = 4 \text{ oder } x = -4$

Summenzeichen

Für Summen mit mehreren Summanden ist es praktisch, das folgende Summenzeichen zu benutzen. Man schreibt

$$a_1 + a_2 + \cdots + a_N = \sum_{n=1}^N a_n,$$

wobei anstelle von n auch eine andere Variable gewählt werden kann. Die Summe muss dabei nicht mit dem Index 1 beginnen. Es gilt zum Beispiel

$$a_3 + a_4 + \cdots + a_N = \sum_{n=3}^N a_n \quad \text{und} \quad a_{-2} + a_{-1} + a_0 + a_1 = \sum_{k=-2}^1 a_k.$$

Für $m > N$ gilt die Vereinbarung

$$\sum_{n=m}^N a_n = 0.$$

Beispiele

$$1. \quad \sum_{n=1}^{10} \frac{1}{n} =$$

$$2. \quad \sum_{n=-1}^2 (3n + 5) =$$

$$3. \quad 5 + 7 + 9 + 11 + 13 + 15 + 17 =$$

$$4. \quad -1 + \frac{1}{4} - \frac{1}{9} + \frac{1}{16} - \cdots + \cdots - \frac{1}{81} + \frac{1}{100} =$$

Wir werden das Summenzeichen vor allem im nächsten Kapitel und dann in der Statistik im nächsten Semester gebrauchen. Die folgenden *Rechenregeln* sind dabei sehr nützlich:

$$1. \quad \sum_{n=1}^N (a_n + b_n) = \sum_{n=1}^N a_n + \sum_{n=1}^N b_n$$

$$2. \quad \sum_{n=1}^N (c \cdot a_n) = c \sum_{n=1}^N a_n$$

$$3. \quad \sum_{n=1}^N a_n = \sum_{n=1}^k a_n + \sum_{n=k+1}^N a_n \quad \text{für } k \in \{1, \dots, N\}$$

1.2 Funktionen und Abbildungen

Der Funktionsbegriff ist einer der zentralsten Begriffe in der Mathematik. Er spielt auch bei allen Anwendungen in den Naturwissenschaften eine fundamentale Rolle. Funktionen treten überall da auf, wo Zusammenhänge zwischen zwei (oder mehreren) Grössen bestehen.

Beispiel

Zustandsgleichung des idealen Gases. Der Druck eines Gases in einem geschlossenen Gefäss hängt von der Temperatur und dem Volumen ab. Mit Hilfe geeigneter Skalen wird der Druck durch eine Grösse $p \in \mathbb{R}$, die (absolute) Temperatur durch $T \in \mathbb{R}$ und das Volumen durch $V \in \mathbb{R}$ angegeben, und jedem bestimmten Wert von V und T entspricht ein Wert von p . Für 1 mol (1 mol = $6,022 \cdot 10^{23}$ Moleküle bzw. Atome (*Avogadrosche Zahl*)) eines idealen Gases gilt die Beziehung

$$p = \frac{RT}{V}, \quad \text{wobei } R \text{ eine (absolute) Konstante ist.}$$

Halten wir das Volumen fest, so gilt $p(T) = aT$ mit $a = \frac{R}{V}$. Halten wir jedoch die Temperatur fest, so gilt $p(V) = b\frac{1}{V}$ mit $b = RT$.

Wie im Beispiel betrachten wir zunächst vor allem Zuordnungen zwischen Grössen, welche Werte in den reellen Zahlen annehmen. Im nächsten Semester werden wir jedoch auch allgemeinere Zuordnungen (sogenannte Abbildungen) studieren.

Definition Seien X und Y zwei Mengen. Eine Vorschrift f , die jedem Element x aus X ($x \in X$) genau ein Element y aus Y ($y \in Y$) zuordnet, heisst *Abbildung*. Wir schreiben

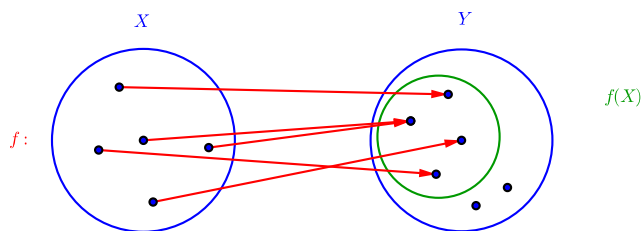
$$\begin{aligned} f &: X \longrightarrow Y \\ x &\mapsto f(x) \end{aligned}$$

Die Menge X heisst *Definitionsbereich* und die Menge Y *Zielbereich* von f .

Eine Abbildung ordnet also jedem $x \in X$ eindeutig ein $y \in Y$ zu. Es können jedoch zwei verschiedene $x, x' \in X$ demselben $y \in Y$ zugeordnet werden. Ausserdem muss nicht zu jedem $y \in Y$ ein $x \in X$ existieren, dem es zugeordnet wird. Man nennt deshalb die Teilmenge

$$f(X) = \{ f(x) \mid x \in X \} = \{ y \in Y \mid \text{es gibt ein } x \in X \text{ mit } y = f(x) \}$$

von Y die *Bildmenge* (oder den *Wertebereich*) von f . Entsprechend ist $f(x)$ das *Bild* von f an der Stelle x .



Ist der Zielbereich Y eine Zahlenmenge, dann nennen wir die Abbildung eine *Funktion* (in vielen Büchern werden Abbildungen und Funktionen synonym verwendet). Funktionen mit $Y = \mathbb{R}$ und $X \subset \mathbb{R}$ heissen *reelle Funktionen*. Für X schreiben wir in diesem Fall D und meinen damit jeweils die grösstmögliche Teilmenge von $\mathbb{R}$, auf der die Funktionsvorschrift f definiert ist.

Beispiele

1. $f : D \longrightarrow \mathbb{R}$, $f(x) = x^2$ definiert eine reelle Funktion.

2. Sei X die Menge aller Studierenden der Pharmazie im ersten Semester und Y die Menge $\{\text{blau, grün, braun}\}$. Dann ist die Zuordnung $f : X \longrightarrow Y$, die jedem Studierenden seine Augenfarbe zuordnet, eine Abbildung. Die Bildmenge ist ganz Y , falls es je mindestens eine*n Studierende*n mit der Augenfarbe blau, grün, bzw. braun gibt.

3. Sei $X = \mathbb{R}^2$. Dann definiert $f : X \longrightarrow \mathbb{R}$ mit $f(x, y) = \sqrt{x^2 + y^2}$ eine Funktion in zwei Variablen. Für einen Punkt $P = (x, y)$ in $\mathbb{R}^2$ gibt $f(x, y)$ den Abstand von P zum Ursprung an. Der Definitionsbereich ist $(D =) X = \mathbb{R}^2$ und die Bildmenge ist gleich $f(X) = \{z \in \mathbb{R} \mid z \geq 0\}$. Funktionen in zwei und mehr Variablen werden wir im nächsten Semester untersuchen.

4. Funktionen können auch abschnittsweise definiert werden. Zum Beispiel ist $f : D \rightarrow \mathbb{R}$ definiert durch

$$f(x) = \begin{cases} x - 3 & \text{falls } x \geq 0 \\ \frac{1}{1-x} & \text{falls } x < 0 \end{cases}$$

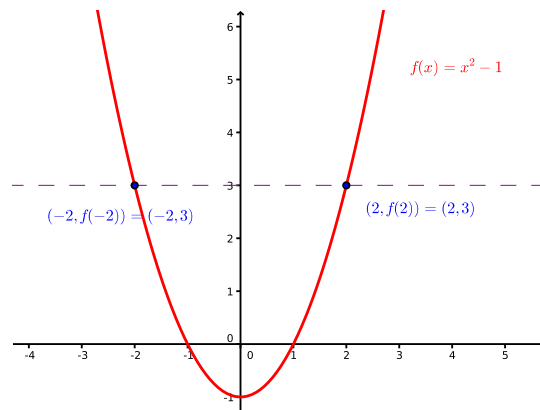
auch eine (reelle) Funktion.

Reelle Funktionen können mit Hilfe ihres Graphen dargestellt werden. Der *Graph* einer reellen Funktion $f : D \rightarrow \mathbb{R}$ ist definiert als die Menge aller Punkte der Ebene mit den Koordinaten $(x, f(x))$ für $x \in D$:

$$\text{Graph}(f) = \{ (x, f(x)) \mid x \in D \} \subset D \times \mathbb{R}$$

Der Graph ist also eine Kurve in $\mathbb{R}^2$ mit der Eigenschaft, dass über jedem $x \in D$ auf der x -Achse genau ein Punkt des Graphen von f liegt, nämlich der Punkt $(x, f(x))$. Umgekehrt können jedoch mehreren $x_1, x_2, \dots \in D$ der gleiche Wert y zugeordnet sein, das heisst, die Parallele zur x -Achse in der Höhe y kann den Graphen in mehreren Punkten schneiden.

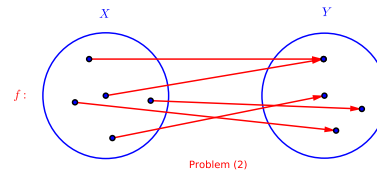
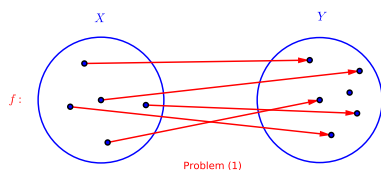
Beispiel



Umkehrbarkeit

Sei $f : D \rightarrow \mathbb{R}$ eine (reelle) Funktion. Dann ist die Umkehrfunktion von f (falls es sie überhaupt gibt) diejenige Vorschrift, die einer Zahl $y \in \mathbb{R}$ diejenige Zahl $x \in D$ zuordnet, für die $f(x) = y$ gilt. Die Klammerbemerkung “falls es sie überhaupt gibt” kommt daher, da zwei Probleme auftreten können:

- (1) Eventuell gibt es nicht zu jedem $y \in \mathbb{R}$ ein $x \in D$ mit $f(x) = y$.
- (2) Eventuell gibt es zu einem $y \in \mathbb{R}$ nicht nur ein $x \in D$ mit $f(x) = y$, sondern ein weiteres $\tilde{x} \neq x$ in D mit $f(\tilde{x}) = y$.



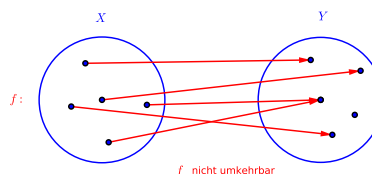
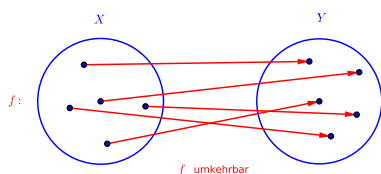
Beispiel

Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$. Dann gibt es zum Beispiel zu $y = -1 \in \mathbb{R}$ kein $x \in \mathbb{R} = D$ mit $f(x) = x^2 = -1$. Und zu $y = 4$ gibt es $x = 2, \tilde{x} = -2$ mit $f(x) = 4 = f(\tilde{x})$ und $x \neq \tilde{x}$.

Wir müssen also zuerst definieren, welche Funktionen eine Umkehrfunktion haben. Das machen wir gleich allgemein für Abbildungen und nicht nur für Funktionen.

Definition Eine Abbildung $f : X \rightarrow Y$ heisst *umkehrbar*, wenn es zu jedem $y \in Y$ genau ein $x \in X$ gibt mit $y = f(x)$.

Die Abbildung $f^{-1} : Y \rightarrow X$, die jedem $y \in Y$ das $x \in X$ mit $y = f(x)$ zuordnet, heisst *Umkehrabbildung* von f (bzw. *Umkehrfunktion*, falls f eine Funktion ist).



Anschaulich gesprochen gilt:

f umkehrbar $\iff$ auf jedes Element in Y zeigt *genau ein* Pfeil

Achtung

$$f^{-1}(x) \neq \frac{1}{f(x)} = (f(x))^{-1} \quad !!!$$

Beispiele

1. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = 3x$ ist umkehrbar mit Umkehrfunktion $f^{-1} : \mathbb{R} \rightarrow \mathbb{R}$, $f^{-1}(x) = \frac{1}{3}x$.

2. $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^2$ ist (gemäss obigen Bemerkungen) nicht umkehrbar, da für f beide Probleme (1) und (2) auftreten.

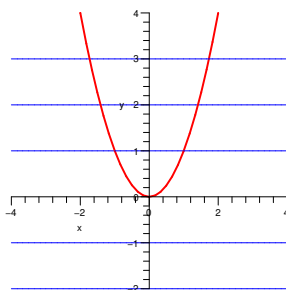
3. $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$, $f(x) = \frac{5}{x}$.

Es folgt, dass $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R} \setminus \{0\}$ (mit verkleinertem Zielbereich) umkehrbar ist. Für die Umkehrfunktion gilt $f^{-1} = f$.

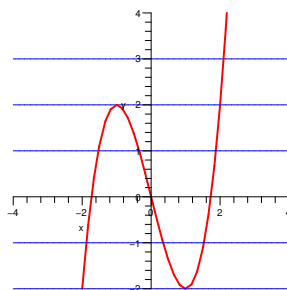
4. Seien $X = \{\text{Studierende der Pharmazie im 1. Semester}\}$, $Y = \{\text{blau, grün, braun}\}$ und $f : X \rightarrow Y$ ordnet jedem*r Studierenden seine*ihre Augenfarbe zu. Nun ist f keine reelle Funktion. Wird hier jedes $y \in Y$ von *genau einem* $x \in X$ getroffen?

Dass f in diesem Beispiel nicht umkehrbar ist, erkennt man auch daran, dass die Menge X aus (viel) mehr Elementen besteht als die Menge Y .

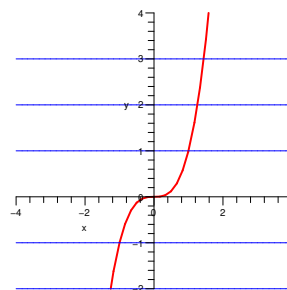
Zur Untersuchung von speziell reellen Funktionen $f : D \rightarrow \mathbb{R}$ gibt es weitere Methoden. Das Problem (1) hängt mit der Bildmenge $f(D)$ von f zusammen. Denn zu $y \in \mathbb{R}$ gibt es genau dann ein $x \in D$ mit $y = f(x)$, wenn $y \in f(D)$. Wir können also den Zielbereich $\mathbb{R}$ auf $f(D)$ verkleinern (wie oben im 3. Beispiel), dann tritt Problem (1) nicht auf. Ob Problem (2) auf $f : D \rightarrow f(D)$ zutrifft oder nicht, kann einfach am Graphen von f abgelesen werden. Wenn jede Parallele zur x -Achse den Graphen in höchstens einem Punkt schneidet, ist f umkehrbar:



$f(x) = x^2$
nicht umkehrbar



$f(x) = x^3 - 3x$
nicht umkehrbar



$f(x) = x^3$
umkehrbar

Dies hängt mit der Monotonie der Funktion zusammen. Eine Funktion $f : D \rightarrow \mathbb{R}$ heisst

<i>monoton wachsend</i>	falls $f(x_1) \leq f(x_2)$
<i>streng monoton wachsend</i>	falls $f(x_1) < f(x_2)$
<i>monoton fallend</i>	falls $f(x_1) \geq f(x_2)$
<i>streng monoton fallend</i>	falls $f(x_1) > f(x_2)$

für alle $x_1, x_2 \in D$ mit $x_1 < x_2$ gilt.

Satz 1.1 Ist eine reelle Funktion $f : D \rightarrow f(D)$ streng monoton wachsend oder streng monoton fallend, dann ist sie umkehrbar.

Ist eine Funktion $f : D \rightarrow f(D)$ nicht streng monoton wachsend oder fallend, so kann man meistens eine Teilmenge $\tilde{D}$ von D finden, auf welcher f streng monoton wachsend oder fallend ist. Damit ist $f : \tilde{D} \rightarrow f(\tilde{D})$ eine umkehrbare Funktion.

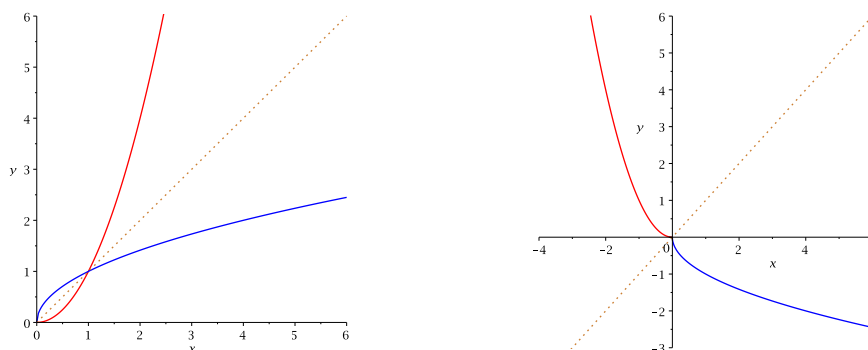
Beispiel

Sei $\mathbb{R}_{\geq 0} = \{x \in \mathbb{R} \mid x \geq 0\}$ und $f : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ mit $f(x) = x^2$. Dann gilt $\mathbb{R}_{\geq 0} = f(\mathbb{R})$. Auf dem ganzen Definitionsbereich $D = \mathbb{R}$ ist f allerdings nicht umkehrbar, wie man dem Graphen ansieht.

Verkleinert man hingegen den Definitionsbereich $D = \mathbb{R}$ auf die Teilmenge $\mathbb{R}_{\geq 0}$, dann ist $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0} = f(\mathbb{R}_{\geq 0})$ streng monoton wachsend, also umkehrbar. Die Umkehrfunktion ist gegeben durch $f^{-1} : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$, $f^{-1}(x) = \sqrt{x}$.

Auf der Teilmenge $\mathbb{R}_{\leq 0} = \{x \in \mathbb{R} \mid x \leq 0\}$ ist f ebenfalls umkehrbar, weil f auf $\mathbb{R}_{\leq 0}$ streng monoton fallend ist. Die Umkehrfunktion ist $f^{-1} : \mathbb{R}_{\leq 0} \rightarrow \mathbb{R}_{\leq 0}$, $f^{-1}(x) = -\sqrt{x}$.

In den folgenden Abbildungen sind die Graphen von $f(x) = x^2$, links auf $\mathbb{R}_{\geq 0}$, rechts auf $\mathbb{R}_{\leq 0}$, mit zugehöriger Umkehrfunktion $f^{-1}(x) = \pm\sqrt{x}$ zu sehen. Man erkennt, dass die Graphen von f und f^{-1} spiegelbildlich zur Geraden $y = x$ liegen.



Um die Funktionsgleichung einer Umkehrfunktion zu bestimmen, geht man vor wie im 1. und 3. Beispiel auf Seite 7. Man setzt $y = f(x)$, löst die Gleichung nach x auf und vertauscht anschliessend die Variablen x und y . Dann ist f^{-1} gegeben durch $f^{-1}(x) = y$. Dieses Vertauschen der Variablen x und y entspricht genau dem Spiegeln des Graphen von f an der Geraden $y = x$.

Wie der Name Umkehrfunktion oder allgemeiner Umkehrabbildung andeutet, ist f^{-1} die "Umkehrung" der Abbildung f . Dies ist im folgenden Sinn gemeint.

Satz 1.2 Für eine umkehrbare Abbildung $f : X \rightarrow Y$ mit Umkehrabbildung f^{-1} gilt

$$f^{-1}(f(x)) = x \quad \text{und} \quad f(f^{-1}(y)) = y$$

für alle $x \in X$ und $y \in Y$.

Die Schreibweise $f^{-1}(f(x))$ bedeutet, dass zuerst f auf x angewendet wird und danach f^{-1} auf $f(x)$.

Beispiel

Sei $f(x) = \frac{5}{x}$ mit $f^{-1}(y) = \frac{5}{y}$ (vgl. 3. Beispiel von Seite 7).

Für $f^{-1}(f(x))$ schreibt man auch $(f^{-1} \circ f)(x)$ und meint damit also die Komposition (d.h. Verkettung) von f und f^{-1} .

Definition Seien $f : X \rightarrow Y$ und $g : Y \rightarrow Z$ zwei Abbildungen. Dann ist die *Komposition* von f und g definiert als die Funktion $g \circ f : X \rightarrow Z$ mit

$$(g \circ f)(x) = g(f(x)) \quad \text{für alle } x \in X.$$

Achtung: Die Komposition $g \circ f$ liest man von rechts nach links, das heisst zuerst f , dann g !

1.3 Elementare algebraische Funktionen

Reelle Funktionen lassen sich in zwei Klassen einteilen. Wir betrachten zuerst die sogenannten *algebraischen* Funktionen: Ihre Funktionsvorschrift $f(x)$ entsteht aus der Variablen x und reellen Konstanten durch Anwendung der Grundoperationen $+$, $-$, $\cdot$, $/$ und $\sqrt[n]{}$. Alle übrigen Funktionen nennt man *transzendent*. Zu den wichtigsten transzendenten Funktionen gehören die trigonometrischen Funktionen sowie die Exponential- und Logarithmusfunktionen. Diese behandeln wir in den folgenden beiden Abschnitten.

Polynomfunktionen

Definition Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit der Gleichung

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

wobei $a_0, \dots, a_n \in \mathbb{R}$, $a_n \neq 0$ und $n \in \mathbb{N} \cup \{0\}$, heisst *Polynomfunktion* (oder kurz *Polynom*). Man nennt a_n den *Leitkoeffizienten* von f und n den *Grad* von f .

Einige spezielle Polynomfunktionen kennen Sie aus der Schule:

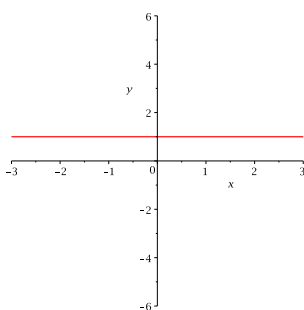
Konstante Funktionen $f(x) = c$ mit $c \in \mathbb{R}$ (Polynom vom Grad 0)

Lineare Funktionen $f(x) = ax + b$ mit $a, b \in \mathbb{R}$, $a \neq 0$ (Grad 1)

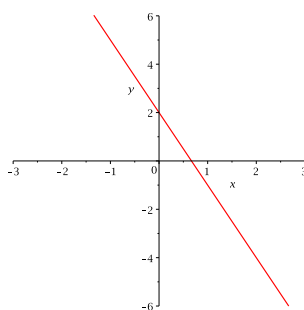
Quadratische Funktionen $f(x) = ax^2 + bx + c$ mit $a, b, c \in \mathbb{R}$, $a \neq 0$ (Grad 2)

Potenzfunktionen $f(x) = x^n$ (Grad n)

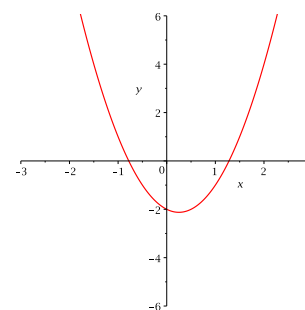
Hier sind die Graphen einiger Beispiele:



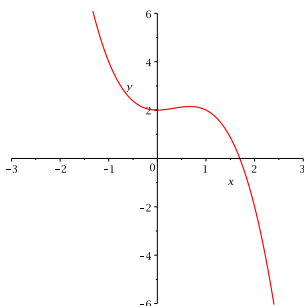
$$f_1(x) = 1$$



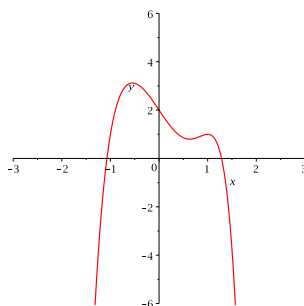
$$f_2(x) = -3x + 2$$



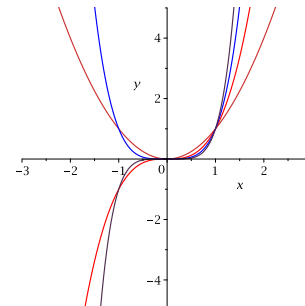
$$f_3(x) = 2x^2 - x - 2$$



$$f_4(x) = -x^3 + x^2 + 2$$



$$f_5(x) = -x^6 + 3x^3 - 3x + 2$$



$$f_6(x) = x^n \quad (n = 2, 3, 4, 5)$$

Offensichtlich verhalten sich die verschiedenen Polynomfunktionen sehr unterschiedlich, so dass kaum allgemeine Aussagen über diese Funktionen gemacht werden können. Was wir hingegen leicht von der Funktionsvorschrift ablesen können, ist das Verhalten der Funktion (oder des Graphen) für x gegen $\pm\infty$.

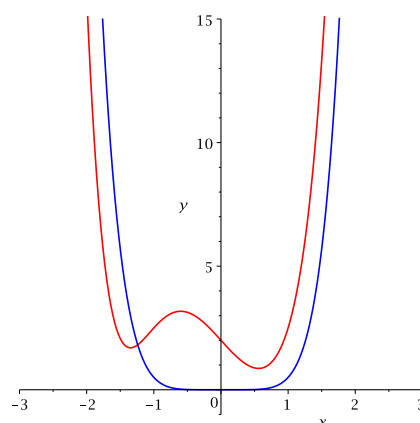
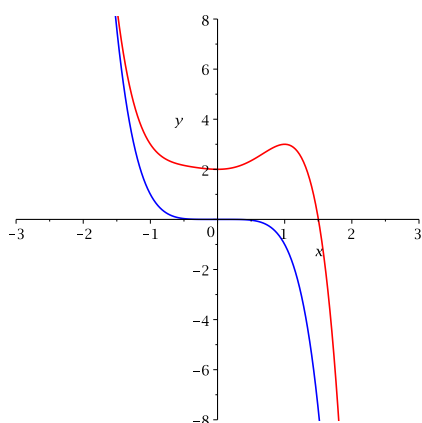
Sei $f(x) = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$ die Funktionsvorschrift. Klammern wir den Ausdruck $a_n x^n$ aus, so erhalten wir

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

Für x gegen $\pm\infty$ geht der Ausdruck in der Klammer gegen 1, das heisst, der Graph von f nähert sich immer mehr dem Graphen der Funktion $g(x) = a_n x^n$.

Beispiele

$f(x) = -x^5 + x^3 + x^2 + 2$ mit $g(x) = -x^5$ und $f(x) = \frac{1}{2}x^6 + 3x^3 - 3x + 2$ mit $g(x) = \frac{1}{2}x^6$



Rationale Funktionen

Definition Eine Funktion $f : D \rightarrow \mathbb{R}$ mit der Vorschrift

$$f(x) = \frac{p(x)}{q(x)},$$

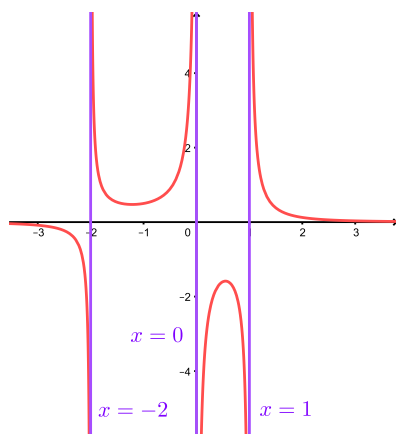
wobei $p(x)$ und $q(x)$ Polynome sind, heisst *rationale Funktion*.

Für den maximalen Definitionsbereich D einer rationalen Funktion f gilt

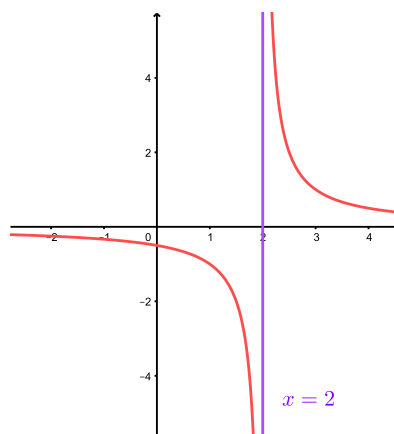
$$D = \mathbb{R} \setminus \{x_1, \dots, x_k\},$$

wobei $x_1, \dots, x_k$ die verschiedenen Nullstellen des Nennerpolynoms $q(x)$ sind. Gilt für eine solche Nullstelle x_i gleichzeitig $p(x_i) \neq 0$, dann bildet diese eine sogenannte *Polstelle* von f . Dies bedeutet, dass die Gerade $x = x_i$ eine *senkrechte Asymptote* für den Graphen von f ist. Anschaulich ist eine Asymptote eine Gerade, an die sich der Funktionsgraph anschmiegt.

Im folgenden ersten Beispiel mit $D = \mathbb{R} \setminus \{0, 1, -2\}$ ist dies für $0, 1, -2$ der Fall, das heisst, es gibt drei senkrechte Asymptoten $x = 0$, $x = 1$, $x = -2$. Im zweiten Beispiel mit $D = \mathbb{R} \setminus \{-1, 2\}$ ist $x = 2$ eine Asymptote, aber $x = -1$ nicht, denn $x + 1 = 0$ im Zähler von f für $x = -1$.

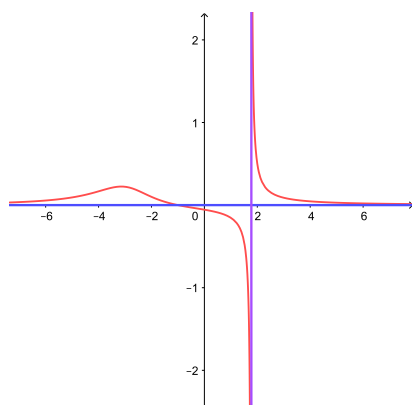


$$f(x) = \frac{1}{x^3 + x^2 - 2x} = \frac{1}{x(x-1)(x+2)}$$

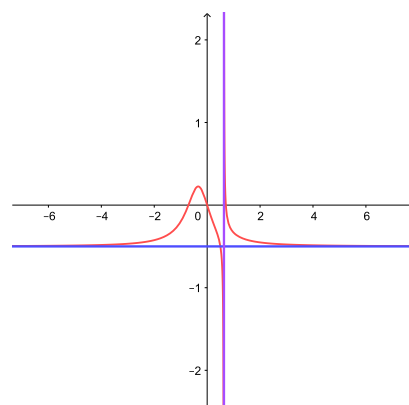


$$f(x) = \frac{x+1}{x^2 - x - 2} = \frac{x+1}{(x+1)(x-2)}$$

Ist $n \leq m$ für den Grad n von $p(x)$ und m von $q(x)$, dann gibt es eine *waagrechte Asymptote* für den Graphen von f (der Funktionsgraph schmiegt sich für x gegen $\pm\infty$ an diese waagrechte Gerade).

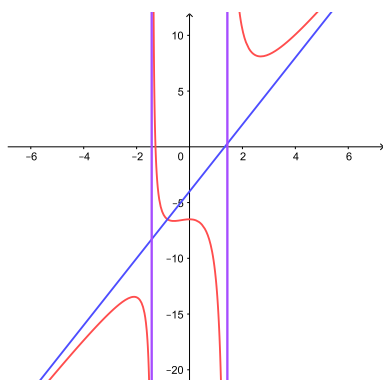


$$f(x) = \frac{x+1}{x^3 + 4x^2 - 18}, \text{ Asymptote: } y = 0$$

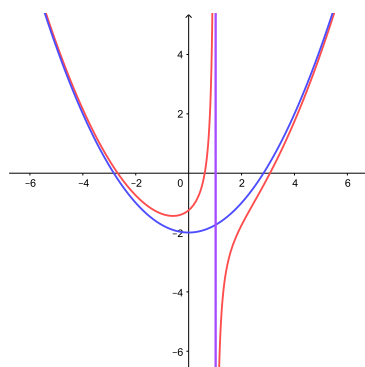


$$f(x) = \frac{-2x^4 + x^2}{4x^4 - x}, \text{ Asymptote: } y = -\frac{1}{2}$$

Im Fall $n > m$ muss die Asymptote nicht notwendigerweise eine Gerade sein.



$$f(x) = \frac{3x^3 - 4x^2 + 13}{x^2 - 2} = 3x - 4 + \frac{6x+5}{x^2-2}$$



$$f(x) = \frac{x^3 - x^2 - 8x + 5}{4x - 4} = \frac{1}{4}x^2 - 2 - \frac{3}{4x-4}$$

Die Asymptoten $g(x) = 3x - 4$, bzw. $g(x) = \frac{1}{4}x^2 - 2$ in den beiden Beispielen findet man durch Polynomdivision.

Für das erste Beispiel führen wir hier die Polynomdivision durch:

Allgemein für $f(x) = \frac{p(x)}{q(x)}$ mit $n = \text{Grad}(p) > \text{Grad}(q) = m$ gibt es Polynome $g(x)$ und $r(x)$ mit $\text{Grad}(g) = n - m > 0$ und $\text{Grad}(r) < \text{Grad}(q)$, so dass

$$p(x) = g(x) \cdot q(x) + r(x)$$

(analog zur Division mit Rest bei den natürlichen Zahlen). Division durch $q(x)$ ergibt

$$f(x) = \frac{p(x)}{q(x)} = g(x) + \frac{r(x)}{q(x)}.$$

Für x gegen $\pm\infty$ geht der zweite Term gegen 0 (da $\text{Grad}(r) < \text{Grad}(q)$), so dass das Polynom $g(x)$ eine Asymptote für den Graphen von f bildet.

Allgemeine Potenzen

Für n in $\mathbb{N}$ ist die Potenzfunktion $x^n = x \cdot x \cdots x$ (n Faktoren) für alle x in $\mathbb{R}$ definiert.

Für $x \neq 0$ definieren wir weiter

$$x^0 = 1 \quad \text{und} \quad x^{-n} = \frac{1}{x^n} \quad \text{für } n \in \mathbb{N}.$$

Damit ist x^n für alle n in $\mathbb{Z}$ und ($x \neq 0$) definiert und es gelten die üblichen Regeln:

$$x^{m+n} = x^m \cdot x^n, \quad x^{m-n} = x^{m+(-n)} = x^m \cdot x^{-n} = \frac{x^m}{x^n}, \quad (x^n)^m = (x^m)^n = x^{m \cdot n}.$$

Achtung: Die Klammern bei der dritten Regel sind notwendig, denn x^{n^m} wird als $x^{(n^m)}$ interpretiert, was im Allgemeinen ungleich $(x^n)^m = x^{n \cdot m}$ ist. Zum Beispiel ist

$$\begin{aligned} (2^3)^4 &= 8^4 = 4096 \\ 2^{(3^4)} &= 2^{81} \approx 2,42 \cdot 10^{24}. \end{aligned}$$

Wie können wir weiter $x^{\frac{1}{q}}$ (für $q \in \mathbb{N}$) sinnvoll definieren? Stellen wir die Bedingung, dass die vorhergehenden Regeln auch für rationale Exponenten gelten sollen, dann folgt

$$(x^{\frac{1}{q}})^q = x^{\frac{1}{q}} \cdot x^{\frac{1}{q}} \cdots x^{\frac{1}{q}} = x^{\frac{1}{q} + \frac{1}{q} + \cdots + \frac{1}{q}} = x^{q \cdot \frac{1}{q}} = x^1 = x,$$

wobei im Produkt q Faktoren und in der Summe q Summanden vorkommen. Dies führt zur folgenden Definition.

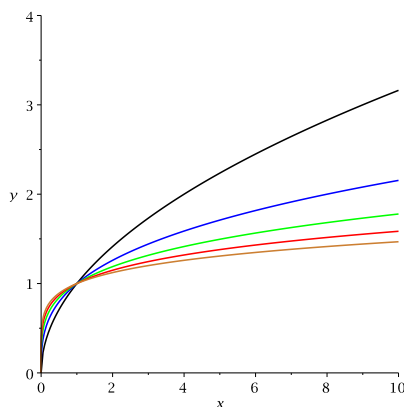
Definition Sei $x > 0$ und q in $\mathbb{N}$. Die q -te Wurzel von x ist definiert durch

$$x^{\frac{1}{q}} = \sqrt[q]{x} = \text{diejenige positive reelle Zahl, deren } q\text{-te Potenz gleich } x \text{ ist.}$$

Zusätzlich gilt $\sqrt[q]{0} = 0$.

Insbesondere ist also $x^{\frac{1}{2}} = \sqrt{x}$ für $x > 0$ mit der üblichen Konvention, dass $\sqrt{x}$ stets die *positive* Quadratwurzel bezeichnet.

Die Graphen von $f(x) = \sqrt[q]{x}$ für $q = 2, 3, 4, 5, 6$ sehen so aus:



Allgemein definieren wir nun für $\frac{p}{q}$ in $\mathbb{Q}$ (mit q in $\mathbb{N}$) und $x > 0$

$$x^{\frac{p}{q}} = \left(x^{\frac{1}{q}}\right)^p = (\sqrt[q]{x})^p = \sqrt[q]{x^p} = (x^p)^{\frac{1}{q}}$$

wobei dies alles verschiedene Schreibweisen für dasselbe sind.

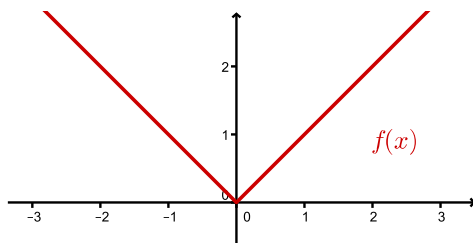
Auf die Frage, wie wir zum Beispiel $x^{\sqrt{2}}$ oder x^π sinnvoll definieren können, kommen wir im übernächsten Abschnitt 1.5 zurück.

Die Betragsfunktion

Definition Die *Betragsfunktion* ist definiert durch

$$f(x) = |x| = \begin{cases} x & \text{falls } x \geq 0 \\ -x & \text{falls } x < 0 \end{cases}.$$

Anschaulich ist $|x|$ der Abstand zwischen der Zahl x und dem Ursprung auf der x -Achse. Hier ist der Graph von f :



Beträge kommen auch in Gleichungen und Ungleichungen vor. Oft tritt dann nicht $|x|$ auf, sondern eine ganze Seite der (Un-)Gleichung steht im Betrag, das heisst, eine Seite der (Un-)Gleichung ist von der Form $|f(x)|$ für irgendeine reelle Funktion f . Um eine solche (Un-)Gleichungen umformen oder lösen zu können, muss zuerst der Betrag aufgelöst werden. Dies geschieht mit Hilfe einer Fallunterscheidung, denn gemäss der Definition des Betrags gilt

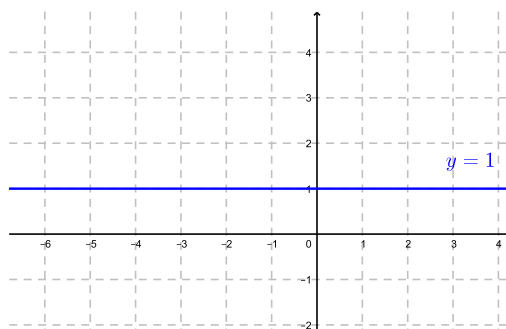
$$|f(x)| = \begin{cases} f(x) & \text{falls } f(x) \geq 0 \\ -f(x) & \text{falls } f(x) < 0 \end{cases}$$

für jede reelle Funktion f .

Beispiele

1. Gesucht sind alle $x \in \mathbb{R}$ mit $|x + 3| = 1$.

Graphisch:



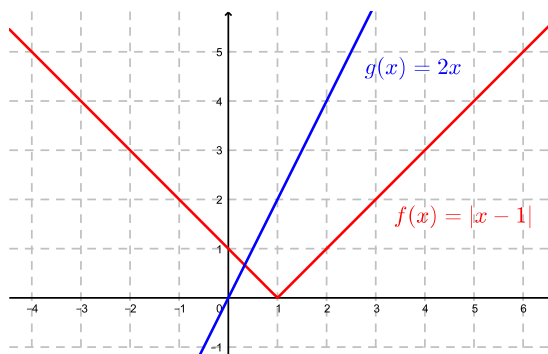
Allgemein beschreibt die Gleichung $|x - a| = c$ die Menge aller $x \in \mathbb{R}$, die auf der x -Achse den Abstand c zur reellen Zahl a haben.

2. Gesucht sind alle $x \in \mathbb{R}$ mit $|x - 1| = 2x$.

- $x - 1 \geq 0$: Dann gilt $x - 1 = 2x \implies -1 = x$ Widerspruch zu $x - 1 \geq 0$!
- $x - 1 < 0$: Dann gilt $-(x - 1) = 2x \implies 1 = 3x \implies \underline{\underline{x = \frac{1}{3}}}$

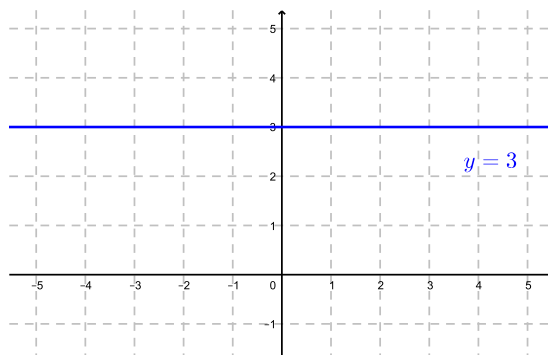
Die Gleichung $|x - 1| = 2x$ hat also nur eine Lösung: $x = \frac{1}{3}$

Graphisch:



3. Gesucht sind alle $x \in \mathbb{R}$ mit $|x^2 - 1| > 3$.

Graphisch:



Gleichwertige Schreibweisen für die Lösungsmenge sind:

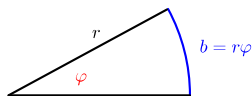
1.4 Trigonometrische Funktionen

Der *Winkel* ist das Mass einer Rotation, und zwar einer Drehung im positiven Drehsinn (d.h. Gegenurzeigersinn). Wir zählen die Drehungen: ganze Drehungen plus Teile von ganzen Drehungen. Die üblichen Skalen für Winkel sind das Gradmass und das Bogenmass:

Gradmass: Mass für die Drehung = 360°

Bogenmass: Mass für die Drehung = 2π = Umfang des Kreises vom Radius 1

Das Bogenmass entspricht der Länge eines Bogens auf dem Kreis vom Radius 1 mit dem entsprechenden Zentriwinkel.



Für den Radius $r = 1$ folgt $b = \varphi$.

Drehungen	$-\frac{1}{4}$	0	$\frac{1}{4}$	$\frac{1}{2}$	$\frac{3}{4}$	1	$1\frac{1}{2}$	2
Gradmass	-90°	0°	90°	180°	270°	360°	540°	720°
Bogenmass	$-\frac{\pi}{2}$	0	$\frac{\pi}{2}$	π	$\frac{3\pi}{2}$	2π	3π	4π

Das Verhältnis $\alpha : 360^\circ = \varphi : 2\pi$ führt zu den Umrechnungen:

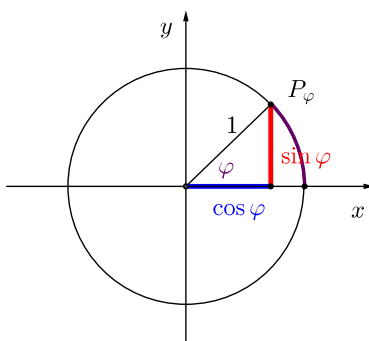
$$\alpha \text{ Winkel im Gradmass} \implies \varphi = \frac{\alpha}{360^\circ} 2\pi \text{ Winkel im Bogenmass}$$

$$\varphi \text{ Winkel im Bogenmass} \implies \alpha = \frac{\varphi}{2\pi} 360^\circ \text{ Winkel im Gradmass}$$

Beispiel

Wir betrachten die Drehung als Prozess, und daher ist $540^\circ \neq 180^\circ$, $360^\circ \neq 0^\circ$ und auch $-180^\circ \neq 180^\circ$. Nimmt man jedoch das Resultat der Drehung, das heisst die Endlage, dann gilt $540^\circ = 180^\circ$, $360^\circ = 0^\circ$ und auch $-180^\circ = 180^\circ$.

Nun starten wir mit dem Punkt $P_0 = (1, 0)$ auf dem Einheitskreis (das ist der Kreis mit Radius 1 und Mittelpunkt $(0, 0)$) und drehen ihn um den Winkel φ . Dies liefert den Punkt P_φ auf der Kreislinie. Wie oben bemerkt, gilt dann $P_{\varphi+2\pi} = P_\varphi$.



Definition Die trigonometrischen Funktionen $\cos : \mathbb{R} \longrightarrow \mathbb{R}$ (*Cosinus*) und $\sin : \mathbb{R} \longrightarrow \mathbb{R}$ (*Sinus*) werden wie folgt definiert:

$$\cos \varphi = x\text{-Koordinate des Punktes } P_\varphi$$

$$\sin \varphi = y\text{-Koordinate des Punktes } P_\varphi$$

Der Winkel φ wird dabei üblicherweise im Bogenmass angegeben.

Es gilt $\cos \varphi \in [-1, 1]$ und $\sin \varphi \in [-1, 1]$ für alle φ in $\mathbb{R}$, und jeder Wert zwischen -1 und 1 wird von beiden Funktionen angenommen, das heisst, die Bildmenge der Funktionen $\cos$ und $\sin$ ist jeweils das ganze Intervall $[-1, 1]$.

Eigenschaften von cos und sin

Die folgenden Eigenschaften können direkt am Einheitskreis abgelesen werden. Sie sollten aus der Schule bekannt sein.

(1) *Periodizität*: Es gilt $P_{\varphi+2\pi} = P_{\varphi}$ und daher

$$\cos(\varphi + 2\pi k) = \cos \varphi \quad \text{und} \quad \sin(\varphi + 2\pi k) = \sin \varphi \quad \text{für alle } k \in \mathbb{Z},$$

das heisst, cos und sin sind *periodische* Funktionen mit der Periode 2π .

(2) *Pythagoras*: $\cos^2 \varphi + \sin^2 \varphi = 1$ für alle φ in $\mathbb{R}$.

(Hier ist $\cos^2 \varphi$ die übliche Kurzschreibweise für $(\cos \varphi)^2$.)

(3) *Nullstellen*:

$$\cos \varphi = 0 \iff P_{\varphi} \text{ auf } y\text{-Achse} \iff \varphi = \frac{\pi}{2} + k\pi \text{ für } k \in \mathbb{Z}$$

$$\sin \varphi = 0 \iff P_{\varphi} \text{ auf } x\text{-Achse} \iff \varphi = k\pi \text{ für } k \in \mathbb{Z}$$

(4) *Symmetrien*:

$$\cos(\varphi + \pi) = -\cos \varphi \qquad \cos(\varphi + \frac{\pi}{2}) = -\sin \varphi$$

$$\sin(\varphi + \pi) = -\sin \varphi \qquad \sin(\varphi + \frac{\pi}{2}) = \cos \varphi$$

und

$$\cos(-\varphi) = \cos \varphi$$

$$\sin(-\varphi) = -\sin \varphi$$

Die letzten zwei Zeilen besagen, dass cos eine gerade und sin eine ungerade Funktion ist.

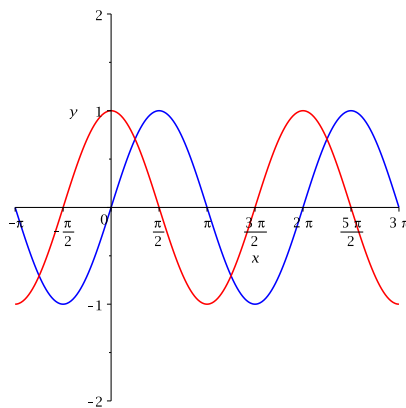
Dabei heisst eine reelle Funktion *gerade*, wenn $f(-x) = f(x)$ für alle x in D und sie heisst *ungerade*, wenn $f(-x) = -f(x)$ für alle x in D . Der Graph einer geraden Funktion ist achsensymmetrisch zur y -Achse, der Graph einer ungeraden Funktion ist punktsymmetrisch zum Ursprung.

(5) *Additionstheoreme*:

$$\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$$

$$\sin(\alpha + \beta) = \sin \alpha \cos \beta + \cos \alpha \sin \beta$$

Die Graphen von **cos** und **sin**:

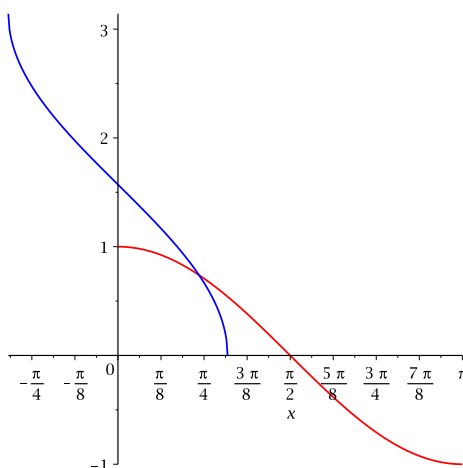


Wegen der Gleichung $\cos \varphi = \sin(\varphi + \frac{\pi}{2})$ ist der Graph von cos gegenüber dem Graphen von sin um $\frac{\pi}{2}$ nach links verschoben.

Umkehrfunktionen von $\cos$ und $\sin$

Die Cosinusfunktion ist als Funktion von $\mathbb{R}$ nach $\mathbb{R}$ nicht umkehrbar. Doch auf gewissen Intervallen ist sie streng monoton fallend, zum Beispiel auf dem Intervall $[0, \pi]$. Verkleinern wir noch den Zielbereich auf die Bildmenge $[-1, 1]$, dann erhalten wir die umkehrbare Funktion $\cos : [0, \pi] \rightarrow [-1, 1]$. Die Umkehrfunktion heisst *Arcuscosinusfunktion*

$$\arccos : [-1, 1] \rightarrow [0, \pi]$$

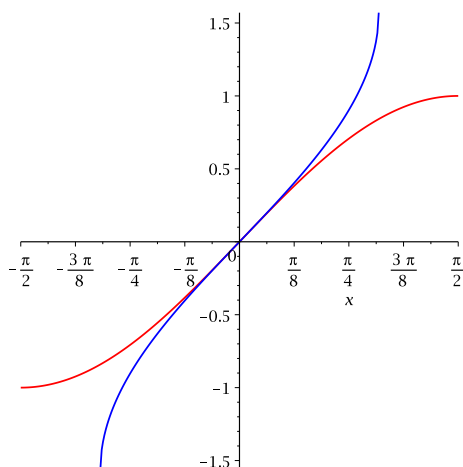


Achtung: Beim Lösen der Gleichung $\cos(x) = a$ ist $x_1 = \arccos(a)$ nur *eine* von unendlich vielen Lösungen, nämlich diejenige im Intervall $[0, \pi]$. Da die Cosinusfunktion 2π -periodisch ist, gibt man jeweils die Lösungen in einem Intervall der Länge 2π an. Bei $\cos$ bietet sich das Intervall $[-\pi, \pi]$ an, denn wegen $\cos(-x) = \cos(x)$ ist $x_2 = -x_1 = -\arccos(a) \in [-\pi, 0]$ eine weitere Lösung von $\cos(x) = a$.

$\implies$ Lösungen von $\cos(x) = a$ im Intervall $[-\pi, \pi]$ sind $x_{1,2} = \pm \arccos(a)$!

Analog ist die Sinusfunktion streng monoton wachsend auf dem Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$, das heisst $\sin : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow [-1, 1]$ ist umkehrbar. Die Umkehrfunktion nennt man *Arcussinusfunktion*

$$\arcsin : [-1, 1] \rightarrow [-\frac{\pi}{2}, \frac{\pi}{2}]$$



Achtung: Beim Lösen der Gleichung $\sin(x) = b$ ist $x_1 = \arcsin(b)$ nur *eine* Lösung, nämlich diejenige im Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$. Bei der Sinusfunktion bietet sich $[-\frac{\pi}{2}, \frac{3\pi}{2}]$ als Intervall der Länge 2π an, denn wegen $\sin(\pi - x) = -\sin(-x) = \sin(x)$ ist $x_2 = \pi - x_1 = \pi - \arcsin(b) \in [\frac{\pi}{2}, \frac{3\pi}{2}]$ eine weitere Lösung von $\sin(x) = b$.

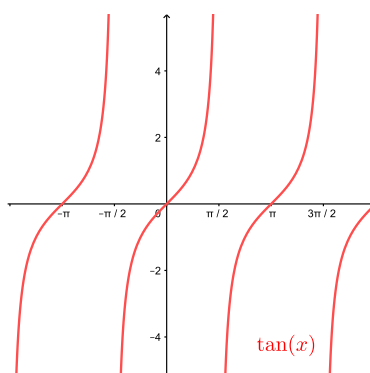
$\Rightarrow$ Lösungen von $\sin(x) = b$ im Intervall $[-\frac{\pi}{2}, \frac{3\pi}{2}]$ sind $x_1 = \arcsin(b)$ und $x_2 = \pi - \arcsin(b)$!

Die Tangensfunktion

Definition Die trigonometrische Funktion $\tan : D \rightarrow \mathbb{R}$ (*Tangens*) ist definiert durch

$$\tan \varphi = \frac{\sin \varphi}{\cos \varphi}.$$

Der Definitionsbereich ist $D = \{\varphi \in \mathbb{R} \mid \cos \varphi \neq 0\} = \mathbb{R} \setminus \{\frac{\pi}{2} + k\pi \mid k \in \mathbb{Z}\}$.

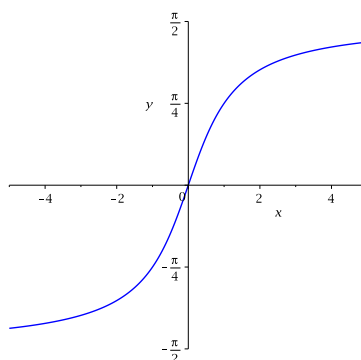


Eigenschaften der Tangensfunktion:

- (1) Es gilt $\tan(-\varphi) = -\tan \varphi$, d.h. $\tan$ ist eine ungerade Funktion.
- (2) Es gilt $\tan(\varphi + k\pi) = \tan \varphi$ für alle k in $\mathbb{Z}$, d.h. $\tan$ ist periodisch mit der Periode π .
- (3) Auf dem offenen Intervall $(-\frac{\pi}{2}, \frac{\pi}{2})$ ist $\tan$ streng monoton wachsend und die Bildmenge ist gleich $\mathbb{R}$.

Wegen der dritten Eigenschaft ist $\tan : (-\frac{\pi}{2}, \frac{\pi}{2}) \rightarrow \mathbb{R}$ umkehrbar. Die Umkehrfunktion heisst *Arcustangensfunktion*

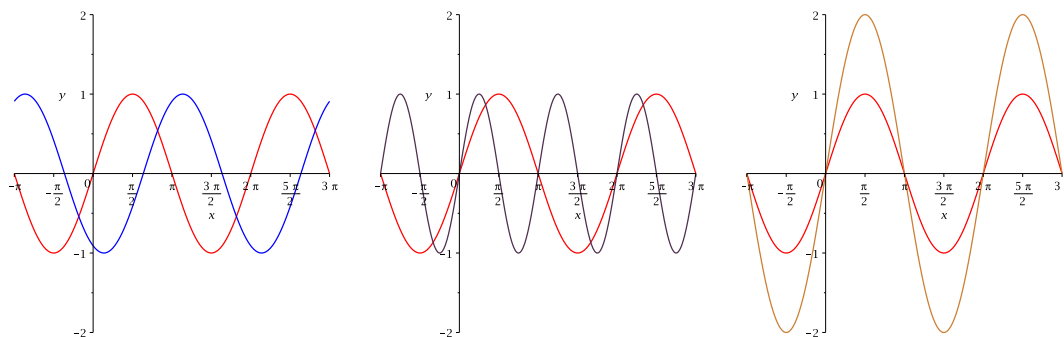
$$\arctan : \mathbb{R} \rightarrow (-\frac{\pi}{2}, \frac{\pi}{2})$$



Modifizierte trigonometrische Funktionen

Die Graphen der Funktionen $f(x) = \sin(x-u)$, $g(x) = \sin(\alpha x)$ und $h(x) = A \sin(x)$ für u, α, A in $\mathbb{R}$ sind gegenüber der Sinuskurve um u nach rechts verschoben, bzw. in der x -Richtung gestaucht (für $|\alpha| > 1$) oder gestreckt (für $|\alpha| < 1$), bzw. in der y -Richtung gestreckt (für $|A| > 1$) oder gestaucht (für $|A| < 1$).

Hier die Graphen von $\sin(x)$, $\sin(x-2)$, $\sin(2x)$, $2\sin(x)$:

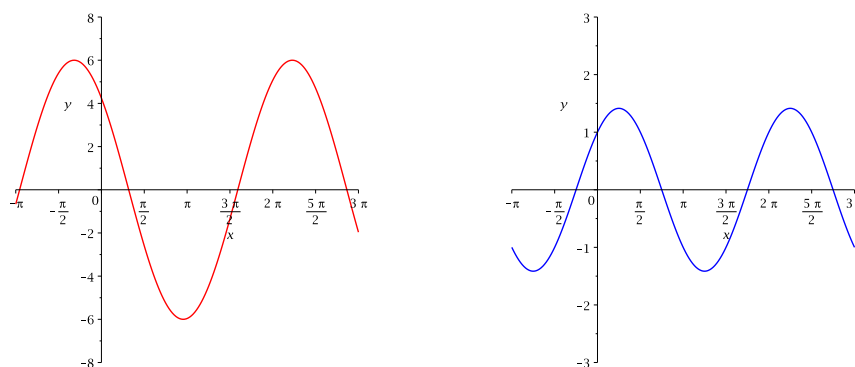


Kombinationen von diesen Modifikationen tauchen bei Schwingungsproblemen in der Physik auf und haben meistens die Form

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

Dabei ist $t \in \mathbb{R}$ die Zeit, ω die Kreisfrequenz, $A (> 0)$ die Amplitude und φ die Phase. Diese Funktion ist periodisch mit der Periode (= Schwingungsdauer) $T = \frac{2\pi}{\omega}$. Ihre Werte liegen zwischen $-A$ und A und ihr Graph ist um φ nach links verschoben.

Hier der Graph von $y(t) = A \sin(\omega t + \varphi)$ mit $A = 6$, $T = \frac{2\pi}{\omega} = 8$, $\varphi = \frac{3\pi}{4}$:



Die blaue Kurve auf der rechten Seite ist der Graph der Funktion $f(x) = \sin(x) + \cos(x)$. Er sieht aus wie eine verschobene und in Richtung der y -Achse gestreckte Sinuskurve.

Satz 1.3 Jede Linearkombination $A \sin(x) + B \cos(x)$ ist eine (modifizierte) trigonometrische Funktion und lässt sich in der Form

$$A \sin(x) + B \cos(x) = C \sin(x + u)$$

darstellen. Dabei ist $C = \sqrt{A^2 + B^2}$, und u ist bestimmt durch die Gleichungen $A = C \cos(u)$, $B = C \sin(u)$. Umgekehrt lässt sich jede modifizierte trigonometrische Funktion $C \sin(x + u)$ (oder $C \cos(x + v)$) als Linearkombination $A \sin(x) + B \cos(x)$ darstellen.

Die Gleichungen für u folgen direkt aus dem Additionstheorem für die Sinusfunktion,

$$C \sin(x+u) = \underbrace{C \cos(u)}_{=A} \sin(x) + \underbrace{C \sin(u)}_{=B} \cos(x) = A \sin(x) + B \cos(x),$$

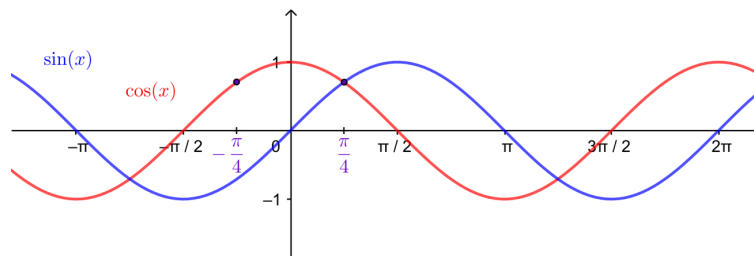
und die Gleichung für C gilt, da

$$A^2 + B^2 = C^2 \cos^2(u) + C^2 \sin^2(u) = C^2 \underbrace{(\cos^2(u) + \sin^2(u))}_{=1} = C^2.$$

Beispiel

$$\sin(x) + \cos(x) = ?$$

Graphisch sieht die Situation im Beispiel so aus:



Zur eindeutigen Bestimmung des Winkels u in Satz 1.3 sind also tatsächlich beide Gleichungen $A = C \cos u$ und $B = C \sin u$ notwendig!

1.5 Exponential- und Logarithmusfunktionen

Definition Die *natürliche Exponentialfunktion* $f(x) = e^x = \exp(x)$ ist definiert durch

$$e^x = 1 + \frac{x}{1!} + \frac{x^2}{2!} + \frac{x^3}{3!} + \cdots = \sum_{k=0}^{\infty} \frac{x^k}{k!}$$

für alle x in $\mathbb{R}$.

Der Ausdruck auf der rechten Seite ist eine sogenannte Potenzreihe. Potenzreihen und Reihen im Allgemeinen werden wir in Abschnitt 3.2 näher betrachten.

Für $x = 1$ in der Definition erhält man die *Eulersche Zahl*

$$e = 1 + \frac{1}{1!} + \frac{1^2}{2!} + \frac{1^3}{3!} + \cdots = \sum_{k=0}^{\infty} \frac{1}{k!} \approx 2,7182818284590452354.$$

Die Zahl e kann auch als Grenzwert einer Zahlenfolge definiert werden,

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n.$$

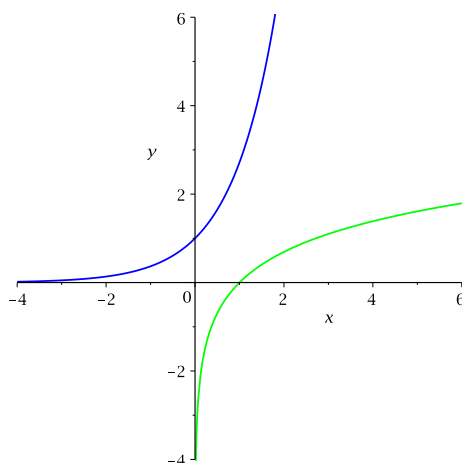
Wir werden dies jedoch aus der obigen Definition in Abschnitt 4.3 herleiten.

Die natürliche Exponentialfunktion ist auf ganz $\mathbb{R}$ streng monoton wachsend und ihre Bildmenge ist $\mathbb{R}_{>0}$. Die Funktion $e^x : \mathbb{R} \rightarrow \mathbb{R}_{>0}$ ist also umkehrbar.

Definition Die Umkehrfunktion der natürlichen Exponentialfunktion wird *natürliche Logarithmusfunktion* genannt:

$$\begin{array}{ccc} \ln : \mathbb{R}_{>0} & \longrightarrow & \mathbb{R} \\ x & \mapsto & \ln(x) \end{array}$$

Die Graphen von e^x und $\ln(x)$:



Es gibt noch weitere Exponentialfunktionen.

Definition Sei $a > 0$ in $\mathbb{R}$, $a \neq 1$. Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit der Gleichung

$$f(x) = a^x$$

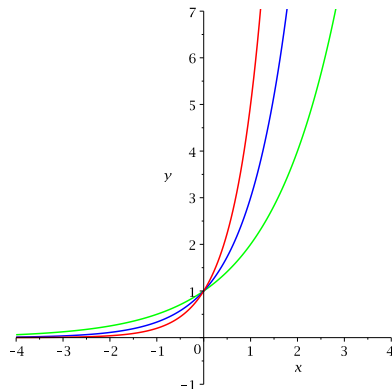
heißt *Exponentialfunktion zur Basis a*.

Wie a^x für x in $\mathbb{Q}$ definiert ist, haben wir in Abschnitt 1.3 (Seite 14) gesehen. Aber was bedeutet nun a^x für beispielsweise $a = 3$ und $x = \sqrt{5} \notin \mathbb{Q}$? Nun, da e^x die Umkehrfunktion von $\ln(x)$ ist, gilt (mit Hilfe der Rechenregel (3) für den natürlichen Logarithmus, s.u.)

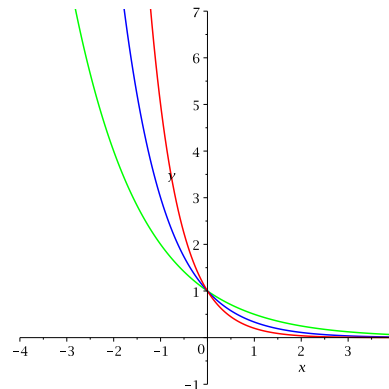
$$a^x = e^{\ln(a^x)} = e^{x \ln(a)}$$

für alle x in $\mathbb{Q}$. Für irrationale x definieren wir nun a^x durch diese Gleichung. Damit ist $3^{\sqrt{5}} = e^{\sqrt{5} \ln(3)}$, wofür man mit der Potenzreihe von e^x problemlos eine gute Näherung erhält.

Hier die Graphen einiger Exponentialfunktionen:



$$f(x) = 2^x, f(x) = 3^x, f(x) = 5^x$$



$$f(x) = \left(\frac{1}{2}\right)^x, f(x) = \left(\frac{1}{3}\right)^x, f(x) = \left(\frac{1}{5}\right)^x$$

Exponentialfunktionen sind sehr wichtige Funktionen, beispielsweise um Wachstumsprozesse mathematisch beschreiben zu können (vgl. Kapitel 6).

Die Funktion $f(x) = a^x$ ist streng monoton wachsend für $a > 1$ und streng monoton fallend für $0 < a < 1$. Die Bildmenge ist gleich $\mathbb{R}_{>0}$ für jedes a . Es folgt wie für e^x , dass die Funktion $f(x) = a^x : \mathbb{R} \rightarrow \mathbb{R}_{>0}$ umkehrbar ist.

Definition Die Umkehrfunktion von $f(x) = a^x$ wird *Logarithmusfunktion zur Basis a* genannt:

$$\begin{aligned} \log_a : \mathbb{R}_{>0} &\longrightarrow \mathbb{R} \\ x &\mapsto \log_a(x) \end{aligned}$$

Zur Erinnerung seien die Logarithmusgesetze erwähnt.

Sei $a > 0$, $a \neq 1$ reell. Dann gilt für alle x, y in $\mathbb{R}_{>0}$:

- (1) $\log_a(x \cdot y) = \log_a(x) + \log_a(y)$
- (2) $\log_a\left(\frac{x}{y}\right) = \log_a(x) - \log_a(y)$
- (3) $\log_a(x^y) = y \cdot \log_a(x)$
- (4) $\log_a(x) = \frac{\ln(x)}{\ln(a)}$

2 Komplexe Zahlen

Für alle reellen Zahlen x gilt $x^2 \geq 0$. Es gibt also keine reelle Zahl, welche Lösung der Gleichung $x^2 + 1 = 0$ ist. Allgemein hat die quadratische Gleichung

$$ax^2 + bx + c = 0, \quad a, b, c \in \mathbb{R}$$

nur dann reelle Lösungen, wenn $b^2 - 4ac \geq 0$ gilt. Im Bereich der reellen Zahlen ist also nicht jede algebraische Gleichung, das heisst

$$a_n x^n + a_{n-1} x^{n-1} + \cdots + a_1 x + a_0 = 0, \quad \text{mit } a_0, \dots, a_n \in \mathbb{R},$$

lösbar. Dieser Mangel motiviert eine Erweiterung des Zahlbereichs $\mathbb{R}$ zu einem Zahlbereich, in dem alle algebraischen Gleichungen lösbar sind. Tatsächlich reicht es aus, den Bereich $\mathbb{R}$ so zu erweitern, dass

1. die spezielle Gleichung $x^2 + 1 = 0$ lösbar ist, das heisst es soll eine *imaginäre Zahl* i geben, so dass $i^2 = -1$;
2. alle Grundrechenarten $+$, $-$, $\cdot$ und $/$ uneingeschränkt durchführbar sind und die Rechenregeln für $\mathbb{R}$ erhalten bleiben.

Der Bereich $\mathbb{C}$ der *komplexen Zahlen* ist der kleinste Bereich mit diesen Eigenschaften. In diesem neuen Bereich sind automatisch alle algebraischen Gleichungen lösbar!

Die Entwicklung der komplexen Zahlen war ein langer Prozess, der mehr als drei Jahrhunderte lang dauerte. Seit dem 19. Jh. werden die komplexen Zahlen jedoch in vielen Gebieten der Mathematik eingesetzt und auch in den Naturwissenschaften sind sie heute unverzichtbar.

2.1 Definitionen und Rechenregeln

Definition Eine *komplexe Zahl* z ist ein Ausdruck der Form

$$z = a + bi \quad \text{mit } a, b \in \mathbb{R}.$$

Die Zahl i , genannt *imaginäre Einheit*, ist definiert durch die Gleichung

$$i^2 = i \cdot i = -1.$$

Die *Menge aller komplexen Zahlen* wird mit $\mathbb{C}$ bezeichnet.

Die Zahl $a = \operatorname{Re}(z)$ heisst *Realteil*, die Zahl $b = \operatorname{Im}(z)$ *Imaginärteil* von z .

Ist $b = 0$, so ist $z = a$ *reell* (insbesondere ist $\mathbb{R} \subset \mathbb{C}$).

Ist $a = 0$, so ist $z = bi$ *rein imaginär*.

Die Darstellung $z = a + bi$ ist eindeutig, das heisst, zwei komplexe Zahlen $z_1 = a + bi$ und $z_2 = c + di$ sind gleich genau dann, wenn $a = c$ und $b = d$.

Definition Sei $z = a + bi$ eine komplexe Zahl.

Dann heisst $\bar{z} = a - bi$ die zu z *konjugiert komplexe Zahl*.

Der *Betrag* $|z|$ von z ist definiert als die reelle Zahl $|z| = \sqrt{a^2 + b^2} \geq 0$.

Die komplexen Zahlen kann man als Punkte in der *Gaußschen Zahlenebene* darstellen. Man erhält $\bar{z}$, indem man z an der reellen Achse spiegelt. Der Betrag $|z|$ entspricht dem Abstand des Punktes z vom Ursprung der Zahlenebene.

Definition (Addition und Subtraktion)

$$(a + bi) \pm (c + di) = (a \pm c) + (b \pm d)i$$

Die Addition und Subtraktion komplexer Zahlen entspricht der Vektoraddition und -subtraktion in der Zahlenebene.

Wie multipliziert man zwei komplexe Zahlen?

Definition (Multiplikation)

$$(a + bi) \cdot (c + di) = (ac - bd) + (ad + bc)i$$

Insbesondere gilt

$$z \cdot \bar{z} = |z|^2.$$

Ist $z \neq 0$, dann können wir diese Gleichung auf beiden Seiten durch die reelle Zahl $|z|^2$ dividieren und wir erhalten $z \cdot \frac{\bar{z}}{|z|^2} = 1$. Wir folgern, dass

$$\frac{1}{z} = \frac{\bar{z}}{|z|^2} = \frac{\bar{z}}{z \cdot \bar{z}} \in \mathbb{C}.$$

Es kann also durch jede komplexe Zahl $z \neq 0$ dividiert werden (man schreibt auch $\frac{1}{z} = z^{-1}$).

Definition (Division) Für $w, z \in \mathbb{C}$, $z \neq 0$, gilt:

$$\frac{w}{z} = \frac{w \cdot \bar{z}}{z \cdot \bar{z}} = \frac{w \cdot \bar{z}}{|z|^2}$$

Das heisst, wir erweitern den Bruch mit der konjugiert komplexen Zahl des Nenners.

Beispiele

$$1. \frac{4}{1+2i} =$$

$$2. \frac{2+i}{3-4i} =$$

$$3. \frac{1}{i} =$$

Die uns bekannten Rechenregeln für reelle Zahlen bleiben erhalten und das Kommutativ- und Assoziativgesetz sowie die Distributivgesetze gelten auch für alle komplexen Zahlen. Während sich die reellen Zahlen mit Hilfe der $<$ -Relation auf dem Zahlenstrahl anordnen lassen, ist dies für die komplexen Zahlen jedoch nicht möglich.

Es ist praktisch, die folgenden Rechenregeln zu kennen.

Satz 2.1 Für $w, z \in \mathbb{C}$ gilt:

$$(a) \overline{w \pm z} = \overline{w} \pm \overline{z}, \quad \overline{w \cdot z} = \overline{w} \cdot \overline{z}, \quad \overline{\left(\frac{w}{z}\right)} = \frac{\overline{w}}{\overline{z}}$$

$$(b) |w \cdot z| = |w| \cdot |z|, \quad \left|\frac{w}{z}\right| = \frac{|w|}{|z|}$$

$$(c) |w + z| \leq |w| + |z| \quad (\text{Dreiecksungleichung})$$

2.2 Algebraische Gleichungen

Nun sind wir schon fähig, jede *quadratische Gleichung* zu lösen.

Beispiele

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

$$2. x^2 = -4$$

$$3. x^2 = -c \quad \text{für eine positive reelle Zahl } c. \text{ Die beiden Lösungen sind}$$

$$x_{1,2} = \pm \sqrt{c} i.$$

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

Hier muss man aufpassen. Für reelles $a > 0$ ist $\sqrt{a}$ die eindeutige positive Wurzel aus a . Hingegen

ist $\sqrt{-a}$ nicht eindeutig, $\sqrt{-a}$ steht für die *zwei* Lösungen der Gleichung $x^2 = -a$. Deshalb ist $\sqrt{-1}$ gleich i oder gleich $-i$. Wir kommen später beim Wurzelziehen (Seite 34) nochmals darauf zurück.

Aufpassen muss man auch beim Multiplizieren von Wurzeln. Für zwei negative reelle Zahlen a, b gilt $\sqrt{ab} \neq \sqrt{a}\sqrt{b}$! Nach obiger Bemerkung ist die linke Seite eindeutig definiert (da $ab > 0$), die rechte Seite jedoch nicht (d.h. nur bis aufs Vorzeichen).

Auf ähnliche Weise wie im 4. Beispiel findet man die Lösungen der allgemeinen quadratischen Gleichung

$$ax^2 + bx + c = 0,$$

wobei a, b, c in $\mathbb{R}$. Die Lösungsformel liefert die Lösungen

$$x_{1,2} = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}.$$

Abhängig von der *Diskriminante* $D = b^2 - 4ac$ tritt nun eine von drei möglichen Situationen ein:

- $D > 0 \implies$ zwei reelle Lösungen x_1, x_2
- $D = 0 \implies$ eine reelle Lösung $x_1 = x_2$
- $D < 0 \implies$ zwei komplexe (nicht-reelle) Lösungen x_1, x_2 , wobei $x_2 = \overline{x_1}$

Beispiel

Gesucht sind alle Lösungen der quadratischen Gleichung $4x^2 + 2x + 1 = 0$.

Die obige Lösungsformel gilt auch für eine quadratische Gleichung $az^2 + bz + c = 0$ mit Koeffizienten a, b, c in $\mathbb{C}$. Nun ist $D = b^2 - 4ac$ eine komplexe Zahl und man muss zwei Fälle unterscheiden: Im Fall $D \neq 0$ gibt es zwei (komplexe) Lösungen und im Fall $D = 0$ gibt es genau eine (komplexe) Lösung. Wir werden am Ende dieses Kapitels eine solche quadratische Gleichung lösen (wenn wir gelernt haben, wie man Wurzeln aus komplexen Zahlen zieht).

Betrachten wir nun eine allgemeine *algebraische Gleichung* über $\mathbb{C}$,

$$a_n z^n + a_{n-1} z^{n-1} + \cdots + a_1 z + a_0 = 0, \quad a_0, \dots, a_n \in \mathbb{C}.$$

Die natürliche Zahl n (falls $a_n \neq 0$) heisst der *Grad* der Gleichung (die linke Seite ist ja ein Polynom). Der Mathematiker CARL FRIEDRICH GAUSS (1777 – 1855) bewies, dass jede algebraische Gleichung über $\mathbb{C}$ mindestens eine (komplexe) Lösung hat. Aus dieser Aussage erhält man den folgenden fundamentalen Satz.

Fundamentalsatz der Algebra Jede algebraische Gleichung über $\mathbb{C}$ vom Grad n hat bei geeigneter Zählweise genau n Lösungen in $\mathbb{C}$.

Im Beweis wird gezeigt, dass jedes Polynom $p(z) = z^n + a_{n-1}z^{n-1} + \dots + a_1z + a_0$ mit $a_0, \dots, a_n \in \mathbb{C}$ als Produkt von genau n Linearfaktoren geschrieben werden kann,

$$p(z) = (z - z_1) \cdot (z - z_2) \cdot \dots \cdot (z - z_n)$$

mit $z_1, \dots, z_n$ in $\mathbb{C}$. Dann ist

$$p(z) = 0 \iff z \in \{z_1, \dots, z_n\}.$$

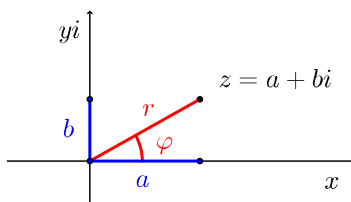
Falls zwei oder mehr der $z_1, \dots, z_n$ übereinstimmen, werden diese Nullstellen doppelt oder mehrfach gezählt.

Beispiel

Für Gleichungen vom Grad 2, 3 und 4 gibt es Lösungsformeln. Es gibt jedoch keine Formeln für allgemeine Gleichungen vom Grad ≥ 5 . In diesen Fällen benutzt man Näherungsverfahren zur Bestimmung der Lösungen.

2.3 Polarkoordinaten und exponentielle Darstellung

Ein Punkt $z = a + bi$ der Gaußschen Zahlenebene ist durch seine *kartesischen Koordinaten* a und b eindeutig festgelegt. Man kann jedoch auch zwei andere Größen zur Beschreibung von z wählen.



Die Größen r und φ nennt man *Polarkoordinaten* des Punktes $z \in \mathbb{C}$. Die reelle Zahl r ist der Abstand von z zum Ursprung in der Gaußschen Zahlenebene und der Winkel φ zwischen positiver x -Achse und z nennt man das *Argument* von z . Das Argument von z ist eindeutig bestimmt, falls $0 \leq \varphi < 2\pi$ verlangt wird. Eine komplexe Zahl kann also entweder in kartesischen Koordinaten (a, b) oder in Polarkoordinaten (r, φ) eindeutig beschrieben werden. Wie hängen die Polar- und die kartesischen Koordinaten zusammen?

Da r der Abstand des Punktes $z = a + bi$ zum Ursprung ist, gilt

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

Damit liegt die komplexe Zahl $\frac{z}{r}$ auf dem Einheitskreis, denn $|\frac{z}{r}| = \frac{|z|}{r} = 1$. Der Real- und der Imaginärteil von z sind demnach gegeben durch $\cos \varphi$, bzw. $\sin \varphi$. Das heisst, $\frac{z}{r} = \cos \varphi + i \sin \varphi$, und so

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

Diese Darstellung nennt man *Polarform* der komplexen Zahl z .

In Polarform kann die zu $z = r(\cos \varphi + i \sin \varphi)$ konjugiert komplexe Zahl $\bar{z}$ geschrieben werden als $\bar{z} = r(\cos \varphi - i \sin \varphi) = r(\cos(-\varphi) + i \sin(-\varphi))$.

Umrechnung $(r, \varphi) \longrightarrow (a, b)$

Seien r, φ die Polarkoordinaten von z , dann gilt

$$z = a + bi \quad \text{mit } a = r \cos \varphi \text{ und } b = r \sin \varphi .$$

Beispiele

1. $r = 2, \varphi = 90^\circ = \frac{\pi}{2}.$

2. $r = \sqrt{3}, \varphi = 230^\circ.$

Umrechnung $(a, b) \longrightarrow (r, \varphi)$

Sei $z = a + bi$, dann gilt

$$\begin{aligned} r &= |z| = \sqrt{a^2 + b^2} \\ \varphi &= \begin{cases} \arccos\left(\frac{a}{r}\right) & \text{falls } b \geq 0 \\ -\arccos\left(\frac{a}{r}\right) & \text{falls } b < 0 \end{cases} . \end{aligned}$$

Warum? Für φ haben wir von oben zunächst die Gleichungen $\cos \varphi = \frac{a}{r}$ und $\sin \varphi = \frac{b}{r}$. Lösungen der ersten Gleichung sind $\varphi = \pm \arccos\left(\frac{a}{r}\right)$, wobei $\arccos\left(\frac{a}{r}\right) \in [0, \pi]$ (vgl. Seite 19). Da $\sin \varphi \geq 0$ für $\varphi \in [0, \pi]$ und $\sin \varphi < 0$ für $\varphi \in (-\pi, 0)$, gilt für $\varphi \in (-\pi, \pi]$

$$\varphi = +\arccos\left(\frac{a}{r}\right) \iff \varphi \geq 0 \iff \sin \varphi \geq 0 \iff \frac{b}{r} \geq 0 \iff b \geq 0 .$$

Beispiele

1. $z = 1 - i$

2. $w = -2 + i$

Mit Hilfe der Exponentialfunktion kann die Polarform einer komplexen Zahl noch eleganter geschrieben werden. Dazu erweitern wir die Exponentialfunktion auf komplexe Zahlen. Tatsächlich ist für jede komplexe Zahl z die unendliche Summe

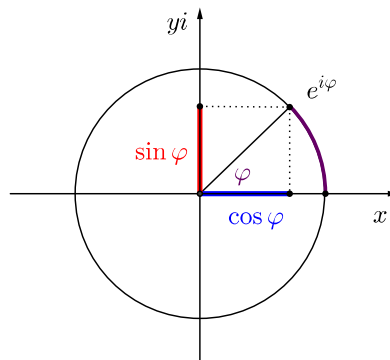
$$e^z = \sum_{k=0}^{\infty} \frac{z^k}{k!}$$

wieder eine komplexe Zahl. Wir haben also eine Funktion $f : \mathbb{C} \rightarrow \mathbb{C}$, $f(z) = e^z$. Wie im Reellen gilt $e^{z+w} = e^z \cdot e^w$ für alle $z, w \in \mathbb{C}$ und weiter ist $e^{\bar{z}} = \overline{e^z}$ für alle $z \in \mathbb{C}$.

Wir betrachten nun e^z für die rein imaginäre Zahl $z = i\varphi$, das heisst für $\varphi \in \mathbb{R}$. Es gilt

$$|e^{i\varphi}|^2 = e^{i\varphi} \cdot \overline{e^{i\varphi}} = e^{i\varphi} \cdot e^{-i\varphi} = e^{i\varphi - i\varphi} = e^0 = 1$$

und damit $|e^{i\varphi}| = 1$. Dies bedeutet, dass $e^{i\varphi}$ auf dem Einheitskreis in der Gaußschen Zahlenebene liegt. Nun kann man zeigen, dass φ gerade das Argument (im Bogenmass) der komplexen Zahl $e^{i\varphi}$ ist. Es folgt, dass der Realteil von $e^{i\varphi}$ gleich $\cos \varphi$ und der Imaginärteil von $e^{i\varphi}$ gleich $\sin \varphi$ ist.



Das heisst, es gilt die

Eulersche Identität $e^{i\varphi} = \cos \varphi + i \sin \varphi$

Als Spezialfall ($\varphi = \pi$) ergibt sich die wunderschöne Beziehung:

$$e^{i\pi} + 1 = 0$$

Ist nun $z = r(\cos \varphi + i \sin \varphi)$, wobei r der Betrag und φ das Argument von z sind, so folgt aus der Eulerschen Identität die

Exponentielle Darstellung $z = r \cdot e^{i\varphi}$

Beispiel

$$z = 1 - i$$

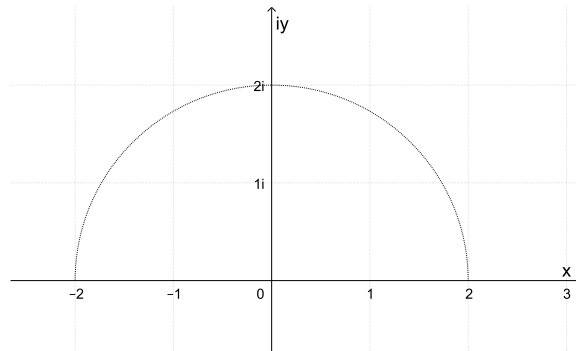
Die Multiplikation und Division von komplexen Zahlen in dieser Darstellung sind nun ganz einfach zu rechnen und erst noch geometrisch interpretierbar!

Satz 2.2 Seien $z_1 = r_1 e^{i\varphi_1}$ und $z_2 = r_2 e^{i\varphi_2}$ zwei komplexe Zahlen. Dann gilt

$$z_1 \cdot z_2 = r_1 r_2 e^{i(\varphi_1 + \varphi_2)} \quad \text{und} \quad \frac{z_1}{z_2} = \frac{r_1}{r_2} e^{i(\varphi_1 - \varphi_2)}$$

Beispiel

$$z_1 = 2 e^{i\frac{\pi}{6}}, \quad z_2 = e^{i\frac{\pi}{2}}$$



Die Multiplikation mit i entspricht also einer Drehung um 90° um den Ursprung der Gaußschen Zahlenebene. Allgemein entspricht die Multiplikation mit einer komplexen Zahl $z = r e^{i\varphi}$ einer Drehstreckung (mit dem Streckfaktor r und dem Drehwinkel φ).

2.4 Potenzen und Wurzeln

Potenzen komplexer Zahlen werden wie im Reellen definiert, das heisst

$$z^0 = 1, \quad z^1 = z, \quad z^n = z^{n-1} \cdot z, \quad \text{und} \quad z^{-n} = \frac{1}{z^n}$$

für alle $n \in \mathbb{N}$. In der exponentiellen Darstellung ist das Potenzieren einfach:

$$z = r e^{i\varphi} \quad \implies \quad z^n = r^n e^{in\varphi}$$

In Polarkoordinaten erhält man

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

Formel von de Moivre $(\cos \varphi + i \sin \varphi)^n = \cos(n\varphi) + i \sin(n\varphi) \quad \text{für } n \in \mathbb{Z}$

Beispiel

$$(1 - i)^{16} = ?$$

Einheitswurzeln

Gesucht sind alle Lösungen der speziellen Gleichung

$$z^n = 1.$$

Wir schreiben die rechte Seite exponentiell, $z^n = 1 = e^{2\pi i}$. Beim Wurzelziehen, als Umkehrung des Potenzierens, dividieren wir das Argument $\varphi = 2\pi$ durch n und erhalten die Lösung

$$\zeta = e^{i\frac{2\pi}{n}} = \cos\left(\frac{2\pi}{n}\right) + i \sin\left(\frac{2\pi}{n}\right).$$

Weiter ist jede Potenz dieser Zahl auch wieder eine Lösung.

Satz 2.3 *Die komplexen Zahlen*

$$\zeta^k = e^{i\frac{2\pi k}{n}} = \cos\left(\frac{2\pi k}{n}\right) + i \sin\left(\frac{2\pi k}{n}\right), \quad k = 0, 1, \dots, n-1,$$

sind die n verschiedenen Lösungen der Gleichung $z^n = 1$. Man nennt sie n -te Einheitswurzeln.

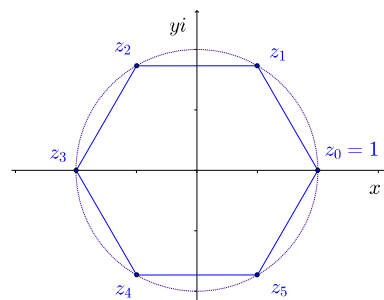
In der Gaußschen Zahlenebene liegen die n -ten Einheitswurzeln genau auf den Ecken des dem Einheitskreis eingeschriebenen regelmäßigen n -Ecks, wobei die eine Ecke bei $z = 1$ liegt.

Beispiele

1. $z^4 = 1$.

2. $z^6 = 1$. Die Lösungen sind $z_k = e^{i\frac{\pi}{3}k}$ für $k = 0, 1, \dots, 5$.

$$\begin{aligned} z_0 &= 1 \\ z_1 &= e^{i\frac{\pi}{3}} \\ z_2 &= e^{i\frac{2\pi}{3}} \\ z_3 &= e^{i\pi} = -1 \\ z_4 &= e^{i\frac{4\pi}{3}} \\ z_5 &= e^{i\frac{5\pi}{3}} \end{aligned}$$



Allgemeine Wurzeln

Gesucht sind nun alle Lösungen der allgemeineren Gleichung

$$z^n = w$$

für $w \in \mathbb{C}$, $w \notin \mathbb{R}_{\geq 0}$. Wir schreiben w in exponentieller Form $w = re^{i\varphi}$ (wie oben bei den Einheitswurzeln die Zahl 1). Ist ζ wie vorher die n -te Einheitswurzel $\zeta = e^{i\frac{2\pi}{n}}$, dann gilt der folgende Satz.

Satz 2.4 *Die komplexen Zahlen*

$$z_0 = \sqrt[n]{r}e^{i\frac{\varphi}{n}} \quad \text{und} \quad z_0\zeta, z_0\zeta^2, \dots, z_0\zeta^{n-1}$$

sind die n verschiedenen Lösungen der Gleichung $z^n = w = re^{i\varphi}$. Man nennt jede dieser Lösungen eine n -te Wurzel von w und bezeichnet sie mit $\sqrt[n]{w}$.

Die Wurzel aus einer komplexen Zahl ist also nicht eindeutig, wie schon auf Seite 28 für $n = 2$ bemerkt.

Beispiele

- Wir betrachten die Gleichung $z^3 = -2$. Mit $-2 = 2e^{i\pi}$ und $\zeta = e^{i\frac{2\pi}{3}}$ finden wir die drei Lösungen

$$z_0 = \sqrt[3]{2}e^{i\frac{\pi}{3}}, \quad z_0\zeta = \sqrt[3]{2}e^{i\pi} = -\sqrt[3]{2}, \quad z_0\zeta^2 = \sqrt[3]{2}e^{i\frac{5\pi}{3}}.$$

- Gesucht sind die Lösungen der quadratischen Gleichung

$$z^2 - (1 + 3i)z - 2 + i = 0.$$

Wir haben also eine quadratische Gleichung $az^2 + bz + c = 0$ mit komplexen Koeffizienten

$$a = 1, \quad b = -(1 + 3i), \quad c = -2 + i.$$

3 Grenzwerte und stetige Funktionen

In diesem und den folgenden zwei Kapiteln behandeln wir die für die Analysis grundlegenden Begriffe Stetigkeit, Differenzierbarkeit und Integration von Funktionen. Diese Begriffe werden durch Folgen von reellen Zahlen definiert. Auch die natürliche Exponentialfunktion wird als Grenzwert einer Zahlenfolge definiert. Wir müssen deshalb in diesem Kapitel zuerst Folgen und Reihen (das sind Summen von Folgengliedern) untersuchen.

Das Ziel dieses Kapitels ist, Stetigkeit zu verstehen und die wichtigsten Eigenschaften von stetigen Funktionen zu kennen. Stetigkeit spielt zum Beispiel bei der Nullstellenbestimmung einer reellen Funktion eine Rolle. Sei zum Beispiel $f(x) = x^5 + 5x + 1$. Gesucht sind alle x in $\mathbb{R}$ mit $f(x) = 0$. Wir könnten einzelne Werte für x ausprobieren, beispielsweise $f(0) = 1 > 0$ und $f(-1) = -5 < 0$. Wir kennen also zwei Punkte auf dem Graphen von f , $(0, 1)$ oberhalb und $(-1, -5)$ unterhalb der x -Achse. Wenn der Graph von f eine ununterbrochene Linie ist, muss der Graph die x -Achse an mindestens einer Stelle schneiden, was einer Nullstelle entspricht. Die Existenz einer Nullstelle im Intervall $(-1, 0)$ hängt also davon ab, ob der Graph von f eine ununterbrochene Linie ist. Ist dies der Fall, nennt man f stetig. Mit Hilfe von Zahlenfolgen kann die Stetigkeit mathematisch präzise beschrieben werden.

3.1 Folgen

Als erstes Beispiel betrachten wir die Folge der Zahlen

$$\frac{1}{2}, \frac{2}{3}, \frac{3}{4}, \dots, \frac{999\,999}{1\,000\,000}, \dots$$

Wir erkennen, dass die n -te Zahl dieser Folge von der Form

$$a_n = \frac{n}{n+1}$$

ist.

Definition Eine *Folge* von reellen Zahlen (oder *Zahlenfolge*) ist eine Funktion $\mathbb{N} \rightarrow \mathbb{R}$, die jedem n in $\mathbb{N}$ ein $a_n = a(n)$ in $\mathbb{R}$ zuordnet (man nennt a_n das *n -te Glied* der Folge). Wir schreiben $(a_n)_{n \in \mathbb{N}}$ oder $(a_n)_{n \geq 1}$ für die Folge.

In vielen Beispielen beginnt die Folge mit einem nullten Glied a_0 . Es sind also auch Folgen $(a_n)_{n \geq 0}$ zulässig.

Beispiele

1. $1, 2, 3, \dots, n, \dots$
2. $1, \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{n}, \dots$
3. $a_n = n^2 - 1$ für $n \geq 1$
4. $a_n = \frac{1}{\sqrt{n+1}}$ für $n \geq 0$
5. $a_n = (-1)^n$ für $n \geq 0$

Wir betrachten nun nochmals das Beispiel ganz oben mit $a_n = \frac{n}{n+1}$, d.h. die Folge $\frac{1}{2}, \frac{2}{3}, \frac{3}{4}, \dots$



Die Folgenglieder a_n kommen der Zahl 1 mit wachsendem n immer näher. Das heisst, der Abstand $|a_n - 1|$ wird beliebig klein. Dieses “beliebig klein” kann wie folgt präzisiert werden.

Wir geben uns ein beliebiges Intervall um die Zahl 1 herum vor und überprüfen, ob ab einem bestimmten Index N alle Folgenglieder a_n ($n \geq N$), das heisst $a_N, a_{N+1}, a_{N+2}, \dots$, in diesem Intervall liegen. Ist dies für jedes noch so kleine Intervall der Fall, dann heisst die Folge *konvergent gegen 1*.

Betrachten wir die folgenden Intervalle:

- $(1 - \frac{1}{4}, 1 + \frac{1}{4}) = (\frac{3}{4}, \frac{5}{4})$ (offenes Intervall, ohne die Grenzen $\frac{3}{4}$ und $\frac{5}{4}$)

- $(1 - \frac{1}{8}, 1 + \frac{1}{8}) = (\frac{7}{8}, \frac{9}{8})$

- $(1 - \varepsilon, 1 + \varepsilon)$ für ein beliebiges $\varepsilon > 0$

Wir werden im 3. Beispiel unten sehen, dass wir jedes N mit $N \geq \frac{1}{\varepsilon}$ wählen können. Es gilt dann, dass a_n für alle $n \geq N$ im Intervall $(1 - \varepsilon, 1 + \varepsilon)$ liegen.

Eine allgemeine Folge $(a_n)_{n \in \mathbb{N}}$ heisst nun *konvergent gegen a* in $\mathbb{R}$, wenn es zu jedem $\varepsilon > 0$ ein N in $\mathbb{N}$ gibt, so dass a_n für alle $n \geq N$ im Intervall $(a - \varepsilon, a + \varepsilon)$ liegen. Man nennt das Intervall $(a - \varepsilon, a + \varepsilon)$ auch ε -Umgebung von a .

Man kann dies auch mit Hilfe des Betrages oder Abstands beschreiben, denn

$$a_n \in (a - \varepsilon, a + \varepsilon) \iff a - \varepsilon < a_n < a + \varepsilon \iff |a_n - a| < \varepsilon$$

Definition Eine Folge $(a_n)_{n \in \mathbb{N}}$ heisst *konvergent gegen a* in $\mathbb{R}$, wenn gilt:

Zu jedem $\varepsilon > 0$ gibt es ein N in $\mathbb{N}$, so dass

$$|a_n - a| < \varepsilon \quad \text{für alle } n \geq N.$$

Die Zahl a heisst *Grenzwert* der Folge. Man schreibt

$$\lim_{n \rightarrow \infty} a_n = a \quad \text{oder} \quad a_n \rightarrow a \quad \text{für } n \rightarrow \infty.$$

Ist die Folge nicht konvergent, dann nennt man sie *divergent*.

In Worten: Eine Folge $(a_n)_{n \in \mathbb{N}}$ ist konvergent gegen a , wenn in jeder beliebig kleinen ε -Umgebung von a fast alle Folgenglieder liegen (nämlich alle bis auf $a_1, a_2, \dots, a_{N-1}$).

Wie im Beispiel oben hängt die Zahl N von ε ab (je kleiner ε , desto grösser muss N im Allgemeinen gewählt werden).

Beispiele

1. Sei a in $\mathbb{R}$ und $a_n = a$ für alle n . Dann gilt $\lim_{n \rightarrow \infty} a_n = a$.

Denn $|a_n - a| = 0 < \varepsilon$ für jedes $\varepsilon > 0$ und alle n in $\mathbb{N}$.

2. $\lim_{n \rightarrow \infty} \frac{1}{n} = 0$.

Sei $\varepsilon > 0$. Gibt es ein N (abhängig von ε), so dass $|\frac{1}{n} - 0| < \varepsilon$ für alle $n \geq N$?

3. $\lim_{n \rightarrow \infty} \frac{n}{n+1} = 1$.

Denn sei $\varepsilon > 0$. Wähle $N \geq \frac{1}{\varepsilon}$. Dann folgt für alle $n \geq N$, dass

$$\left| \frac{n}{n+1} - 1 \right| = \left| \frac{n - (n+1)}{n+1} \right| = \left| \frac{-1}{n+1} \right| = \frac{1}{n+1} < \frac{1}{n} \leq \frac{1}{N} \leq \varepsilon.$$

4. $a_n = (-1)^n$ ist divergent.

Denn unendlich viele Folgenglieder sind gleich 1, aber auch unendlich viele Folgenglieder sind gleich -1 . Für $\varepsilon = \frac{1}{2}$ zum Beispiel gibt es also kein N , so dass die a_n für alle $n \geq N$ im Intervall $(1 - \frac{1}{2}, 1 + \frac{1}{2}) = (\frac{1}{2}, \frac{3}{2})$ liegen. Die Folge konvergiert daher nicht gegen 1. Analog konvergiert sie auch nicht gegen -1 . Ein anderer Grenzwert kommt nicht in Frage.

Bei den divergenten Folgen unterscheidet man Folgen, die gegen $+\infty$ oder $-\infty$ streben, und Folgen, die dies nicht tun.

Definition Eine Folge $(a_n)_{n \in \mathbb{N}}$ heisst

- *bestimmt divergent gegen ∞* , wenn gilt: Zu jedem K in $\mathbb{R}$ gibt es ein N in $\mathbb{N}$ mit

$$a_n \geq K \quad \text{für alle } n \geq N.$$

- *bestimmt divergent gegen $-\infty$* , wenn gilt: Zu jedem K in $\mathbb{R}$ gibt es ein N in $\mathbb{N}$ mit

$$a_n \leq K \quad \text{für alle } n \geq N.$$

Man schreibt $\lim_{n \rightarrow \infty} a_n = \infty$, bzw. $\lim_{n \rightarrow \infty} a_n = -\infty$.

Beispiele

1. Die Folge $(n)_{n \in \mathbb{N}}$ ist bestimmt divergent gegen ∞ .

2. Die Folge $((-1)^n)_{n \in \mathbb{N}}$ ist divergent, aber *nicht* bestimmt divergent.
3. Ein weiteres wichtiges Beispiel ist die *geometrische Folge* $(q^n)_{n \in \mathbb{N}}$. Es gilt

$$\lim_{n \rightarrow \infty} q^n = \begin{cases} 0 & \text{für } |q| < 1 & (\text{konvergent}) \\ 1 & \text{für } q = 1 & (\text{konvergent}) \\ \infty & \text{für } q > 1 & (\text{bestimmt divergent}) \end{cases}$$

Für $q \leq -1$ ist die Folge unbestimmt divergent.

Satz 3.1 (Rechenregeln) Seien $(a_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ konvergente Folgen und c in $\mathbb{R}$.

- (1) $\lim_{n \rightarrow \infty} (a_n + b_n) = \lim_{n \rightarrow \infty} a_n + \lim_{n \rightarrow \infty} b_n$
- (2) $\lim_{n \rightarrow \infty} (c \cdot a_n) = c \cdot \lim_{n \rightarrow \infty} a_n$
- (3) $\lim_{n \rightarrow \infty} (a_n b_n) = (\lim_{n \rightarrow \infty} a_n) \cdot (\lim_{n \rightarrow \infty} b_n)$
- (4) $\lim_{n \rightarrow \infty} \left(\frac{a_n}{b_n} \right) = \frac{\lim_{n \rightarrow \infty} a_n}{\lim_{n \rightarrow \infty} b_n}, \quad \text{falls } b_n \neq 0 \text{ und } \lim_{n \rightarrow \infty} b_n \neq 0$

Diese Rechenregeln kann man teilweise auf bestimmt divergente Folgen erweitern.

- (5) Ist $\lim_{n \rightarrow \infty} a_n = \infty$ und $\lim_{n \rightarrow \infty} b_n = b$, so gilt

$$\lim_{n \rightarrow \infty} (a_n \pm b_n) = \infty, \quad \lim_{n \rightarrow \infty} (a_n b_n) = \begin{cases} \infty & \text{für } b > 0 \\ -\infty & \text{für } b < 0 \end{cases}, \quad \lim_{n \rightarrow \infty} \left(\frac{b_n}{a_n} \right) = 0.$$

- (6) Ist $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n = \infty$, so gilt

$$\lim_{n \rightarrow \infty} (a_n + b_n) = \infty \quad \text{und} \quad \lim_{n \rightarrow \infty} (a_n b_n) = \infty.$$

Diese Rechenregeln sind sehr wichtig und nützlich. Ausgehend vom Wissen, dass

$$\lim_{n \rightarrow \infty} \frac{1}{n} = 0 \quad \text{bzw.} \quad \lim_{n \rightarrow \infty} \frac{1}{n^k} = 0 \quad \text{für } k \geq 1,$$

kann mit Hilfe dieser Rechenregeln der Grenzwert von vielen Zahlenfolgen bestimmt werden. Dabei hat man gleichzeitig elegant bewiesen, dass die Zahlenfolgen konvergent sind.

Beispiele

Typische Beispiele sind Folgen mit *rationalen Ausdrücken*. Man dividiert zuerst Zähler und Nenner durch die höchste im Nenner vorkommende Potenz von n und wendet anschliessend die Rechenregeln an.

$$1. \quad a_n = \frac{n}{n+1}$$

$$2. \quad a_n = \frac{n+2}{5n^2-2n+1}$$

$$3. \quad a_n = \frac{n^2-2n}{5n^2+1} = \frac{1-\frac{2}{n}}{5+\frac{1}{n^2}} \quad \implies \quad \lim_{n \rightarrow \infty} a_n = \frac{\lim_{n \rightarrow \infty} (1-\frac{2}{n})}{\lim_{n \rightarrow \infty} (5+\frac{1}{n^2})} = \frac{1-0}{5+0} = \frac{1}{5}$$

$$4. \quad a_n = \frac{n^4-2n^3+1}{5n^2+1} = \frac{n^2(1-\frac{2}{n}+\frac{1}{n^4})}{5+\frac{1}{n^2}} = n^2 \cdot b_n \quad \text{mit} \quad \lim_{n \rightarrow \infty} b_n = \lim_{n \rightarrow \infty} \frac{1-\frac{2}{n}+\frac{1}{n^4}}{5+\frac{1}{n^2}} = \frac{1}{5}$$

Mit Rechenregel (5) folgt $\lim_{n \rightarrow \infty} a_n = \infty$.

Weiter kann man Aussagen über die Konvergenz einer Folge machen, wenn man die Folge auf Monotonie und Beschränktheit untersucht.

Definition Eine Folge $(a_n)_{n \in \mathbb{N}}$ heisst

- *monoton wachsend* bzw. *fallend*, wenn gilt

$$a_{n+1} \geq a_n \quad \text{bzw.} \quad a_{n+1} \leq a_n \quad \text{für alle } n \text{ in } \mathbb{N},$$

- *beschränkt*, wenn es eine reelle Zahl K gibt mit

$$|a_n| \leq K \quad \text{für alle } n \text{ in } \mathbb{N}.$$

Die Zahl K nennt man eine *Schranke* für die Folge.

Beispiele

1. Die Folge $a_n = \frac{1}{n}$ ist (streng) monoton fallend.
2. Die Folge $a_n = \sqrt{n} - 10$ ist (streng) monoton wachsend.
3. Die Folge $a_n = \frac{1}{n}$ ist beschränkt.
4. Die Folge $a_n = (-1)^n$ ist beschränkt.
5. Die Folge $a_n = \sqrt{n} - 10$ ist nicht beschränkt.

Satz 3.2 Ist eine Zahlenfolge konvergent, dann ist sie beschränkt.

Gleichbedeutend mit Satz 3.2 ist die Aussage, dass jede nicht beschränkte Folge nicht konvergent (d.h. divergent) ist, wie das zum Beispiel bei der Folge $a_n = \sqrt{n} - 10$ der Fall ist. Die Umkehrung ist allerdings falsch. Es gibt Folgen, welche beschränkt, jedoch nicht konvergent sind. Die Folge $a_n = (-1)^n$ ist eine solche Folge.

Satz 3.3 Ist eine Zahlenfolge beschränkt und monoton wachsend oder fallend, dann ist sie konvergent.

Satz 2.3 ist vor allem dann nützlich, wenn der Grenzwert einer Folge unbekannt ist oder wenn die Folge rekursiv definiert ist (s. unten).

Schliesslich gibt es noch eine wichtige Regel zum Merken:

Exponentiell ist stärker als polynomial.

Es gilt also beispielsweise

$$\lim_{n \rightarrow \infty} \left(\frac{n^2 + 2n}{2^n} \right) = 0 \quad \text{und} \quad \lim_{n \rightarrow \infty} \left(\frac{3^n}{n^{2025}} \right) = \infty.$$

Dies kann man mit Hilfe der Regeln von Bernoulli-de l'Hôpital herleiten (vgl. Kapitel 4).

Rekursiv definierte Folgen

Bis jetzt haben wir Zahlenfolgen betrachtet, deren Folgenglieder *explizit* durch eine Formel, abhängig von n , gegeben waren. Man kann Folgen jedoch auch *rekursiv* definieren, das heisst, man definiert das n -te Folgenglied mit Hilfe von vorangehenden Folgengliedern. Dabei darf nicht vergessen werden, das erste Glied explizit anzugeben.

Beispiele

1. Sei $a_{n+1} = \frac{1}{2} \left(a_n + \frac{2}{a_n} \right)$ für $n \geq 1$ und $a_1 = 1$. Durch Einsetzen erhält man

$$a_2 = 1,5, \quad a_3 = 1,416666\dots, \quad a_4 = 1,414215\dots, \quad a_5 = 1,414213\dots$$

Diese Folge ist monoton fallend für $n \geq 2$ und beschränkt (es gilt $|a_n| = a_n \leq 1,5$). Nach Satz 3.3 ist die Folge also konvergent. Sei a der Grenzwert der Folge. Lassen wir nun n gegen ∞ gehen auf beiden Seiten der Gleichung für a_{n+1} , dann erhalten wir mit Hilfe der Rechenregeln (Satz 3.1) und wegen $\lim_{n \rightarrow \infty} a_{n+1} = \lim_{n \rightarrow \infty} a_n = a$, dass

$$a = \frac{1}{2} \left(a + \frac{2}{a} \right) \implies$$

2. Wir wollen die durchschnittliche Population einer Bakterienkultur in einem Experiment durch eine Folge $(a_n)_{n \geq 0}$ beschreiben und zwar so, dass a_n die (ungefähre) Anzahl von Bakterien nach n Tagen darstellt. Wir nehmen an, dass am Anfang 5000 Bakterien vorhanden sind, die sich mit einer täglichen Rate von 4 % vermehren, das heisst, jeden Tag kommen 4 % der Population vom Vortag hinzu. Ausserdem sterben durch äussere Einflüsse täglich 100 Bakterien. Damit hängt die Anzahl a_{n+1} von Bakterien nach $n + 1$ Tagen von der Anzahl a_n von Bakterien am Vortag wie folgt ab:

$$a_{n+1} = 1,04 \cdot a_n - 100$$

Berechnen wir die ersten Folgenglieder:

$$\begin{aligned} a_0 &= 5000 \\ a_1 &= 1,04 \cdot 5000 - 100 = 5100 \\ a_2 &= \end{aligned}$$

Und weiter

$$\begin{aligned} a_3 &= 1,04 \cdot (1,04^2 \cdot 5000 - 100 \cdot (1 + 1,04)) - 100 \\ &= 1,04^3 \cdot 5000 - 100 \cdot (1 + 1,04 + 1,04^2) \end{aligned}$$

Allgemein gilt also

$$a_n = 1,04^n \cdot 5000 - 100 \cdot (1 + 1,04 + \cdots + 1,04^{n-1}).$$

Dies ist nun immerhin eine explizite Formel für a_n , mit der wir uns im Moment begnügen müssen. Doch im nächsten Abschnitt (Seite 42) werden wir sehen (oder Sie wissen es aus der Schule), dass die Summe $1 + 1,04 + \cdots + 1,04^{n-1}$ eine geometrische Reihe ist, die wir noch einmal vereinfachen können.

3.2 Reihen

Sei $a_0, a_1, a_2, \dots$ eine beliebige Folge. Wir bilden daraus eine neue Folge durch sukzessives Aufsummieren:

$$\begin{aligned} s_0 &= a_0 \\ s_1 &= a_0 + a_1 = s_0 + a_1 \\ s_2 &= a_0 + a_1 + a_2 = s_1 + a_2 \\ &\vdots \\ s_n &= a_0 + a_1 + \cdots + a_n = s_{n-1} + a_n \end{aligned}$$

Diese Folge $s_0, s_1, s_2, \dots$ der *Partialsummen* kann also explizit oder rekursiv definiert werden. Explizit kann sie elegant mit Hilfe des Summenzeichens geschrieben werden:

$$s_n = a_0 + a_1 + a_2 + \cdots + a_n = \sum_{k=0}^n a_k$$

Definition Die *unendliche Reihe*

$$\sum_{k=0}^{\infty} a_k$$

ist *konvergent*, falls die Folge $(s_n)_{n \in \mathbb{N}}$ der Partialsummen konvergiert. Man schreibt kurz

$$\sum_{k=0}^{\infty} a_k = s, \quad \text{falls } s = \lim_{n \rightarrow \infty} s_n = \lim_{n \rightarrow \infty} \sum_{k=0}^n a_k.$$

Andernfalls heisst die unendliche Reihe *divergent*.

Beispiele

1. Die *harmonische Reihe* $\sum_{k=1}^{\infty} \frac{1}{k}$ ist divergent.

In der Partialsumme s_n fassen wir die Summanden wie folgt zusammen:

$$s_n = 1 + \frac{1}{2} + \left(\frac{1}{3} + \frac{1}{4}\right) + \left(\frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8}\right) + \left(\frac{1}{9} + \frac{1}{10} + \cdots + \frac{1}{16}\right) + \cdots$$

Die einzelnen Klammerausdrücke sind jeweils $> \frac{1}{2}$. Also ist die Folge der Partialsummen nicht beschränkt und daher (gemäss Satz 3.2) nicht konvergent.

Bemerkung Aus $\lim_{k \rightarrow \infty} a_k = 0$ folgt nicht, dass die Reihe $\sum_{k=0}^{\infty} a_k$ konvergent ist.

2. Die unendliche Reihe $\sum_{k=1}^{\infty} \frac{1}{k^2}$ ist konvergent.

LEONHARD EULER zeigte 1735, dass $\sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{\pi^2}{6} \approx 1,644934068$ (“Basler Problem”).

Bei allgemeinen Reihen ist es schwierig zu entscheiden und beweisen, ob sie konvergent sind oder nicht. Wir betrachten hier deshalb nur noch die geometrischen Reihen, für die es ein einfaches Konvergenzkriterium gibt.

Geometrische Reihen

Sei q eine reelle Zahl. Dann ist $(q^k)_{k \geq 0}$ eine geometrische Folge und die zugehörige unendliche Reihe $\sum_{k=0}^{\infty} q^k$ nennt man *geometrische Reihe*.

Für $q \neq 1$ gilt die folgende Formel für die Partialsumme s_n ,

$$s_n = \frac{1 - q^{n+1}}{1 - q}.$$

Ist $|q| < 1$, dann gilt $\lim_{n \rightarrow \infty} q^n = 0$ (vgl. Seite 38, 3. Beispiel). Damit gilt der folgende Satz.

Satz 3.4 Die geometrische Reihe $\sum_{k=0}^{\infty} q^k$ konvergiert für $|q| < 1$ und es gilt

$$s = \sum_{k=0}^{\infty} q^k = \frac{1}{1 - q}.$$

Beispiel

Am Ende des Abschnitts 3.1 fanden wir für die Anzahl a_n von Bakterien nach n Tagen, dass

$$a_n = 1,04^n \cdot 5000 - 100 \cdot (1 + 1,04 + \cdots + 1,04^{n-1}).$$

Mit obiger Formel können wir nun den Klammerausdruck vereinfachen zu

$$1 + 1,04 + \cdots + 1,04^{n-1} =$$

und wir finden

$$a_n = 1,04^n \cdot 5000 - 100 \cdot \frac{1,04^n - 1}{0,04} = 1,04^n \cdot 5000 - 2500 \cdot (1,04^n - 1) = 2500 \cdot (1,04^n + 1).$$

Wegen $\lim_{n \rightarrow \infty} 1,04^n = \infty$ folgt $\lim_{n \rightarrow \infty} a_n = \infty$, die Bakterienpopulation wächst also unbegrenzt.

Eine geometrische Reihe muss nicht mit dem nullten Glied $q^0 = 1$ beginnen. Manchmal beginnt sie mit $q^1 = q$, manchmal mit q^2 , oder allgemein mit q^m für ein $m \geq 0$. Nehmen wir an, dass $|q| < 1$. Dann gilt

Bemerkung *Es gilt*

$$\sum_{k=m}^{\infty} q^k = \frac{q^m}{1-q},$$

wobei q^m das erste Glied der Reihe ist.

Potenzreihen

Wir haben vorher gesehen, dass die Reihe $\sum_{k=0}^{\infty} x^k$ für alle x mit $|x| < 1$, das heisst, für x im offenen Intervall $I = (-1, 1)$, konvergent ist mit Summe $\frac{1}{1-x}$. Betrachten wir die Funktion $f(x) = \frac{1}{1-x} : I \rightarrow \mathbb{R}$, dann können wir also die Polynome (die Partialsummen)

$$s_n(x) = \sum_{k=0}^n x^k = 1 + x + x^2 + x^3 + \cdots + x^n$$

als Näherungen von $f(x)$ auffassen. Für kleine $|x|$ gilt in erster Näherung $\frac{1}{1-x} \approx 1 + x$, in zweiter Näherung $\frac{1}{1-x} \approx 1 + x + x^2$, usw.

Definition Eine *Potenzreihe* ist eine Reihe der Gestalt

$$\sum_{k=0}^{\infty} a_k x^k,$$

wobei $(a_k)_{k \geq 0}$ eine beliebige Folge ist.

Das zentrale Problem bei Potenzreihen besteht darin, alle x in $\mathbb{R}$ anzugeben, für welche die Reihe konvergiert. Es stellt sich heraus, dass die Menge dieser x stets ein offenes Intervall ist, das heisst, auf diesem Intervall definiert $f(x) = \sum_{k=0}^{\infty} a_k x^k$ eine reelle Funktion. Auf diese Weise erhalten wir wichtige Funktionen. Das wichtigste Beispiel ist die natürliche Exponentialfunktion

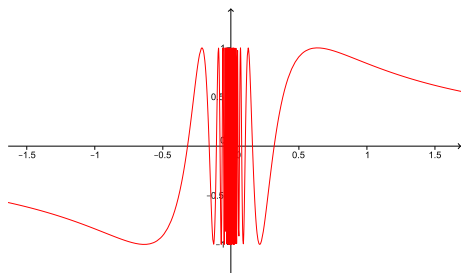
$$f(x) = e^x = 1 + \frac{x}{1!} + \frac{x^2}{2!} + \frac{x^3}{3!} + \cdots = \sum_{k=0}^{\infty} \frac{x^k}{k!},$$

konvergent für alle $x \in \mathbb{R}$, die wir schon in Abschnitt [1.5](#) kennengelernt haben.

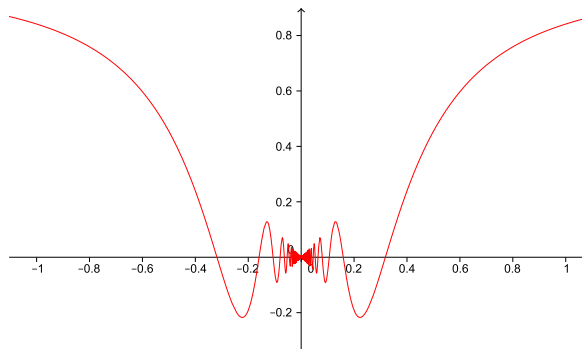
3.3 Grenzwerte bei Funktionen

Betrachten wir die beiden reellen Funktionen:

$$f(x) = \sin\left(\frac{1}{x}\right)$$



$$g(x) = x \sin\left(\frac{1}{x}\right)$$



Beide Funktionen sind für $x_0 = 0$ nicht definiert. Doch während f in der Nähe von 0 zwischen den Werten -1 und 1 oszilliert, nähern sich die Funktionswerte von g dem Wert 0 für x gegen 0. Man sagt, dass g an der Stelle $x_0 = 0$ den Grenzwert $a = 0$ hat. Präzise kann die Näherung von x gegen 0 durch Folgen beschrieben werden.

Definition Sei $f : D \rightarrow \mathbb{R}$ eine Funktion und x_0 eine reelle Zahl, die Grenzwert einer Zahlenfolge $(x_n)_{n \in \mathbb{N}}$ mit $x_n \in D$ ist. Dann ist a der Grenzwert von f an der Stelle x_0 , falls für jede Folge (x_n) mit $x_n \in D$, $x_n \neq x_0$ für alle n , und $\lim_{n \rightarrow \infty} x_n = x_0$ gilt, dass

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

Wir schreiben dann

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

Die reelle Zahl x_0 kann dabei in D liegen, muss aber nicht.

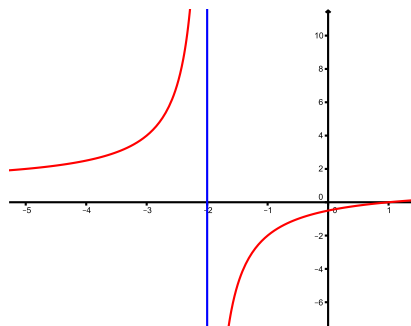
Der Grenzwert einer Funktion ist also durch den Grenzwert von Zahlenfolgen definiert. Wir können deshalb die praktischen Rechenregeln für Folgen (Satz 3.1) anwenden.

Beispiel

Wir betrachten die rationale Funktion

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

definiert auf $D = \mathbb{R} \setminus \{-2, -1\}$.



- $x_0 = -1$: Sei (x_n) eine Folge mit $\lim_{n \rightarrow \infty} x_n = -1$, $x_n \in D$ und $x_n \neq -1$ für alle n .

- $x_0 = -2$: Da für $x \rightarrow -2$ insbesondere $x \neq -1$ ist, gilt

$$\lim_{x \rightarrow -2} f(x) = \lim_{x \rightarrow -2} \frac{(x+1)(x-1)}{(x+1)(x+2)} = \lim_{x \rightarrow -2} \frac{x-1}{x+2}.$$

Ist nun (x_n) eine Folge mit Grenzwert -2 und alle Folgenglieder x_n sind kleiner als -2 , dann strebt $f(x_n)$ gegen $+\infty$ (denn Zähler < 0 , Nenner < 0). Gilt jedoch $x_n > -2$ für alle n , dann strebt $f(x_n)$ gegen $-\infty$ (denn Zähler < 0 , Nenner > 0). Der Grenzwert $\lim_{x \rightarrow -2} f(x)$ existiert also nicht.

Eine Stelle x_0 , in deren unmittelbarer Nähe die Funktionswerte über alle Grenzen hinaus wachsen oder fallen, nennt man eine *Polstelle*. Im Beispiel oben ist $x_0 = -2$ eine Polstelle.

Wie in diesem Beispiel kann sich eine Funktion unterschiedlich verhalten, je nachdem, ob sich x der Stelle x_0 von links oder von rechts nähert.

Definition Für den *linksseitigen Grenzwert* betrachtet man nur Folgen x_n mit $x_n < x_0$, man schreibt

$$\lim_{x \uparrow x_0} f(x).$$

Für den *rechtsseitigen Grenzwert* betrachtet man nur Folgen x_n mit $x_n > x_0$, man schreibt

$$\lim_{x \downarrow x_0} f(x).$$

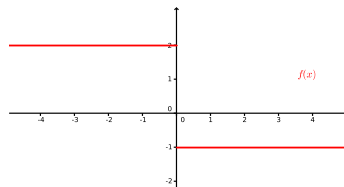
Satz 3.5 Sei $f : D \rightarrow \mathbb{R}$ eine Funktion und a in $\mathbb{R}$. Dann ist $a = \lim_{x \rightarrow x_0} f(x)$ genau dann, wenn gilt

$$\lim_{x \uparrow x_0} f(x) = \lim_{x \downarrow x_0} f(x) = a.$$

Beispiele

1. Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$f(x) = \begin{cases} -1 & \text{falls } x \geq 0 \\ 2 & \text{falls } x < 0 \end{cases}$$



Der Grenzwert $\lim_{x \rightarrow 0} f(x)$ existiert also nicht. Der Graph macht in $x_0 = 0$ einen Sprung.

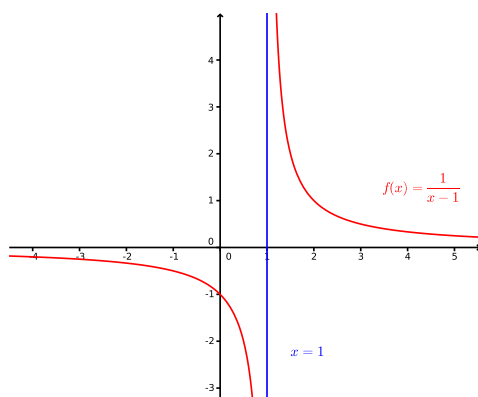
2. Sei $f : \mathbb{R} \setminus \{1\} \rightarrow \mathbb{R}$, $f(x) = \frac{1}{x-1}$.

Für $x_0 \neq 1$ gilt

$$\lim_{x \rightarrow x_0} f(x) = \lim_{x \rightarrow x_0} \frac{1}{x-1} = \frac{1}{\lim_{x \rightarrow x_0} x - 1} = \frac{1}{x_0 - 1} = f(x_0).$$

Und $x_0 = 1$ ist eine Polstelle, denn

$$\lim_{x \downarrow 1} \frac{1}{x-1} = \infty \quad \text{und} \quad \lim_{x \uparrow 1} \frac{1}{x-1} = -\infty$$

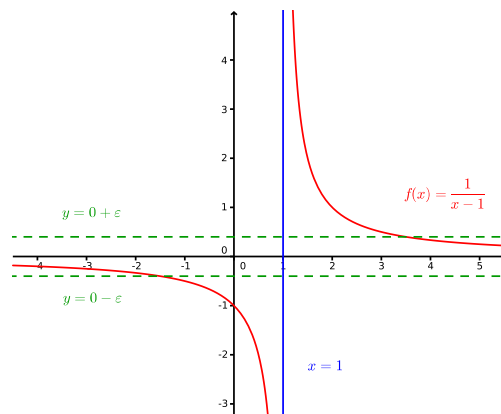


Analog zum Grenzwert von f an einer Stelle x_0 in $\mathbb{R}$ definiert man den Grenzwert von f für $x \rightarrow \infty$ und für $x \rightarrow -\infty$. Betrachtet man den Graphen von f , so bedeutet $\lim_{x \rightarrow \infty} f(x) = a$, dass die horizontale Gerade $y = a$ eine *Asymptote* für $x \rightarrow \infty$ ist. Genauer heisst dies, dass der Graph für genügend grosse x im horizontalen Streifen zwischen den Geraden $y = a + \varepsilon$ und $y = a - \varepsilon$ liegt.

Beispiele

1. Sei $f : \mathbb{R} \setminus \{1\} \rightarrow \mathbb{R}$, $f(x) = \frac{1}{x-1}$ wie vorher. Dann gilt

$$\lim_{x \rightarrow \infty} \frac{1}{x-1} = \lim_{x \rightarrow -\infty} \frac{1}{x-1} = 0.$$



2. Sei $f(x) = \frac{x^2 - 1}{x^2 + 3x + 2}$ die Funktion von Seite 44.

3.4 Stetige Funktionen

Für die Funktion $f(x) = \frac{1}{x-1}$ von vorher gilt also für alle $x_0 \neq 1$, dass

$$\lim_{x \rightarrow x_0} f(x) = \frac{1}{x_0 - 1} = f(x_0).$$

Tatsächlich strebt in den meisten Fällen $f(x)$ gegen $f(x_0)$ für $x \rightarrow x_0$. Man nennt in diesen Fällen die Funktion f *stetig in x_0* .

Definition Sei $f : D \rightarrow \mathbb{R}$ eine Funktion und x_0 in D . Dann heisst f *stetig in x_0* , wenn gilt

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

Die Funktion heisst *stetig in D* , wenn f in jedem Punkt $x_0 \in D$ stetig ist.

Anschaulich bedeutet die Stetigkeit von f in einem Punkt x_0 , dass der Wert $f(x)$ nahe bei $f(x_0)$ ist, sobald x genügend nahe bei x_0 ist.

Beispiele

1. $f : \mathbb{R} \rightarrow \mathbb{R}, f(x) = x^2$.

Sei $x_0 \in \mathbb{R}$ und $(x_n)_{n \geq 1}$ eine beliebige Folge mit $\lim_{n \rightarrow \infty} x_n = x_0$. Dann gilt

$$\lim_{x \rightarrow x_0} f(x) = \lim_{n \rightarrow \infty} f(x_n) = \lim_{n \rightarrow \infty} x_n^2 = \left(\lim_{n \rightarrow \infty} x_n \right)^2 = x_0^2 = f(x_0),$$

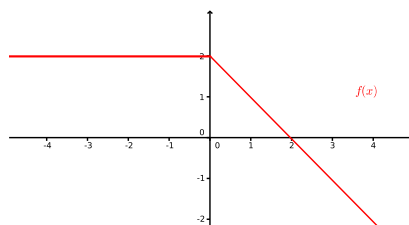
wobei wir für die dritte Gleichung die Rechenregel (3) über Zahlenfolgen (Satz 3.1) benutzt haben. Damit ist f stetig in x_0 , und da x_0 beliebig war, ist f stetig in ganz $\mathbb{R}$.

2. Betrachten wir noch einmal $f : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch $f(x) = -1$ für $x \geq 0$ und $f(x) = 2$ für $x < 0$ (vgl. Seite 45).

In allen $x \neq 0$ ist f stetig. In $x_0 = 0$ hingegen ist f nicht stetig, denn der Grenzwert $\lim_{x \rightarrow 0} f(x)$ existiert nicht; der Graph macht in $x_0 = 0$ einen *Sprung*.

3. Macht der Graph an einer Stelle nur einen *Knick*, dann ist die Funktion an dieser Stelle immer noch stetig. Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$f(x) = \begin{cases} -x + 2 & \text{falls } x \geq 0 \\ 2 & \text{falls } x < 0 \end{cases}$$



In allen $x \neq 0$ ist f (offensichtlich) stetig. Wir überprüfen nun, dass f auch in $x_0 = 0$ stetig ist:

Eine Funktion ist also stetig, wenn man ihren Graphen ohne abzusetzen (d.h. der Graph macht keine Sprünge) “anständig” zeichnen kann, das heisst, wenn der Graph eine ununterbrochene Linie ist. Alle elementaren Funktionen haben diese Eigenschaft.

Satz 3.6 *Alle elementaren Funktionen (Polynome, rationale Funktionen, Potenzfunktionen, Exponential- und Logarithmusfunktionen, trigonometrische Funktionen) sind stetig in ihrem Definitionsbereich.*

Aus diesen elementaren Funktionen kann man eine Vielzahl von weiteren stetigen Funktionen konstruieren.

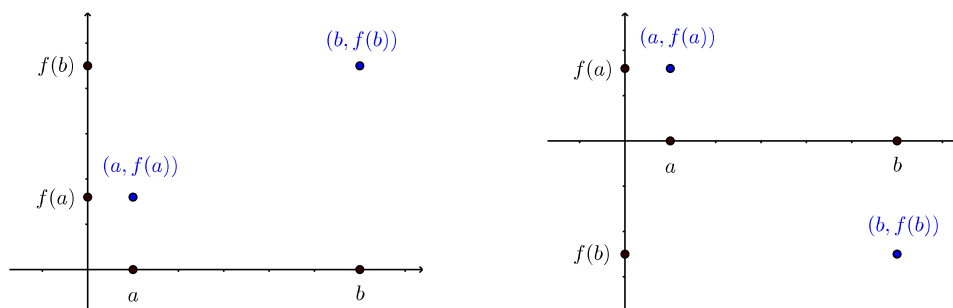
Satz 3.7 *Seien $f, g : D \rightarrow \mathbb{R}$ zwei in einem Punkt $x_0 \in D$ stetige Funktionen. Dann gilt:*

- (1) *Die Summe $f + g$ ist stetig in x_0 . Hierbei ist die Funktion $f + g : D \rightarrow \mathbb{R}$ definiert durch $(f + g)(x) = f(x) + g(x)$.*
- (2) *Das Produkt $f \cdot g$ ist stetig in x_0 . Hierbei ist die Funktion $f \cdot g : D \rightarrow \mathbb{R}$ definiert durch $(f \cdot g)(x) = f(x) \cdot g(x)$.*
- (3) *Ist $g(x_0) \neq 0$, so ist auch der Quotient $\frac{f}{g}$ stetig in x_0 . Hierbei ist die Funktion $\frac{f}{g}$ definiert durch $\frac{f}{g}(x) = \frac{f(x)}{g(x)}$.*
- (4) *Seien $f : D \rightarrow \mathbb{R}$ und $g : D' \rightarrow \mathbb{R}$ zwei Funktionen mit $f(D) \subset D'$. Sei f stetig in x_0 und g stetig in $f(x_0)$. Dann ist die Komposition $g \circ f$ stetig in x_0 .*

Beispiel

Die beiden Funktionen $f(x) = \sin(\frac{1}{x})$ und $g(x) = x \sin(\frac{1}{x})$ von Seite 44 sind gemäss den Sätzen 3.6 und 3.7 (2) & (4) stetig in $D = \mathbb{R} \setminus \{0\}$.

Wir haben oben erwähnt, dass der Graph einer stetigen Funktion eine ununterbrochene Linie ist. Dies bedeutet, dass die Funktion alle Werte zwischen zwei Funktionswerten annimmt (sofern sie auf dem ganzen Intervall dazwischen definiert ist). Auf dieser Eigenschaft beruht der folgende Nullstellensatz.



Satz 3.8 (Nullstellensatz) *Ist f stetig auf dem Intervall $[a, b]$ und haben $f(a)$ und $f(b)$ unterschiedliche Vorzeichen, so hat f eine Nullstelle zwischen a und b , das heisst, es gibt ein $x_0 \in (a, b)$ mit $f(x_0) = 0$.*

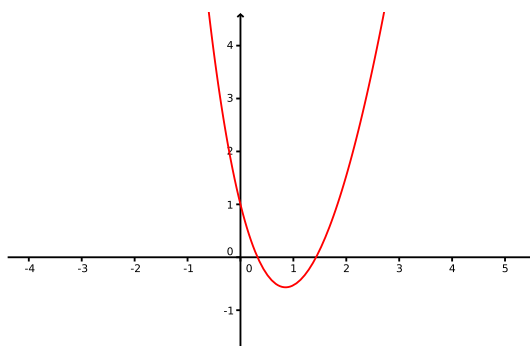
Wie schon in der Einleitung dieses Kapitels bemerkt (Seite 35), ist die Stetigkeit von f zentral in diesem Satz.

Beispiel

Nach den Sätzen 3.6 und 3.7 ist die Funktion $f(x) = 4e^{-x} + x^2 - 3$ stetig für alle x in $\mathbb{R}$. Es gilt

$$f(0) = 1 > 0 \quad \text{und} \quad f(1) = -0,528 < 0.$$

Gemäss Nullstellensatz hat f also mindestens eine Nullstelle im (offenen) Intervall $(0, 1)$.



Wie wir eine Näherung für diese Nullstelle finden können, werden wir in Abschnitt 4.6 sehen.

4 Differentiation

Alle elementaren Funktionen (und damit auch Summen, Produkte und Quotienten davon) sind in ihrem Definitionsbereich nicht nur stetig, sondern auch differenzierbar. Mit Hilfe der Ableitung können wir ihre Maxima und Minima bestimmen, die Nullstellen näherungsweise berechnen sowie komplizierte Funktionen durch einfachere Funktionen näherungsweise beschreiben.

4.1 Die Ableitung einer Funktion

Der Begriff der Ableitung ist aus einem praktischen Problem entstanden. Nehmen wir an, wir fahren mit dem Zug von Basel nach Chur. Für die rund 200 km lange Strecke benötigen wir 2 Stunden und 19 Minuten, das heisst, 2,32 Stunden. Wie schnell sind wir gefahren, das heisst, wie gross war unsere Geschwindigkeit?

Nun, Geschwindigkeit ist gleich Weg durch Zeit, genauer: zurückgelegter Weg durch verstrichene Zeitspanne. Für unsere Geschwindigkeit folgt also

$$\text{Geschwindigkeit} = \frac{200 \text{ km}}{2,32 \text{ h}} \approx 86 \frac{\text{km}}{\text{h}}.$$

Dies ist unsere mittlere oder durchschnittliche Geschwindigkeit.

Sei allgemein $s(t)$ der zurückgelegte Weg zum Zeitpunkt t . Gehen wir vom zurückgelegten Weg $s(t_0)$ zu einem Zeitpunkt t_0 aus, dann können wir die *mittlere Geschwindigkeit* während der folgenden Zeitspanne h (d.h. von t_0 bis $t_0 + h$) berechnen als

$$\frac{s(t_0 + h) - s(t_0)}{h}.$$

Nun wollen wir aber wissen, wie gross unsere Geschwindigkeit zu einem bestimmten Zeitpunkt war. Wie kann die momentane Geschwindigkeit zu einem genauen Zeitpunkt t_0 berechnet werden? Die Idee ist, die momentane Geschwindigkeit durch mittlere Geschwindigkeiten anzunähern.

Beispiel

Nach dem langen Halt im Bahnhof Zürich fährt der Zug wieder los. Dieser Anfahrvorgang sei durch die Wegfunktion $s(t) = t^2$ (in Metern) beschrieben. Wie gross ist dann unsere Geschwindigkeit nach genau 3 Sekunden nach dem Anfahren? Wir betrachten also den Zeitpunkt $t_0 = 3$ und lassen die Zeitspanne h (in Sekunden) immer kleiner werden.

Zeitspanne	h	1	0,1	0,01	0,001
mittlere Geschwindigkeit	$\frac{s(t_0+h)-s(t_0)}{h}$	7	6,1	6,01	6,001

Wir erkennen, dass sich für $h \rightarrow 0$ die momentane Geschwindigkeit $6 \frac{\text{m}}{\text{s}}$ nähert.

Allgemein berechnet sich die *momentane Geschwindigkeit* zur Zeit t_0 durch

$$\lim_{h \rightarrow 0} \frac{s(t_0 + h) - s(t_0)}{h}.$$

Diese Idee der Näherung geht auf ISAAC NEWTON (1643 – 1727) zurück. Doch damals kannte man noch keine präzise Definition des Grenzwerts. Diese wurde erst im 19. Jahrhundert eingeführt. Heute nennen wir diesen Grenzwert die *Ableitung* der Funktion $s(t)$.

Definition Sei $f : D \rightarrow \mathbb{R}$ eine reelle Funktion. Dann heisst f an der Stelle $x_0 \in D$ *differenzierbar*, wenn der Grenzwert

$$\lim_{h \rightarrow 0} \frac{f(x_0 + h) - f(x_0)}{h}$$

existiert. Er heisst *Ableitung* (oder *Differentialquotient*) von f in x_0 und wird mit $f'(x_0)$ bezeichnet. Der Ausdruck

$$\frac{f(x_0 + h) - f(x_0)}{h}$$

wird *Differenzenquotient* genannt.

Andere Notationen: Für die Ableitung sind auch die folgenden Schreibweisen üblich:

$$f'(x_0) = \frac{df}{dx}(x_0) = \left. \frac{df}{dx} \right|_{x=x_0}$$

Ist die (Weg-)Funktion $s(t)$ abhängig von der Zeit t , dann wird speziell in der Physik die Ableitung mit $\dot{s}(t)$ bezeichnet.

Anstelle von h ist auch Δx üblich, bzw. $x = x_0 + \Delta x$. Damit ist

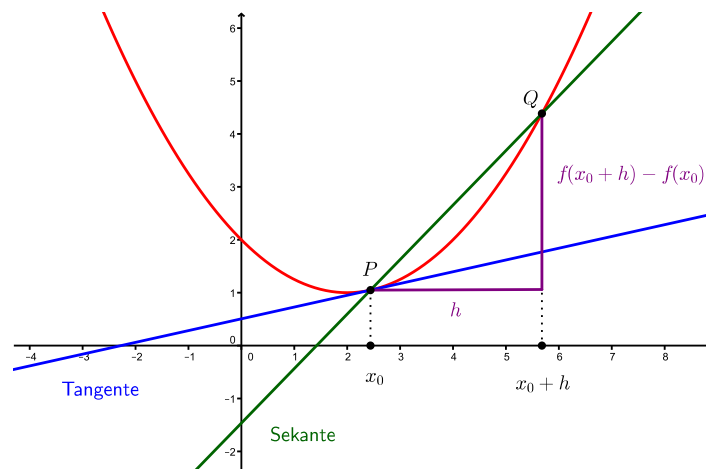
$$f(x_0 + h) - f(x_0) = f(x_0 + \Delta x) - f(x_0) = f(x) - f(x_0) = \Delta f$$

und man erhält für $f'(x_0)$ die äquivalenten Schreibweisen

$$\lim_{h \rightarrow 0} \frac{f(x_0 + h) - f(x_0)}{h} = \lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x} = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} = \lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x}.$$

Geometrische Deutung

Der Differenzenquotient ist gleich der Steigung der Sekante durch die Punkte $P = (x_0, f(x_0))$ und $Q = (x_0 + h, f(x_0 + h))$ auf dem Graphen von f . Für $h \rightarrow 0$ (d.h. lässt man Q gegen P wandern) geht die Sekante in die Tangente an den Graphen von f in x_0 über. Die Ableitung $f'(x_0)$ beschreibt also die Steigung der Tangente an den Funktionsgraphen an der Stelle x_0 (was auch als Steigung des Funktionsgraphen in x_0 bezeichnet wird).



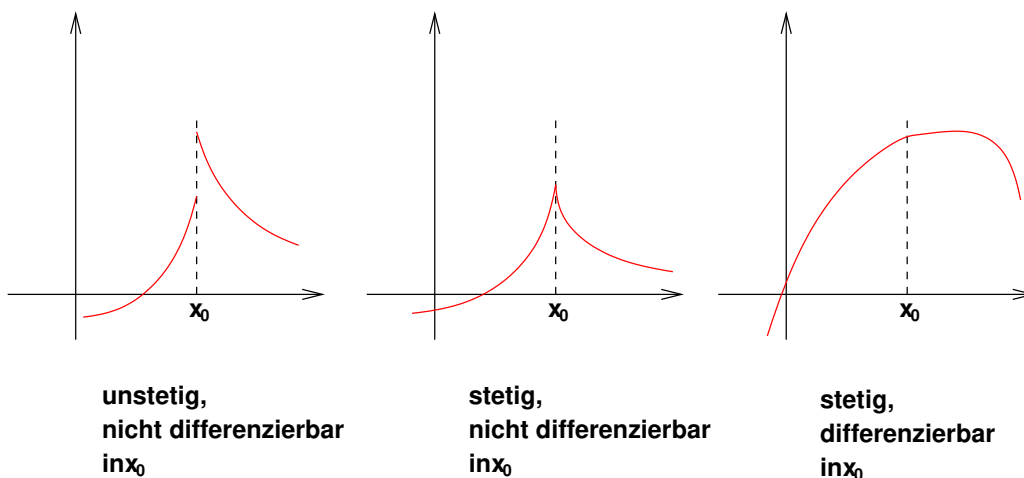
Auf diese Interpretation der Ableitung stiess unabhängig von Newton GOTTFRIED WILHELM LEIBNIZ (1646 – 1716), als er sich fragte, wie man die Steigung einer Kurve (insbesondere eines Funktionsgraphen) in einem Punkt erklären kann.

Definition Eine Funktion $f : D \rightarrow \mathbb{R}$ heisst *differenzierbar* (in D), falls f in allen $x \in D$ differenzierbar ist. Die Funktion $f' : D \rightarrow \mathbb{R}$, die jedem $x \in D$ die Zahl $f'(x)$ zuordnet, heisst *Ableitung* von f .

Anschaulich gesehen ist eine Funktion differenzierbar, wenn sich an jedem Punkt des Graphen in eindeutiger Weise eine Tangente anlegen lässt. Eine Funktion ist also insbesondere an einer Stelle nicht differenzierbar, wenn sie dort einen Sprung macht, eine Polstelle hat oder oszilliert; das heisst, wenn sie dort nicht stetig ist. Es kann aber auch Stellen geben, an welchen die Funktion wohl stetig ist, jedoch nicht differenzierbar, beispielsweise wenn die Funktion einen Knick macht.

Bemerkung f differenzierbar in $x_0 \implies f$ stetig in x_0

Die umgekehrte Richtung “ $\Leftarrow$ ” ist hingegen falsch.



Beispiele

1. Sei $f(x) = c$ eine konstante Funktion. Dann ist $f(x_0 + h) - f(x_0) = c - c = 0$ und damit $f'(x_0) = 0$ für alle $x_0 \in \mathbb{R}$. Dies erkennt man auch direkt am Graphen von f .
2. Sei $s(t) = t^2$ die Wegfunktion vom Beispiel auf Seite 50. Wie gross ist die (momentane) Geschwindigkeit zur Zeit t_0 ?

Zum Zeitpunkt $t_0 = 3$ s beträgt die Geschwindigkeit also tatsächlich $s'(3) = 6 \frac{\text{m}}{\text{s}}$.

Die folgende allgemeinere Ableitungsregel werden wir auf Seite 55 beweisen.

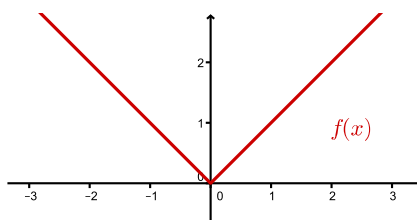
Bemerkung Für jede reelle Zahl r gilt

$$f(x) = x^r \quad \implies \quad f'(x) = r x^{r-1}$$

Insbesondere gilt für $f(x) = \sqrt{x} = x^{\frac{1}{2}}$, dass

$$f'(x) = \frac{1}{2} x^{-\frac{1}{2}} = \frac{1}{2\sqrt{x}}.$$

4. Sei $f(x) = |x|$.



Diese Funktion ist differenzierbar in allen $x \in \mathbb{R}$, ausser in $x_0 = 0$. Man erkennt dies direkt am Graphen von f , denn in $x_0 = 0$ macht er einen Knick und es ist unmöglich, hier eine Tangente auf eindeutige Weise anzulegen. Tatsächlich existiert der Differentialquotient in $x_0 = 0$ nicht, denn

$$\lim_{h \downarrow 0} \frac{f(x_0 + h) - f(x_0)}{h} =$$

aber

$$\lim_{h \uparrow 0} \frac{f(x_0 + h) - f(x_0)}{h} =$$

Um für eine konkrete Funktion f die Ableitung f' zu bestimmen, sind die folgenden Ableitungsregeln sehr nützlich.

Satz 4.1 (Ableitungsregeln) (a) Seien f, g differenzierbar. Dann sind auch $f + g$, λf für $\lambda \in \mathbb{R}$, $f \cdot g$ und $\frac{f}{g}$ (überall wo $g(x) \neq 0$) differenzierbar.

(i) Es gilt

$$(f + g)' = f' + g' \quad \text{und} \quad (\lambda f)' = \lambda f' \quad \text{für alle } \lambda \in \mathbb{R}.$$

(ii) Es gilt die Produktregel

$$(f \cdot g)' = f' \cdot g + f \cdot g'.$$

(iii) Es gilt die Quotientenregel

$$\left(\frac{f}{g}\right)' = \frac{f' \cdot g - f \cdot g'}{g^2}.$$

(b) Ist $f(x) = \sum_{k=0}^{\infty} a_k x^k$ eine auf dem (offenen) Intervall I konvergente Potenzreihe, dann ist f auf I differenzierbar und es gilt

$$f'(x) = \sum_{k=1}^{\infty} k a_k x^{k-1} \quad \text{für alle } x \in I.$$

Die Regeln von Satz 4.1 (a) sollten aus der Schule bekannt sein. Wir geben hier deshalb nur Beispiele zu Satz 4.1 (b).

Beispiele

1. Sei $f(x) = e^x = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \frac{x^4}{4!} + \dots$. Mit Hilfe von Satz 4.1 (b) folgt

Also gilt $(e^x)' = e^x$. Bis auf Vielfache ist $f(x) = e^x$ die einzige reelle Funktion f mit $f' = f$.

2. Auch $\cos x$ und $\sin x$ können als Potenzreihen dargestellt werden. Aus der Potenzreihe von e^{ix} und der Eulerschen Identität $e^{ix} = \cos x + i \sin x$ folgt, dass

$$\begin{aligned} \cos x &= 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \dots + \dots = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} \\ \sin x &= x - \frac{x^3}{3!} + \frac{x^5}{5!} - \dots + \dots = \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!} . \end{aligned}$$

Für die Ableitungen können wir nach Satz 4.1 (b) gliedweise ableiten. Wir erhalten

und analog

$$(\sin x)' = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \dots + \dots = \cos x .$$

Es gilt also:

$$(\cos x)' = -\sin x \quad \text{und} \quad (\sin x)' = \cos x$$

Als nächstes betrachten wir zusammengesetzte Funktionen (wie in Kapitel 1, Seite 9, definiert). Sei $h = g \circ f$, wobei $f : D \rightarrow \mathbb{R}$ und $g : \tilde{D} \rightarrow \mathbb{R}$ mit $f(D) \subseteq \tilde{D}$. Das heisst, es gilt $h(x) = (g \circ f)(x) = g(f(x))$ für alle $x \in D$.

Satz 4.2 (Kettenregel) Ist f differenzierbar in x_0 und g differenzierbar in $y_0 = f(x_0)$, so ist die Komposition $g \circ f$ differenzierbar in x_0 und es gilt

$$(g \circ f)'(x_0) = g'(f(x_0)) \cdot f'(x_0) .$$

Da $(g \circ f)(x) = g(f(x))$, nennt man g die *äussere Funktion* und f die *innere Funktion*. Die Kettenregel

$$(g \circ f)'(x_0) = g'(f(x_0)) \cdot f'(x_0)$$

kann man sich (in Kurzform) also so merken:

$$\text{äussere Ableitung} \cdot \text{innere Ableitung}$$

Genauer bedeutet dies: Die Ableitung der äusseren Funktion an der Stelle der inneren Funktion mal die Ableitung der inneren Funktion.

Beispiele

1. Sei $h(x) = (\sin x)^5 = \sin^5 x$.

2. Sei $h(x) = e^{\cos(3x+\pi)}$.

Satz 4.3 (Ableitung einer Umkehrfunktion) Sei $f : D \rightarrow \mathbb{R}$ umkehrbar und differenzierbar in x_0 mit $f'(x_0) \neq 0$. Dann ist die Umkehrfunktion $g = f^{-1}$ differenzierbar in $y_0 = f(x_0)$ und es gilt

$$g'(y_0) = \frac{1}{f'(x_0)} = \frac{1}{f'(g(y_0))}.$$

Beispiele

1. Die Funktion $g(y) = \ln(y)$ ist die Umkehrfunktion von $f(x) = e^x$.

Es gilt also für alle $y > 0$,

$$(\ln(y))' = \frac{1}{y}.$$

Damit und mit Hilfe der Kettenregel können wir die Ableitungsregel

$$(x^r)' = r x^{r-1}$$

für alle reellen Zahlen r von Seite 53 überprüfen:

2. Die Funktion $f(x) = \sin x : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow [-1, 1]$ ist umkehrbar und $f'(x) = \cos x$ ist $\neq 0$ für $-\frac{\pi}{2} < x < \frac{\pi}{2}$. Die Umkehrfunktion $g(y) = \arcsin y$ ist deshalb differenzierbar für $-1 < y < 1$ (denn $\sin(\pm\frac{\pi}{2}) = \pm 1$) und es gilt

Wegen $\cos^2 x + \sin^2 x = 1$ ist $\cos x = \sqrt{1 - \sin^2 x} = \sqrt{1 - (\sin x)^2}$. Damit erhalten wir

Es gilt also für $-1 < y < 1$,

$$(\arcsin y)' = \frac{1}{\sqrt{1 - y^2}}.$$

4.2 Extremalstellen

Sei $f : D \rightarrow \mathbb{R}$ eine differenzierbare Funktion. Falls die Ableitung $f' : D \rightarrow \mathbb{R}$ wieder differenzierbar ist, können wir die Ableitung von f' bilden und erhalten die zweite Ableitung f'' von f , usw. Eine Funktion f , die man auf diese Weise n -mal ableiten kann, heisst *n -mal differenzierbar*. Die *n -te Ableitung* bezeichnet man mit

$$f^{(n)}(x).$$

Alternative Schreibweisen sind

$$f^{(n)}(x) = f''\dots'(x) = \frac{d^n f(x)}{dx^n} = \frac{d^n}{dx^n} f(x).$$

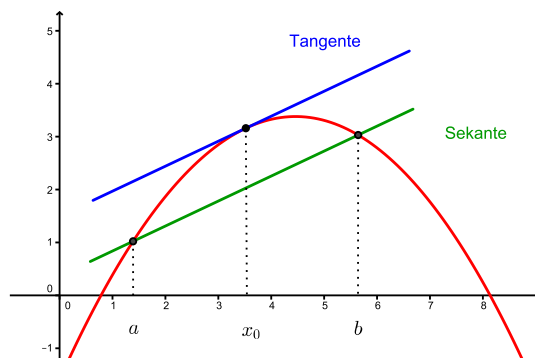
Das Ziel dieses Abschnitts ist, aus den Ableitungen einer Funktion Rückschlüsse auf deren Verlauf zu ziehen.

Das Monotonieverhalten einer Funktion kann direkt an der Ableitung der Funktion abgelesen werden. Dieser Zusammenhang basiert auf dem folgenden wichtigen Satz.

Satz 4.4 (Mittelwertsatz) *Sei f eine reelle Funktion, die in einem Intervall $[a, b]$ differenzierbar ist. Dann gibt es einen Punkt $x_0 \in [a, b]$, so dass*

$$f(b) - f(a) = f'(x_0) \cdot (b - a).$$

Der Mittelwertsatz sagt aus, dass es zu jeder Sekante eine parallele Tangente an den Funktionsgraphen gibt.

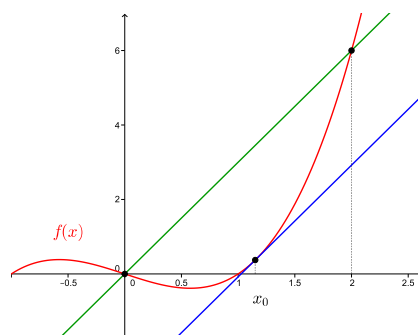


$$\text{Steigung der Sekante} = \frac{f(b) - f(a)}{b - a} = f'(x_0) = \text{Steigung der Tangente in } x_0$$

Beispiel

Wir betrachten $f(x) = x^3 - x$ im Intervall $[a, b] = [0, 2]$. Es gilt

Nach dem Mittelwertsatz gibt es also (mindestens) ein $x_0 \in [0, 2]$ mit $f'(x_0) = 3$. Nun ist



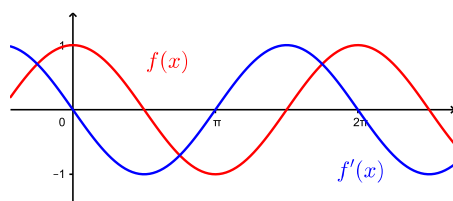
Da x_0 im Intervall $[0, 2]$ liegt, ist also $x_0 = \frac{2}{\sqrt{3}} \approx 1,15$ der gesuchte Punkt.

Satz 4.5 Sei f differenzierbar in $[a, b]$. Dann gilt

- (i) f monoton wachsend in $[a, b] \iff f'(x) \geq 0$ für alle $x \in [a, b]$
- (ii) f monoton fallend in $[a, b] \iff f'(x) \leq 0$ für alle $x \in [a, b]$
- (iii) f konstant in $[a, b] \iff f'(x) = 0$ für alle $x \in [a, b]$

Satz 4.5 ist nachvollziehbar, wenn $f'(x)$ als “momentane Änderungsrate” von f interpretiert wird. Der Mittelwertsatz geht in die Beweise der Richtungen “ $\Leftarrow$ ” ein.

Beispiel



Sei nun $f : [a, b] \rightarrow \mathbb{R}$ eine Funktion. Wir wollen die sogenannten Extremalstellen (bzw. die Extrema) bestimmen.

Die Punkte x im offenen Intervall $(a, b) = \{x \in \mathbb{R} \mid a < x < b\}$ nennen wir die *inneren* Punkte und die Punkte $x = a$ und $x = b$ heissen *Randpunkte* von $[a, b]$.

Definition Ein Punkt $x_0 \in [a, b]$ heisst

- *globale Maximalstelle* von f in $[a, b]$, falls gilt

$$f(x) \leq f(x_0) \quad \text{für alle } x \in [a, b];$$

lokale Maximalstelle, falls x_0 ein innerer Punkt ist und

$$f(x) \leq f(x_0) \quad \text{für alle } x \text{ in einer Umgebung von } x_0.$$

Man nennt $f(x_0)$ ein *globales* (bzw. *lokales*) *Maximum* von f in $[a, b]$.

- *globale Minimalstelle* von f in $[a, b]$, falls gilt

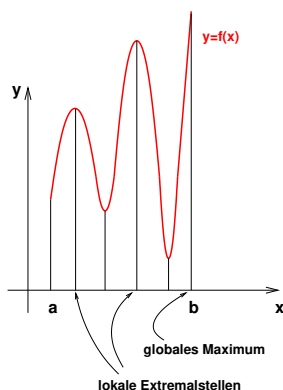
$$f(x) \geq f(x_0) \quad \text{für alle } x \in [a, b];$$

lokale Minimalstelle, falls x_0 ein innerer Punkt ist und

$$f(x) \geq f(x_0) \quad \text{für alle } x \text{ in einer Umgebung von } x_0.$$

Man nennt $f(x_0)$ ein *globales* (bzw. *lokales*) *Minimum* von f in $[a, b]$.

Eine Maximalstelle oder Minimalstelle x_0 heisst auch eine *Extremalstelle* und der Wert $f(x_0)$ ein *Extremum*.



Wenn f auf einem abgeschlossenen Intervall $[a, b]$ definiert und dort stetig ist, dann hat f in $[a, b]$ (globale) Extrema. Wie finden wir diese Extrema?

Beim Graphen oben erkennen wir, dass die Tangente an den Graphen an einer lokalen Maximal- oder Minimalstelle waagrecht ist, das heisst, die Steigung der Tangente Null ist. Für eine lokale Extremalstelle x_0 gilt also $f'(x_0) = 0$.

Nicht jede Stelle x_0 mit $f'(x_0) = 0$ ist jedoch eine Extremalstelle. Auch bei einem Sattelpunkt (s. unten) verschwindet die erste Ableitung. Zum Beispiel gilt $f'(0) = 0$ für die Funktion $f(x) = x^3$, aber f hat in $x_0 = 0$ kein Extremum, sondern einen Sattelpunkt.

Die zweiten Ableitungen in x_0 (falls sie existieren) sagen uns, ob x_0 eine Extremalstelle ist, und wenn ja, von welcher Art.

Satz 4.6 Sei $x_0 \in (a, b)$. Dann gilt:

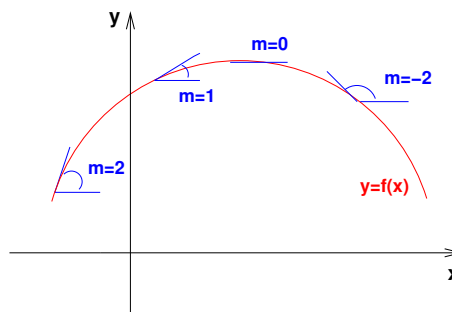
$$\begin{aligned} f'(x_0) = 0 \text{ und } f''(x_0) < 0 &\implies f(x_0) \text{ ist ein lokales Maximum} \\ f'(x_0) = 0 \text{ und } f''(x_0) > 0 &\implies f(x_0) \text{ ist ein lokales Minimum} \end{aligned}$$

Schauen wir uns diese beiden Fälle separat genauer an.

- Sei $f'(x_0) = 0$ und $f''(x_0) < 0$.

Dann gilt, dass $f''(x) < 0$ in einer (kleinen) Umgebung $U(x_0)$ von x_0 .

$$\begin{aligned} \implies f'(x) &\text{ ist monoton fallend in } U(x_0) \\ \implies f &\text{ beschreibt eine Rechtskurve in } U(x_0) \\ \implies f(x_0) &\text{ ist ein lokales Maximum} \end{aligned}$$

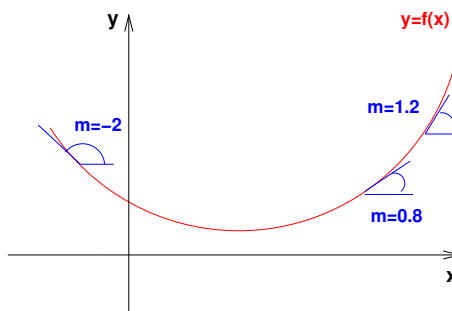


$m = f'(x)$ ist monoton fallend

- Sei $f'(x_0) = 0$ und $f''(x_0) > 0$.

Dann gilt, dass $f''(x) > 0$ in einer (kleinen) Umgebung $U(x_0)$ von x_0 .

$$\begin{aligned} \implies f'(x) &\text{ ist monoton wachsend in } U(x_0) \\ \implies f &\text{ beschreibt eine Linkskurve in } U(x_0) \\ \implies f(x_0) &\text{ ist ein lokales Minimum} \end{aligned}$$



$m = f'(x)$ ist monoton wachsend

Satz 4.6 ist nur für innere Punkte x_0 , wo f zweimal differenzierbar ist, anwendbar. Randpunkte müssen separat betrachtet werden. Beim Graphen auf Seite 58 zum Beispiel gilt für die globale Maximalstelle $x = b$, dass $f'(b) \neq 0$ (bzw. f ist nicht differenzierbar in b).

Weiter gelten die Pfeile in Satz 4.6 nicht in umgekehrter Richtung. Zum Beispiel hat die Funktion $f(x) = x^4$ in $x_0 = 0$ ein lokales (und globales) Minimum. Damit folgt $f'(0) = 0$. Aber für diese Funktion ist die zweite Ableitung ebenfalls Null: $f''(0) = 0$.

Vorgehen zur Bestimmung der Extremalstellen von $f : [a, b] \rightarrow \mathbb{R}$:

- (1) Bestimmung und Untersuchung von allen Stellen x_0 mit $f'(x_0) = 0$.
- (2) Berechnung von $f(a)$, $f(b)$ und den Funktionswerten in den Punkten, wo f nicht differenzierbar ist. Vergleich mit den in (1) erhaltenen Maxima und Minima.

Beispiel

Sei $f : [-1, 10] \rightarrow \mathbb{R}$, $f(x) = (x^2 - 4x)e^{-x}$.

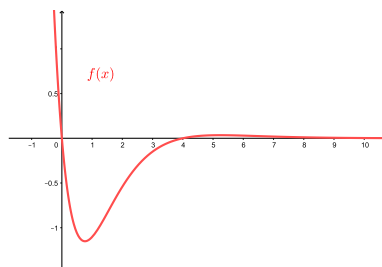
(1)

Es folgt:

$x_1 = 5,24$ ist eine lokale Maximalstelle, $f(x_1) = 0,034$ ein lokales Maximum

$x_2 = 0,76$ ist eine lokale Minimalstelle, $f(x_2) = -1,15$ ein lokales Minimum

- (2) f ist überall differenzierbar. Vorher unter (1) haben wir demnach alle lokalen Extremalstellen gefunden. Für die globalen Extremalstellen müssen wir nur noch die Funktionswerte $f(-1)$ und $f(10)$ mit dem in (1) gefundenen Maximum und Minimum vergleichen.

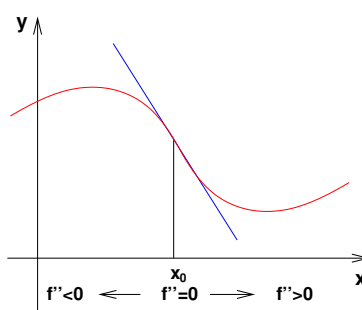


Wir schließen, dass $f(5,24) = 0,034$ ein lokales Maximum, $f(-1) = 13,59$ ein globales Maximum und $f(0,76) = -1,15$ ein lokales und globales Minimum ist.

Wir haben oben gesehen, dass der Funktionsgraph bei einem lokalen Maximum eine Rechtskurve beschreibt (es gilt dort $f''(x) < 0$) und bei einem lokalen Minimum eine Linkskurve (es gilt dort $f''(x) > 0$). Ist f'' stetig, dann gibt es (nach dem Nullstellensatz) eine Stelle x_0 dazwischen mit $f''(x_0) = 0$. An dieser Stelle geht die Rechts- in eine Linkskurve (bzw. die Links- in eine Rechtskurve) über.

Definition Eine Funktion f hat in x_0 einen *Wendepunkt*, wenn im Punkt $(x_0, f(x_0))$ die Krümmung des Graphen von f wechselt (die Rechts- in eine Linkskurve übergeht oder umgekehrt). Gilt zusätzlich $f'(x_0) = 0$ (waagrechte Tangente in x_0), dann nennt man den Wendepunkt einen *Sattelpunkt*.

An einem Wendepunkt “schneidet” die **Tangente** den Funktionsgraphen.

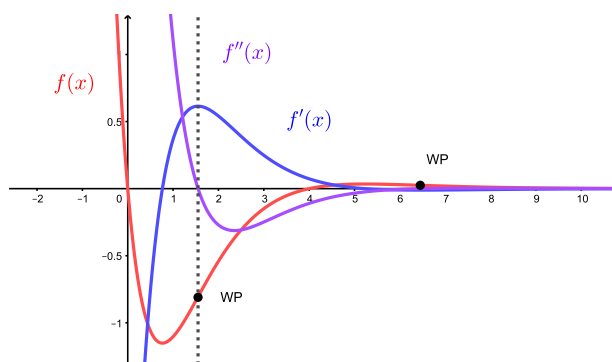


Satz 4.7 Sei $x_0 \in (a, b)$. Dann gilt:

$$f''(x_0) = 0 \text{ und } f'''(x_0) \neq 0 \implies f \text{ hat in } x_0 \text{ einen Wendepunkt}$$

Beispiel

Wir betrachten nochmals $f(x) = (x^2 - 4x)e^{-x}$. Die Gleichung $f''(x) = (x^2 - 8x + 10)e^{-x} = 0$ hat die beiden Lösungen $x_1 = 4 + \sqrt{6} \approx 6,45$ und $x_2 = 4 - \sqrt{6} \approx 1,55$. Da $f'''(x_1) \neq 0$ und $f'''(x_2) \neq 0$, sind $(6,45; 0,025)$ und $(1,55; -0,81)$ zwei Wendepunkte von f .



Hat f in x_0 einen Wendepunkt, dann hat die Ableitung f' in x_0 ein Extremum. Dies erklärt die Bedingung $f'''(x_0) \neq 0$ (vgl. Satz 4.6).

Die Bedingung $f'''(x_0) \neq 0$ in Satz 4.7 ist jedoch nicht notwendig für einen Wendepunkt in x_0 (da $f''(x_0) \neq 0$ in Satz 4.6 nicht notwendig ist). Zum Beispiel hat die Funktion $f(x) = x^5$ einen Wendepunkt in $x_0 = 0$ mit $f'''(0) = 0$.

4.3 Die Regeln von Bernoulli-de l'Hôpital

In Kapitel 2 haben wir gesehen, wie man Grenzwerte der Form $\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)}$ berechnet. Es gibt jedoch Grenzwerte, die wir mit den dort angegebenen Methoden nicht bestimmen können. Betrachten wir zum Beispiel den Grenzwert

$$\lim_{x \rightarrow 0} \frac{\sin x}{x}.$$

Bei diesem Beispiel strebt sowohl der Zähler als auch der Nenner gegen 0 für $x \rightarrow 0$.

Mit Hilfe des Mittelwertsatzes (Satz 4.4) können wir nun diesen Grenzwert bestimmen (wobei es auch mit der Potenzreihe von $\sin x$ ginge). Wir wenden also den Mittelwertsatz für $f(x) = \sin x$ auf dem Intervall $[0, x]$ an:

Dieser Trick funktioniert allgemein und wir erhalten die *Regeln von Bernoulli-de l'Hôpital*.

Satz 4.8 Seien $f, g : (a, b) \rightarrow \mathbb{R}$ differenzierbar, wobei $-\infty \leq a < b \leq \infty$. Weiter gelte

$$\lim_{x \rightarrow a} f(x) = \lim_{x \rightarrow a} g(x) = 0 \quad \text{oder} \quad \lim_{x \rightarrow a} f(x) = \lim_{x \rightarrow a} g(x) = \pm\infty.$$

Falls der Grenzwert $\lim_{x \rightarrow a} \frac{f'(x)}{g'(x)}$ existiert (d.h. $\in \mathbb{R} \cup \{\pm\infty\}$ ist), gilt

$$\lim_{x \rightarrow a} \frac{f(x)}{g(x)} = \lim_{x \rightarrow a} \frac{f'(x)}{g'(x)}.$$

Dies gilt analog für Grenzwerte $x \rightarrow b$.

Beispiele

$$1. \lim_{x \rightarrow 0} \frac{e^x - 1}{x}$$

$$2. \lim_{x \rightarrow \infty} \frac{x^2 + 10^{9999}}{2^x - 10^{2025}}$$

Dies erklärt die Merkregel *Exponentiell ist stärker als polynomial* von Kapitel 3 (Seite 40).

$$3. \lim_{x \rightarrow 0} \frac{\ln(1+cx)}{x} = \lim_{x \rightarrow 0} \frac{(\ln(1+cx))'}{(x)'} = \lim_{x \rightarrow 0} \frac{c}{1+cx} = c.$$

Damit gilt weiter

$$\lim_{x \rightarrow 0} (1+cx)^{\frac{1}{x}} = \lim_{x \rightarrow 0} e^{\frac{1}{x} \ln(1+cx)} = e^{\lim_{x \rightarrow 0} \frac{\ln(1+cx)}{x}} = e^c,$$

wobei wir beim zweiten Gleichheitszeichen die Stetigkeit von e^x benutzt haben. Für die Folge $x_n = \frac{1}{n}$ mit $\lim_{n \rightarrow \infty} x_n = 0$ folgt insbesondere, dass

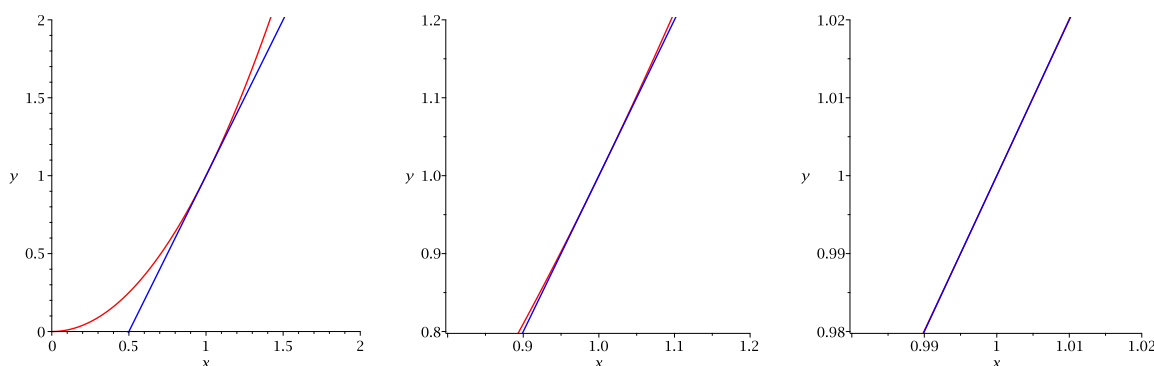
$$e^c = \lim_{x \rightarrow 0} (1+cx)^{\frac{1}{x}} = \lim_{n \rightarrow \infty} (1+cx_n)^{\frac{1}{x_n}} = \lim_{n \rightarrow \infty} \left(1 + \frac{c}{n}\right)^n,$$

wie in Kapitel 1 (Seite 23) für $c = 1$ behauptet.

4.4 Lineare Approximation

An den Graphen einer differenzierbaren Funktion f legen wir die Tangente in einem Punkt $P = (x_0, f(x_0))$ an. Zoomen wir den Punkt P stark heran, so ist die Funktionskurve von der Geraden nicht mehr zu unterscheiden!

Beispiel $f(x) = x^2$, Tangente an f im Punkt $P = (1, 1)$



Die Tangente eignet sich also als lokale (lineare) Näherung der Funktion.

Satz 4.9 *Unter allen Geraden durch den Punkt $(x_0, f(x_0))$ ist die Tangente diejenige Gerade, die f lokal um x_0 (d.h. in der Nähe von x_0) am besten approximiert.*

Wie ist das gemeint? Die Tangente an den Graphen von f in x_0 ist gegeben durch

$$y = t(x) = f(x_0) + f'(x_0)(x - x_0).$$

Für x nahe bei x_0 stimmen die Funktionswerte $f(x)$ ungefähr mit den Werten $t(x)$ überein: $f(x) \approx t(x)$. Dies können wir auch so schreiben:

$$f(x_0 + h) \approx t(x_0 + h) \quad \text{für } h \text{ mit } |h| \text{ klein.}$$

Der Fehler $r(h)$ dieser Näherung ist dabei gegeben durch

$$r(h) = f(x_0 + h) - t(x_0 + h) = f(x_0 + h) - f(x_0) - f'(x_0) \cdot h. \quad (*)$$

Da $t(x_0) = f(x_0)$, geht $r(h)$ gegen 0 für $h \rightarrow 0$. Aber auch $\frac{r(h)}{h}$ geht gegen 0 für $h \rightarrow 0$:

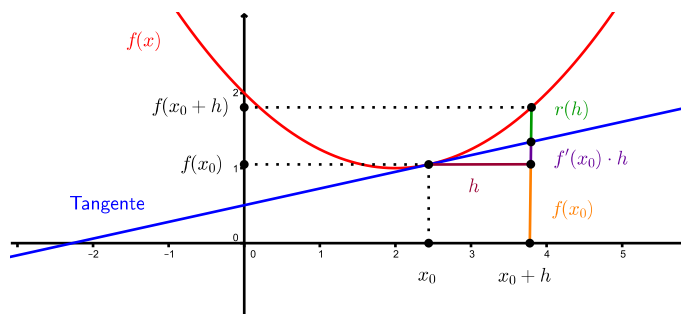
$$\lim_{h \rightarrow 0} \frac{r(h)}{h} = \lim_{h \rightarrow 0} \left(\frac{f(x_0 + h) - f(x_0)}{h} - f'(x_0) \right) = f'(x_0) - f'(x_0) = 0$$

Hier ist entscheidend, dass $t'(x_0) = f'(x_0)$, das heisst, dass die Steigungen von t und f in x_0 übereinstimmen. Deshalb ist die Tangente die einzige Gerade durch den Punkt $(x_0, f(x_0))$, für die nicht nur $r(h)$, sondern auch $\frac{r(h)}{h}$ gegen 0 geht für $h \rightarrow 0$, was bedeutet, dass die Differenz $r(h) = f(x_0 + h) - t(x_0 + h)$ schneller als h gegen 0 geht.

Lösen wir die Gleichung (*) nach $f(x_0 + h)$ auf, erhalten wir die folgende Näherung für f in der Nähe von x_0 .

Satz 4.10 Sei $f : D \rightarrow \mathbb{R}$ differenzierbar in $x_0 \in D$. Dann gilt

$$f(x_0 + h) = f(x_0) + f'(x_0) \cdot h + r(h) \quad \text{mit} \quad \lim_{h \rightarrow 0} \frac{r(h)}{h} = 0.$$



Da das Restglied $r(h)$ für kleine $|h|$ verschwindend klein ist, benutzt man für komplizierte Funktionen f in der Nähe einer Stelle x_0 oft die folgende lineare Näherung.

Näherungsformel (N) $f(x_0 + h) \approx f(x_0) + f'(x_0) \cdot h$ für kleine $|h|$

Beispiel

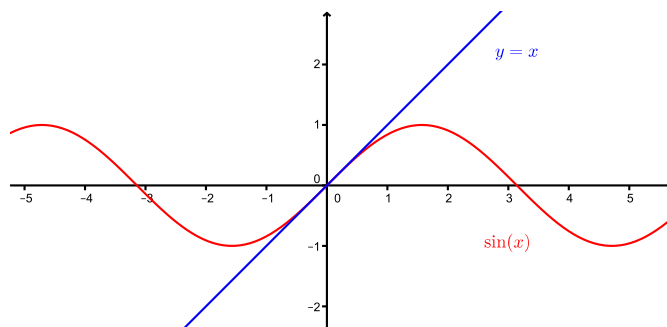
Mit Hilfe der Näherungsformel (N) kann zum Beispiel (im Kopf!) eine gute Näherung von $\sqrt{8,92}$ gegeben werden.

Der auf 7 Nachkommastellen gerundete Wert von $\sqrt{8,92}$ ist 2,9866369, der Näherungsfehler $r(-0,08)$ beträgt also nur $-0,0000298 = 2,98 \cdot 10^{-5}$!

Auf der Formel (N) beruhen auch Näherungen einiger elementarer Funktionen, die häufig (zum Beispiel in der Physik) gebraucht werden. Für kleine $|x|$ gilt:

$$\begin{aligned}(1+x)^n &\approx 1+nx \\ \frac{1}{1-x} &\approx 1+x \\ \sqrt{1+x} &\approx 1+\frac{x}{2} \\ e^x &\approx 1+x \\ \sin x &\approx x\end{aligned}$$

Um zum Beispiel die Näherung $\sin x \approx x$ zu erhalten, nutzt man die Näherungsformel (N) für $f(x) = \sin x$, $h = x$, $x_0 = 0$ mit $f'(x) = \cos x$.



4.5 Taylorpolynome und Taylorreihen

Im vorhergehenden Abschnitt haben wir eine reelle Funktion f in der Nähe eines Punktes $(x_0, f(x_0))$ durch eine lineare Funktion approximiert. Wir können diese Näherung verbessern, indem wir anstelle einer linearen Funktion eine Polynomfunktion von höherem Grad verwenden.

Damit diese Näherung gut ist, soll die Polynomfunktion $p(x)$ in (der Nähe von) x_0 möglichst die gleiche Krümmung wie die Funktion $f(x)$ aufweisen. Das Krümmungsverhalten von f wird durch die Ableitungen von f beschrieben. Ist die Funktion f n -mal differenzierbar, so können wir fordern, dass die ersten n Ableitungen von f und von p in x_0 übereinstimmen, und zusätzlich soll natürlich $p(x_0) = f(x_0)$ sein. Es soll also gelten:

- $p(x_0) = f(x_0)$ (gleiche Funktionswerte in x_0)
- $p'(x_0) = f'(x_0)$ (gleiche Tangente in x_0)
- $p''(x_0) = f''(x_0)$ (gleiche Krümmung in x_0)
- $\vdots$
- $p^{(n)}(x_0) = f^{(n)}(x_0)$

Dies sind $n+1$ Bedingungen. Mit einer Polynomfunktion $p(x) = p_n(x)$ mit $n+1$ Koeffizienten, also vom Grad $\leq n$, können alle Bedingungen erfüllt werden.

Satz 4.11 Das Polynom $p_n(x)$ vom Grad $\leq n$, das die Bedingungen $p^{(k)}(x_0) = f^{(k)}(x_0)$ für $k = 0, \dots, n$ erfüllt, ist gegeben durch

$$p_n(x) = f(x_0) + \frac{f'(x_0)}{1!}(x - x_0) + \dots + \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n = \sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!}(x - x_0)^k.$$

Das Polynom $p_n(x)$ heisst n -tes Taylorpolynom von f um den Entwicklungspunkt x_0 .

Für $n = 1$ ist das Taylorpolynom gegeben durch

$$p_1(x) = f(x_0) + f'(x_0)(x - x_0).$$

Wenig überraschend ist dies die Gleichung der Tangente an f in x_0 !

Die Taylorpolynome sind also Erweiterungen der linearen Näherung vom letzten Abschnitt, bzw. sie verbessern die lineare Näherung. Satz 4.10 lautet für $h = x - x_0$:

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + r(x - x_0) \quad \text{mit} \quad \lim_{x \rightarrow x_0} \frac{r(x - x_0)}{x - x_0} = 0$$

Für die Taylorpolynome $p_n(x)$ gilt nun

$$f(x) = p_n(x) + r_n(x - x_0) \quad \text{mit} \quad \lim_{x \rightarrow x_0} \frac{r_n(x - x_0)}{(x - x_0)^n} = 0.$$

Das Restglied $r_n(x - x_0)$ geht also sehr schnell gegen 0 für $x \rightarrow x_0$. Falls f $(n + 1)$ -mal differenzierbar ist, kann das Restglied (das Lagrangesche Restglied) geschrieben werden als

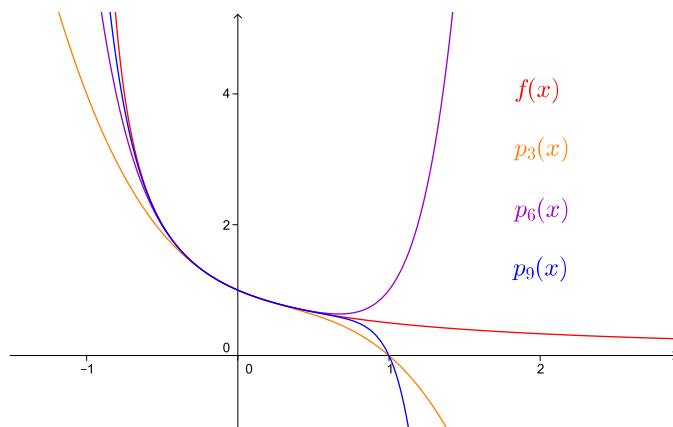
$$r_n(x - x_0) = \frac{f^{(n+1)}(\xi)}{(n + 1)!}(x - x_0)^{n+1} \quad \text{für ein } \xi \text{ zwischen } x_0 \text{ und } x.$$

Beispiele

1. Sei $f(x) = \frac{1}{1+x}$. Gesucht ist das Taylorpolynom $p_3(x)$ im Entwicklungspunkt $x_0 = 0$.

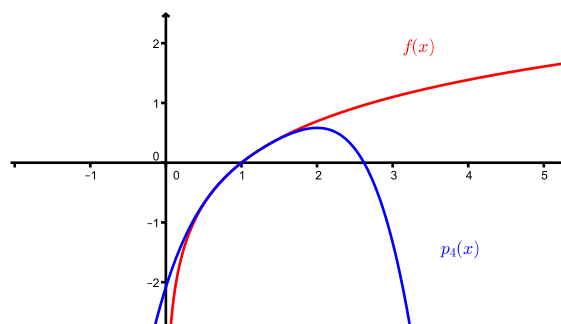
Weiter findet man

$$\begin{aligned} p_6(x) &= 1 - x + x^2 - x^3 + x^4 - x^5 + x^6 \\ p_9(x) &= 1 - x + x^2 - x^3 + x^4 - x^5 + x^6 - x^7 + x^8 - x^9 \\ &\vdots \\ p_n(x) &= \sum_{k=0}^n (-1)^k x^k \end{aligned}$$



2. Sei $f(x) = \ln(x)$. Gesucht ist das Taylorpolynom $p_4(x)$ im Entwicklungspunkt $x_0 = 1$.
Wir berechnen

$$\begin{aligned} f'(x) &= \frac{1}{x} &\implies f'(1) &= 1 \\ f''(x) &= -\frac{1}{x^2} &\implies f''(1) &= -1 \\ f'''(x) &= \frac{2}{x^3} &\implies f'''(1) &= 2 \\ f^{(4)}(x) &= -\frac{6}{x^4} &\implies f^{(4)}(1) &= -6 \end{aligned}$$



Der Entwicklungspunkt $x_0 = 1$ bedeutet, dass $p_4(x)$ die Funktion $f(x) = \ln(x)$ in der Nähe von $x_0 = 1$ gut approximiert. Man löst die Klammern $(x-1)^k$ im Taylorpolynom nicht auf, denn so erhält man für ein konkretes x effizient die Differenz von $\ln(x)$ und $p_4(x)$. Zum Beispiel gilt für $x = 1,1$, dass

$$\ln(1,1) - p_4(1,1) = 0,0953 - \left(0,1 - \frac{0,1^2}{2} + \frac{0,1^3}{3} - \frac{0,1^4}{4}\right) \approx 1,85 \cdot 10^{-6}.$$

Für $x = 2$ ist $p_4(x)$ keine so gute Näherung für $\ln(x)$ mehr; es gilt $\ln(2) - p_4(2) \approx 0,110$.

Ist die Funktion f unendlich oft differenzierbar in einer Stelle x_0 , dann können nicht nur die Taylorpolynome betrachtet werden, sondern die sogenannte *Taylorreihe* um den Entwicklungspunkt x_0 ,

$$\sum_{k=0}^{\infty} \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k.$$

Ob diese Reihe überhaupt konvergent ist, hängt davon ab, wie nahe x bei x_0 ist. Jedoch konvergiert die Taylorreihe auch für x nahe bei x_0 nicht für alle Funktionen f , und wenn, dann muss die Reihe auch gar nicht mit den Funktionswerten $f(x)$ übereinstimmen.

Es gibt aber einige schöne Beispiele, wo die Taylorreihe für gewisse, bzw. alle $x \in \mathbb{R}$ konvergiert und für diese x gleich den Funktionswerten $f(x)$ ist.

Beispiele

1. Betrachten wir nochmals $f(x) = \frac{1}{1+x}$. Diese Funktion ist unendlich oft differenzierbar, und tatsächlich wissen wir ja von den geometrischen Reihen, dass

$$f(x) = \frac{1}{1+x} = \sum_{k=0}^{\infty} (-1)^k x^k \quad \text{für } |x| < 1.$$

2. Sei $f(x) = e^x$. Für diese Funktion gilt $f^{(k)}(x) = e^x$ für alle $k \geq 1$. Wählen wir den Entwicklungspunkt $x_0 = 0$, dann gilt $f^{(k)}(0) = e^0 = 1$ für alle k . Wir erhalten die Taylorreihe

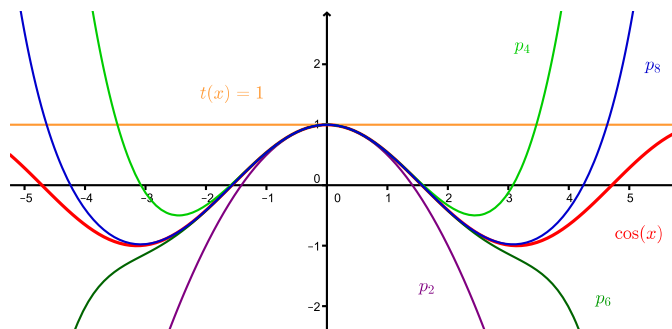
$$1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots = \sum_{k=0}^{\infty} \frac{x^k}{k!} = e^x \quad \text{für alle } x \in \mathbb{R}.$$

Die Potenzreihe, mit welcher wir e^x definiert haben, ist genau die Taylorreihe von e^x !

3. Genauso verhält es sich mit den Funktionen $\sin x$ und $\cos x$. Wir haben auf Seite 54 gesehen, dass diese Funktionen als unendliche Reihen darstellbar sind, nämlich

$$\begin{aligned} \sin x &= x - \frac{x^3}{3!} + \frac{x^5}{5!} - \dots = \sum_{k=0}^{\infty} \frac{(-1)^k}{(2k+1)!} x^{2k+1} \\ \cos x &= 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \dots = \sum_{k=0}^{\infty} \frac{(-1)^k}{(2k)!} x^{2k} \end{aligned}$$

Dies sind die Taylorreihen von $\sin x$ und $\cos x$ um den Entwicklungspunkt $x_0 = 0$. Sie konvergieren also für alle $x \in \mathbb{R}$ gegen die Funktionswerte.



4.6 Näherungsverfahren zur Lösung von Gleichungen

Das Newton-Verfahren

Mit Hilfe des Newton-Verfahrens können Nullstellen einer differenzierbaren Funktion näherungsweise bestimmt werden. Dieses Verfahren beruht auf der Näherungsformel (N) von Abschnitt 4.4, bzw. auf der Näherung der Funktion durch Tangenten.

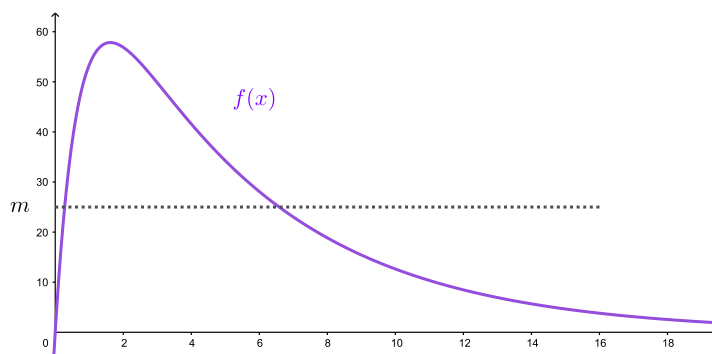
Beispiel

Nehmen wir an, Sie nehmen ein Medikament in Form einer Tablette ein. Dann kann die Konzentration des Wirkstoffes im Blutplasma zum Zeitpunkt x nach der Einnahme durch die Funktion

$$f(x) = C \frac{k_a}{k_a - k_e} (e^{-k_e x} - e^{-k_a x})$$

in μg pro ml beschrieben werden (unter gewissen vereinfachenden Annahmen). Dabei ist k_a die Absorptionskonstante (also ein Mass dafür, wie schnell der Wirkstoff aufgenommen wird), k_e ist die Eliminationskonstante (ein Mass dafür, wie schnell der Wirkstoff abgebaut wird) und C ist eine zeitunabhängige Grösse.

Zur Bestimmung des Zeitpunktes x_{max} , in welchem die Konzentration des Wirkstoffes am grössten ist, berechnet man die Nullstelle der Ableitung von f . Tatsächlich kann die Gleichung $f'(x_{max}) = 0$ problemlos nach x_{max} aufgelöst werden. Anders sieht es aus, wenn man die Zeitspanne berechnen möchte, während der das Medikament wirkt. Für die Wirkung braucht es eine minimale Konzentration $m > 0$. Graphisch sieht dies so aus:



Zu bestimmen sind also x mit $f(x) = m$, bzw. mit $f(x) - m = 0$. Wieder brauchen wir Nullstellen einer Funktion, hier der Funktion $f(x) - m$. Im Gegensatz zu vorher kann die Gleichung $f(x) - m = 0$ jedoch nicht nach x aufgelöst werden! Wir können diese Nullstellen nur mit einem Näherungsverfahren bestimmen. Analog kann man Nullstellen einer allgemeinen Polynomfunktion vom Grad ≥ 5 nicht exakt bestimmen (es gibt keine Lösungsformel, wie wir in Kapitel 2 gesehen haben). Für solche Situationen brauchen wir das Newton-Verfahren.

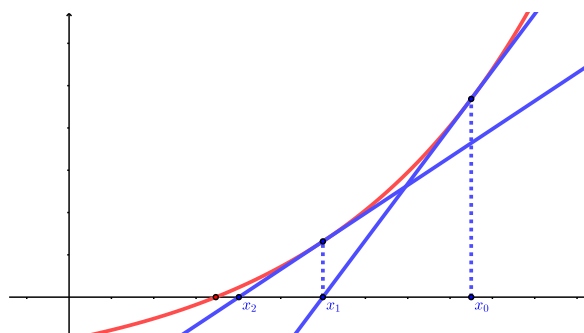
Gegeben ist also eine Funktion $f : [a, b] \rightarrow \mathbb{R}$, von der wir annehmen, dass sie eine Nullstelle hat, es also ein x_{null} in $[a, b]$ mit $f(x_{\text{null}}) = 0$ gibt, und weiter nehmen wir an, dass f differenzierbar ist. Dann nähern wir uns x_{null} wie folgt.

Wir wählen eine Stelle $x_0 \in [a, b]$ aus (die wir nahe bei x_{null} vermuten) und betrachten die lineare Näherung von f um x_0 , das heisst

$$f(x) \approx f(x_0) + f'(x_0) \cdot (x - x_0) = t(x).$$

Wie wir schon gesehen haben, ist die Funktion $t(x)$ nichts anderes als die Tangente an den Graphen von f in x_0 . Als Nächstes bestimmen wir die Nullstelle x_1 dieser Tangente. Um eine waagrechte Tangente $t(x)$ ohne Nullstelle auszuschliessen, sollten wir x_0 mit $f'(x_0) \neq 0$ gewählt haben.

In günstigen Situationen liegt nun x_1 näher bei x_{null} als x_0 . Der Vorgang wird wiederholt mit x_1 anstelle von x_0 , um noch eine bessere Näherung x_2 von x_{null} zu finden, usw.



Definition Das *Newton-Verfahren* ist das wiederholte Anwenden der Formel

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$$

für $n \geq 0$, wobei x_0 ein geschickt gewählter Startpunkt ist (möglichst nahe bei der gesuchten Nullstelle und so, dass $f'(x_0) \neq 0$).

In vielen Fällen konvergiert die Folge (x_n) schnell gegen die Nullstelle x_{null} , wenn x_0 nahe bei x_{null} ist.

Beispiel

Gesucht ist eine Nullstelle der Funktion $f(x) = x^5 + 5x + 1$.

Für einen geeigneten Startwert schaut man sich entweder den Graphen von f an oder man nutzt den Nullstellensatz [3.8](#) (Seite [49](#)):

Mit dem Startwert $x_0 = 0$ findet man:

n	x_n
1	-0,2
2	-0,199936102
3	-0,199936102

Mit dem Startwert $x_0 = 10$ brauchen wir 13 Schritte, um auf die Zahl $-0,199936102$ zu kommen, mit dem Startwert $x_0 = 1000$ sind es 35 Schritte.

Dies war ein harmloses Beispiel, wo das Newton-Verfahren schnell konvergierte. Dies lag vor allem daran, dass f eine streng monoton wachsende Funktion ist und demnach auch nur eine Nullstelle hat.

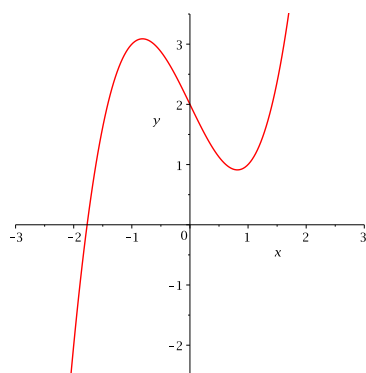
Es gibt aber andere Beispiele, wo bei ungeschickt gewähltem Startwert das Verfahren versagt.

Beispiele

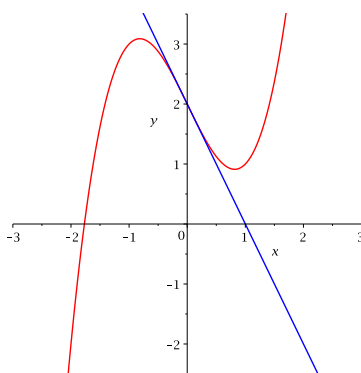
1. Gesucht ist eine Nullstelle der Funktion $f(x) = x^3 - 2x + 2$.

Die Iterationsformel lautet

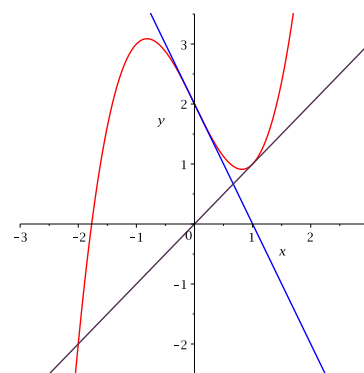
$$x_{n+1} = x_n - \frac{x_n^3 - 2x_n + 2}{3x_n^2 - 2}.$$



$x_0 = 0$



$x_1 = 1$



$x_2 = 0$

Mit dem Startwert $x_0 = 0$ (oder $x_0 = 1$) nimmt x_n abwechselungsweise die Werte 0 und 1 an. Die Folge (x_n) konvergiert also nicht.

Mit den Startwerten $x_0 = -1$ oder $x_0 = -4$ zum Beispiel findet man hingegen (nach 7 bzw. 6 Schritten) die Nullstelle

$$x_{\text{null}} = -1,769292354.$$

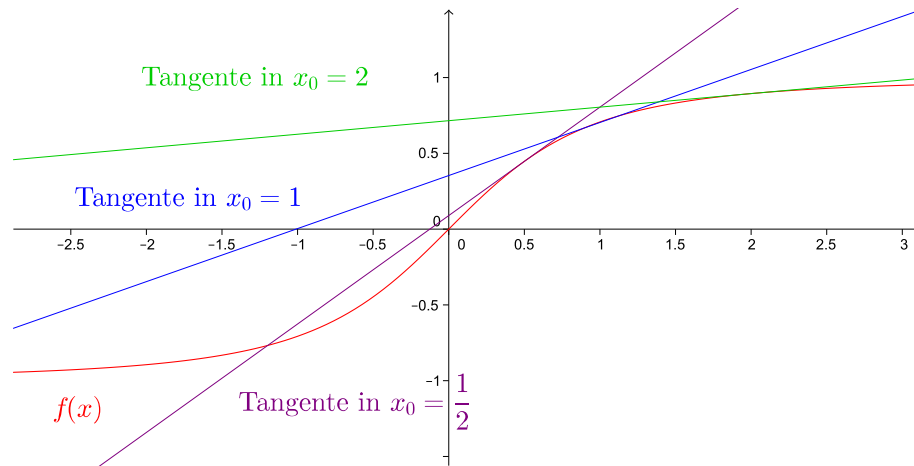
2. Die Funktion $f(x) = x^3 - 2x + \frac{1}{8}$ hat drei verschiedene Nullstellen. Um alle drei Nullstellen mit Hilfe des Newton-Verfahrens zu erhalten, müssen die Startwerte geschickt gewählt werden. Ungünstig sind Startwerte nahe bei einer Extremalstelle, da an diesen Stellen die Tangente an die Funktion fast horizontal verläuft.

3. Die Folge (x_n) des Newton-Verfahrens konvergiert für die Funktion $f(x) = x^2 - 2x + 2$ mit keinem Startwert x_0 . Der einfache Grund ist, dass diese Funktion gar keine Nullstelle hat.

4. Offensichtlich ist $x = 0$ die einzige Nullstelle der Funktion

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

Allerdings konvergiert die Folge (x_n) des Newton-Verfahrens nur für Startwerte x_0 mit $|x_0| < 1$. Für $|x_0| > 1$ divergiert die Folge (x_n) und für $x_0 = \pm 1$ springt die Folge zwischen $+1$ und -1 hin und her (Details dazu siehe Zusatzaufgabe der Übung).



Bemerkung Ist $f : [a, b] \rightarrow \mathbb{R}$ zweimal stetig differenzierbar, dann gilt für die Konvergenz des Newton-Verfahrens:

- Ist x_{null} eine einfache Nullstelle (d.h. $f'(x_{\text{null}}) \neq 0$), dann konvergiert das Newton-Verfahren für alle Startwerte x_0 nahe bei x_{null} .

In diesem Fall gibt es eine reelle Zahl $c > 0$, so dass

$$|x_{n+1} - x_{\text{null}}| \leq c |x_n - x_{\text{null}}|^2 \quad \text{für alle } n.$$

Grob gesagt verdoppelt sich die Anzahl der korrekten Stellen bei jedem Schritt.

- Bei einer mehrfachen Nullstelle x_{null} ist die Konvergenz langsamer.

Sucht man eine Lösung einer beliebigen (algebraischen) Gleichung, so kann man alle Terme der Gleichung auf eine Seite bringen, so dass man eine Gleichung der Form $f(x) = 0$ erhält. Mit Hilfe des Newton-Verfahrens findet man so (im günstigen Fall) näherungsweise eine Lösung der ursprünglichen Gleichung.

Beispiel

Unser Taschenrechner sagt uns, dass

$$\sqrt{2} = 1,4142136 \dots$$

Woher kommt diese Zahl? Tatsächlich kann $\sqrt{2}$ sehr schnell mit Hilfe des Newton-Verfahrens berechnet werden. Die Zahl $\sqrt{2}$ ist ja definiert als positive Lösung der Gleichung

$$x^2 = 2.$$

Diese Gleichung kann zur Gleichung $x^2 - 2 = 0$ umgeformt werden. Wir suchen also die positive Nullstelle der Funktion

$$f(x) = x^2 - 2.$$

Die Iterationsformel des Newton-Verfahrens lautet damit

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} = x_n - \frac{x_n^2 - 2}{2x_n} = \frac{2x_n^2 - x_n^2 + 2}{2x_n} = \frac{1}{2} \left(x_n + \frac{2}{x_n} \right)$$

für $n \geq 0$. Da $1^2 < 2 < 2^2$, gilt $1 < \sqrt{2} < 2$. Wir können also zum Beispiel den Startwert $x_0 = 1$ wählen und erhalten

n	x_n
1	1,5
2	1,41 $\bar{6}$
3	1,414215686
4	1,414213562
5	1,414213562

Die obige Iterationsformel für $\sqrt{2}$ ist übrigens genau die rekursiv definierte Zahlenfolge, die wir schon in Kapitel 3 (Seite 40) untersucht haben.

Fixpunktiteration

Die Fixpunktiteration ist ein weiteres Verfahren, mit dem man eine Lösung einer Gleichung näherungsweise bestimmen kann.

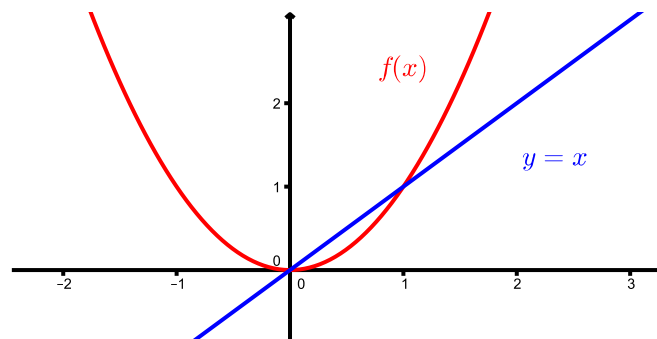
Definition Ist $f : D \rightarrow \mathbb{R}$ eine Funktion, dann heisst $x \in D$ *Fixpunkt* von f , falls gilt

$$f(x) = x.$$

Beispiele

1. Sei $f(x) = \frac{1}{2}x$. Dann ist 0 der einzige Fixpunkt von f .
2. Die Funktion $f(x) = x^2$ hat die beiden Fixpunkte 0 und 1.

Geometrisch betrachtet ist die x -Koordinate von jedem Schnittpunkt des Graphen von f mit der Geraden $y = x$ ein Fixpunkt von f .



Mit der Fixpunktiteration kann man die Fixpunkte einer Funktion bestimmen.

Definition Bei der *Fixpunktiteration* wählt man einen Startwert x_0 und berechnet schrittweise

$$x_1 = f(x_0), \quad x_2 = f(x_1), \quad \dots, \quad x_{n+1} = f(x_n), \quad \dots$$

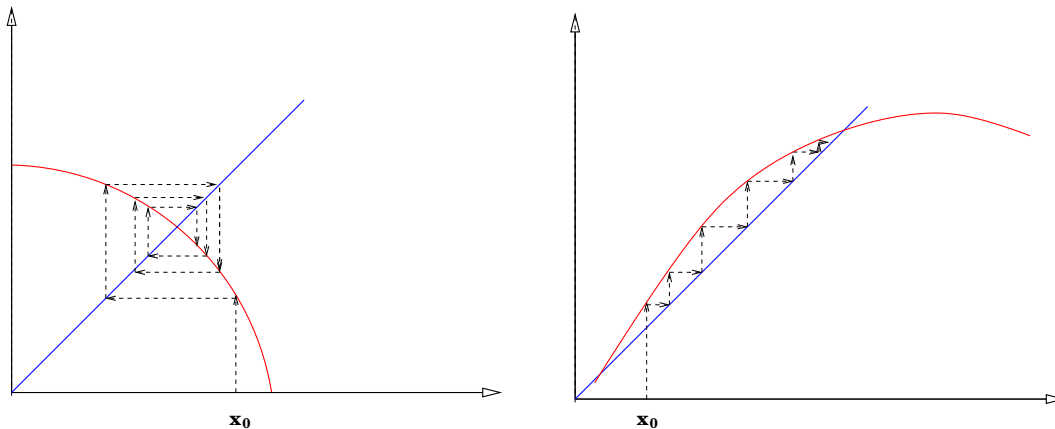
Konvergiert die Folge (x_n) gegen einen Grenzwert x_{fix} und ist f stetig, dann folgt

$$f(x_{\text{fix}}) = x_{\text{fix}},$$

das heisst, x_{fix} ist ein Fixpunkt von f .

Wegen der Stetigkeit von f gilt nämlich

$$x_{\text{fix}} = \lim_{n \rightarrow \infty} x_{n+1} = \lim_{n \rightarrow \infty} f(x_n) = f(\lim_{n \rightarrow \infty} x_n) = f(x_{\text{fix}}).$$



Beispiel

Gesucht ist $x \in \mathbb{R}$ mit $\cos x = x$. Wir wählen den Startwert $x_0 = 1$ und führen die Fixpunktiteration durch. Mit dem Taschenrechner geht das so: Das Bogenmass einstellen (Anzeige RAD), dann fortlaufend die cos-Taste drücken.

Wir brauchen ein wenig Geduld, doch wir erhalten schliesslich nach 53 Schritten den Fixpunkt

$$x_{\text{fix}} = 0,739085133.$$

Unter gewissen Voraussetzungen ist die Konvergenz der Fixpunktiteration garantiert.

Satz 4.12 Sei $D = [a, b]$ ein Intervall und $f : D \rightarrow \mathbb{R}$ differenzierbar mit $f(D) \subset D$. Es gebe ein $K < 1$, so dass $|f'(x)| \leq K$ für alle x in D . Dann besitzt f genau einen Fixpunkt x_{fix} in D und die Folge

$$x_{n+1} = f(x_n)$$

konvergiert mit jedem beliebigen Startwert $x_0 \in D$ gegen x_{fix} .

Die Bedingung $|f'(x)| \leq K$ für ein $K < 1$ bedeutet geometrisch, dass der Graph von f eine flache Kurve ist (die Tangentensteigung ist überall "klein").

Aus dieser Bedingung folgt mit dem Mittelwertsatz (Satz 4.4) die Konvergenz, denn

$$|x_{n+1} - x_{\text{fix}}| = |f(x_n) - f(x_{\text{fix}})| \leq |f'(\xi)| \cdot |x_n - x_{\text{fix}}| \leq K|x_n - x_{\text{fix}}| < |x_n - x_{\text{fix}}|$$

für ein ξ zwischen x_n und x_{fix} . Die Bedingung $|f'(x)| \leq K$ ist auch für die Eindeutigkeit des Fixpunktes verantwortlich, denn ist y_{fix} ein zweiter Fixpunkt, dann gilt

$$|x_{\text{fix}} - y_{\text{fix}}| = |f(x_{\text{fix}}) - f(y_{\text{fix}})| \leq K|x_{\text{fix}} - y_{\text{fix}}|.$$

Wegen $K < 1$ folgt $|x_{\text{fix}} - y_{\text{fix}}| = 0$, also $x_{\text{fix}} = y_{\text{fix}}$.

Bemerkung Mit den Bezeichnungen von Satz 4.12 gilt die Fehlerabschätzung

$$|x_{\text{fix}} - x_n| \leq \frac{K}{1-K}|x_n - x_{n-1}| \leq \frac{K^n}{1-K}|x_1 - x_0|.$$

Die Fixpunktiteration konvergiert umso schneller, je kleiner K ist.

Beispiele

1. Beim Newton-Verfahren haben wir die Gleichung $x^5 + 5x + 1 = 0$ betrachtet. Diese Gleichung können wir in eine Fixpunktgleichung umformen.

Wir wenden also die Fixpunktiteration auf $f(x) = -\frac{1}{5}(x^5 + 1)$ an.

Für die Ableitung gilt

Wählen wir zum Beispiel $D = [-\frac{4}{5}, \frac{4}{5}]$, dann ist $|f'(x)| \leq K$ für $K = (\frac{4}{5})^4 < 1$ erfüllt.

Weiter gilt $f(D) \subset D$. Warum? Nun, $f(-\frac{4}{5}) \approx -0,266 \in D$ und $f(\frac{4}{5}) \approx -0,134 \in D$. Da $f'(x) = -x^4 \leq 0$ für alle $x \in D$, ist f monoton fallend in D . Also liegen alle Funktionswerte für $x \in D$ zwischen $f(-\frac{4}{5})$ und $f(\frac{4}{5})$, d.h. $f(D) \subset [-0,266; -0,134] \subset D$.

Wählen wir nun zum Beispiel den Startwert $x_0 = 0$. Nach Satz 4.12 muss die Fixpunktiteration gegen den Fixpunkt in D konvergieren, was auch der Fall ist:

n	x_n
1	-0,2
2	-0,199936
3	-0,199936102
4	-0,199936102

Im Allgemeinen konvergiert das Newton-Verfahren schneller als die Fixpunktiteration.

2. Die Funktion

$$f(x) = x^2$$

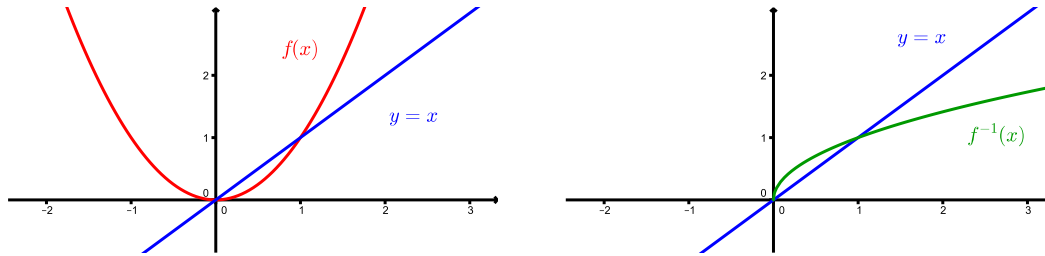
erfüllt die Bedingung $|f'(x)| \leq K$ zum Beispiel auf dem Intervall $D = [-\frac{1}{3}, \frac{1}{3}]$ mit $K = \frac{2}{3} < 1$ und es gilt $f(D) \subset D$. Mit jedem Startwert $x_0 \in D$ konvergiert also die Iteration $x_{n+1} = f(x_n) = x_n^2$ gegen den Fixpunkt $x_{\text{fix}} = 0$.

Für den anderen Fixpunkt $\tilde{x}_{\text{fix}} = 1$ gilt $f'(\tilde{x}_{\text{fix}}) = 2 > 1$ und wir können Satz 4.12 nicht anwenden. Tatsächlich divergiert die Folge x_n gegen ∞ für jeden Startwert $x_0 > 1$.

Nun gibt es einen Trick, und zwar kann die Umkehrfunktion betrachtet werden (falls f umkehrbar ist).

$$x = f(x) \quad \implies \quad f^{-1}(x) = f^{-1}(f(x)) = x$$

Die Fixpunkte von f und f^{-1} sind identisch und wir können den Fixpunktsatz auf f^{-1} (anstelle von f) anwenden. Ist der Funktionsgraph von f steil, dann ist der Graph von f^{-1} flach und umgekehrt.



Die Funktion $f : \mathbb{R}_{\geq 0} \rightarrow \mathbb{R}_{\geq 0}$, $f(x) = x^2$ ist umkehrbar mit $f^{-1}(x) = \sqrt{x}$. Für f^{-1} sind nun tatsächlich die Voraussetzungen von Satz 4.12 im Intervall $D = [\frac{1}{3}, 100]$ beispielsweise erfüllt. Mit der Fixpunktiteration $x_{n+1} = f^{-1}(x_n) = \sqrt{x_n}$ erhalten wir somit den Fixpunkt $\tilde{x}_{\text{fix}} = 1$, und zwar für jeden Startwert in D .

5 Integration

Es gibt zwei verschiedene Arten von Integration: Bei der unbestimmten Integration wird eine Stammfunktion gesucht, bei der bestimmten Integration geht es um die Berechnung eines Flächeninhalts. Der Zusammenhang dieser beiden Arten wird im Hauptsatz der Differential- und Integralrechnung ersichtlich.

5.1 Das unbestimmte Integral

Definition Eine Funktion $F(x)$ heisst *Stammfunktion* von $f(x)$, falls

$$F'(x) = f(x) .$$

Beispiel

Sei $f(x) = 3x^2$. Dann ist $F(x) = x^3$ eine Stammfunktion von $f(x)$. Dies ist jedoch nicht die einzige, denn auch $F(x) = x^3 + 10$ oder $F(x) = x^3 + \pi^{2025}$ ist eine. Die allgemeine Stammfunktion ist $F(x) = x^3 + c$ für eine Konstante c .

Definition Die Menge aller Stammfunktionen von $f(x)$ nennt man *unbestimmtes Integral* und schreibt

$$\int f(x) dx = F(x) + c .$$

Wichtige Beispiele

$$\begin{aligned} \int x^n dx &= \frac{1}{n+1} x^{n+1} + c \\ \int e^x dx &= e^x + c \\ \int \sin(x) dx &= -\cos(x) + c \\ \int \cos(x) dx &= \sin(x) + c \\ \int \frac{1}{x^2+1} dx &= \arctan(x) + c \\ \int \frac{1}{x} dx &= \ln(|x|) + c \end{aligned}$$

In der letzten Gleichung den Betrag nicht vergessen! Für $x > 0$ haben wir gezeigt, dass

$$(\ln(|x|))' = (\ln(x))' = \frac{1}{x} .$$

Für $x < 0$ wenden wir die Kettenregel an und erhalten

Also ist $\ln(|x|)$ eine Stammfunktion von $\frac{1}{x}$ für $x \neq 0$.

5.2 Das bestimmte Integral

Gegeben sei eine auf dem Intervall $[a, b]$ nicht-negative und beschränkte (bei Stetigkeit erfüllt) Funktion f , das heisst es gibt reelle Zahlen m und M , so dass für alle $x \in [a, b]$ die folgende Ungleichung gilt: $0 < m \leq f(x) \leq M$.

Frage: Wie gross ist der Flächeninhalt der zwischen der x -Achse, den vertikalen Grenzen $x = a$ und $x = b$ und der Kurve $y = f(x)$ eingeschlossenen Fläche?

Idee: Beschreibung der Fläche durch Rechtecke (näherungsweise).

Wir teilen deshalb das Intervall $[a, b]$ in n Teilintervalle gleicher Länge $l = \frac{b-a}{n}$ ein, indem wir $n-1$ Zwischenpunkte einfügen:

$$\underbrace{a = a + 0 \cdot l}_{=x_0} < \underbrace{a + 1 \cdot l}_{x_1} < \underbrace{a + 2 \cdot l}_{x_2} < \cdots < \underbrace{a + (n-1) \cdot l}_{x_{n-1}} < \underbrace{b = a + n \cdot l}_{x_n}.$$

Das Intervall wird also in die n Teilintervalle $[x_{i-1}, x_i]$ für $i = 1, \dots, n$ zerlegt.

Für jedes dieser Teilintervalle definieren wir das Minimum und das Maximum der Funktion f auf diesem Intervall

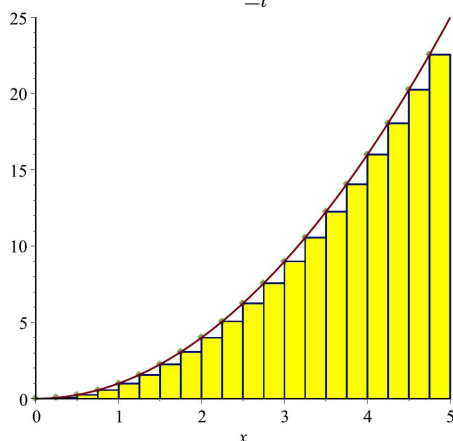
$$m_i = \text{Minimum von } f \text{ auf } [x_{i-1}, x_i]$$

$$M_i = \text{Maximum von } f \text{ auf } [x_{i-1}, x_i]$$

und wir definieren die folgenden beiden Summen, die Näherungen für den gesuchten Flächeninhalt sind.

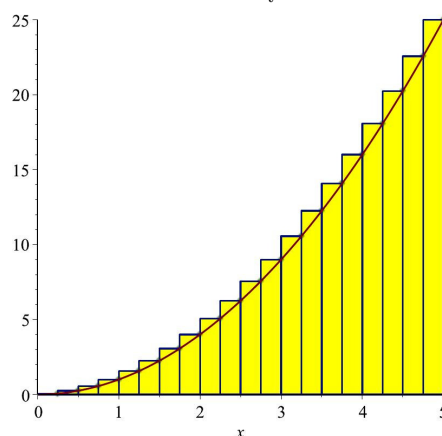
Untersumme

$$U_n = \sum_{i=1}^n m_i \underbrace{(x_i - x_{i-1})}_{=l}$$



Obersumme

$$O_n = \sum_{i=1}^n M_i \underbrace{(x_i - x_{i-1})}_{=l}$$



Mit dem globalen Maximum M und dem globalen Minimum m von f gilt somit

$$m(b-a) \leq U_n \leq \text{Flächeninhalt} \leq O_n \leq M(b-a).$$

Die beiden Folgen $(U_n)_{n \geq 1}$ und $(O_n)_{n \geq 1}$ sind also beschränkt. Ausserdem ist $(U_n)_{n \geq 1}$ monoton wachsend und $(O_n)_{n \geq 1}$ monoton fallend. Beide Folgen sind somit konvergent, müssen aber im Allgemeinen nicht den gleichen Grenzwert haben!

Definition Falls

$$\lim_{n \rightarrow \infty} U_n = \lim_{n \rightarrow \infty} O_n$$

gilt, so heisst dieser gemeinsame Grenzwert (der dann der gesuchte Flächeninhalt ist) das *bestimmte Integral* und man schreibt

$$\int_a^b f(x) dx.$$

Diese Schreibweise geht auf LEIBNIZ zurück. Das Integralzeichen $\int$ steht für $\int$ umme (mit unendlich vielen Summanden) und $dx = \Delta x = x_i - x_{i-1}$ ist ein “unendlich kleiner” Schritt.

Das bestimmte Integral wird nicht nur zur Berechnung des Flächeninhalts einer geometrischen Fläche benutzt. Der Flächeninhalt kann beispielsweise auch eine Weglänge oder einen Zuwachs bedeuten.

Beispiele

1. Wir fahren mit dem Velo zur Mathe-Vorlesung und bewegen uns mit der Geschwindigkeit $v(t)$ fort. Dann haben wir den Weg

$$\int_{\text{Startzeit}}^{\text{Ankunftszeit}} v(t) dt$$

zurückgelegt.

2. Die Basler Bevölkerung wuchs in der ersten Hälfte dieses Jahres. Die Wachstumsrate betrug $g(t)$. Dann ist der Zuwachs in diesem Zeitraum gegeben durch

$$\int_{\text{1. Januar}}^{\text{30. Juni}} g(t) dt.$$

Da wir es meistens mit stetigen Funktionen zu tun haben, ist die folgende Tatsache sehr praktisch.

Satz 5.1 Ist die Funktion f stetig in $[a, b]$, dann gilt

$$\lim_{n \rightarrow \infty} U_n = \lim_{n \rightarrow \infty} O_n = \int_a^b f(x) dx.$$

Insbesondere existieren alle vorkommenden Grenzwerte.

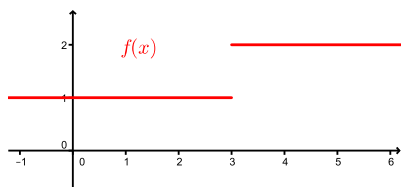
Eigenschaften des bestimmten Integrals

1. Es gilt

$$\int_a^b f(x)dx = \int_a^c f(x)dx + \int_c^b f(x)dx$$

für $c \in [a, b]$. Insbesondere existiert das Integral, auch wenn f in einzelnen Stellen nicht stetig ist (man nennt f in diesem Fall *stückweise stetig*).

Beispiel



Zum Beispiel ist

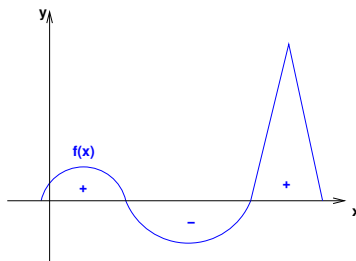
$$\int_1^5 f(x)dx = \int_1^3 f(x)dx + \int_3^5 f(x)dx ,$$

da f in $x = 3$ nicht stetig ist.

2. Nimmt die Funktion auch negative Werte an, so ist das Folgende zu beachten. Das bestimmte Integral

$$\int_a^b f(x)dx$$

ist der Flächeninhalt der Fläche zwischen der Kurve und der x -Achse, wobei der Flächeninhalt von Flächenstücken unterhalb der x -Achse negativ gezählt wird.



Beispiel

$$\int_0^{2\pi} \sin(x) dx$$

3. Man definiert

$$\int_a^b f(x)dx = - \int_b^a f(x)dx .$$

4. Es gilt

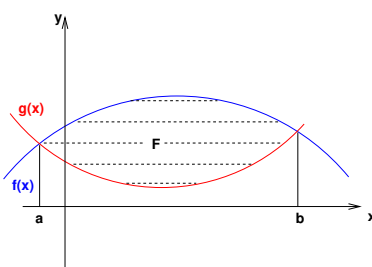
$$\int_a^a f(x)dx = 0 .$$

5. Es gilt

$$\begin{aligned}\int_a^b \lambda \cdot f(x) dx &= \lambda \cdot \int_a^b f(x) dx && \text{für } \lambda \in \mathbb{R} \\ \int_a^b (f(x) + g(x)) dx &= \int_a^b f(x) dx + \int_a^b g(x) dx\end{aligned}$$

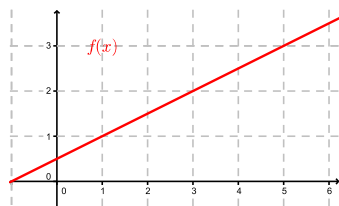
6. Für den Inhalt A der Fläche, die von den Graphen der beiden Funktionen $f(x)$ und $g(x)$ eingeschlossen wird, gilt

$$A = \int_a^b (f(x) - g(x)) dx .$$



Integralmittelwert

Beginnen wir mit einem Beispiel. Wir wollen das Integral der Funktion $f(x) = \frac{1}{2}x + \frac{1}{2}$ zwischen den Grenzen 1 und 5 bestimmen.



Satz 5.2 (Integralmittelwert) Sei f stetig in $[a, b]$. Dann gibt es ein $x_0 \in [a, b]$ mit

$$\int_a^b f(x) dx = f(x_0)(b - a) .$$

Die Bezeichnung Integralmittelwert kommt von der Gleichung im Satz dividiert durch $b - a$,

$$\frac{1}{b - a} \int_a^b f(x) dx = f(x_0) . \quad (\text{IM})$$

Beschreibt zum Beispiel $f(x)$ die Tagestemperatur für x zwischen $a = 0$ und $b = 24$ (Stunden), dann wird durch die linke Seite der Gleichung (IM) die mittlere Tagestemperatur berechnet. Die rechte Seite der Gleichung (IM) sagt weiter aus, dass diese mittlere Tagestemperatur zu einem bestimmten Zeitpunkt x_0 tatsächlich auch angenommen wird. Hierfür ist die Stetigkeit von f wesentlich.

Hauptsatz der Differential- und Integralrechnung

Ein bestimmtes Integral der Funktion f kann mit Hilfe einer Stammfunktion von f berechnet werden.

Beispiel

Sei $f(t)$ die Menge an R  pli (= Konfetti), die pro (infinitesimal kleine) Zeiteinheit w  hrend des Cort  ge (= Fasnachtsumzug) auf die Mittlere Br  cke in Basel f  llt. Sei $F(t)$ die Gesamtmenge an R  pli (seit Beginn des Cort  ge zur Zeit $t = 0$) auf der Mittleren Br  cke. Dann gilt

$$F(t) = \int_0^t f(x) dx .$$

Die   nderungsrate der Gesamtmenge an R  pli entspricht dem Zuwachs pro Zeiteinheit

$$F'(t) = f(t) .$$

Die Menge an R  pli, die in einem Zeitintervall $[t_1, t_2]$ auf die Mittlere Br  cke f  llt, entspricht der Differenz der Gesamtmengen,

$$\int_{t_1}^{t_2} f(x) dx = F(t_2) - F(t_1) .$$

Die obigen drei Gleichungen gelten f  r allgemeine Funktionen f mit einer Stammfunktion F . Integrieren ist also die Umkehrung des Ableitens.

Satz 5.3 (Hauptsatz der Differential- und Integralrechnung)

Sei f stetig in $[a, b]$. Dann gilt

$$\int_a^b f(x) dx = F(b) - F(a)$$

f  r eine beliebige Stammfunktion F von f , das heisst, f  r F mit $F'(x) = f(x)$.

Man schreibt

$$\int_a^b f(x) dx = F(x) \Big|_a^b .$$

F  r die Herleitung dieses Satzes zeigt man zun  chst, dass

$$F_a(x) = \int_a^x f(t) dt$$

eine Stammfunktion von f ist, und zwar diejenige, f  r die $F_a(a) = 0$ gilt. Mit einer kurzen Rechnung erh  lt man daraus die Behauptung des Satzes.

Beispiel

$$\int_0^\pi (2 \cos(x) - e^x) dx =$$

5.3 Integrationstechniken

In Kapitel 4 über die Differentiation haben wir praktische Regeln kennengelernt, mit welchen jede aus elementaren Funktionen zusammengesetzte Funktion abgeleitet werden kann. Für die Integration gibt es auch Regeln und Tricks (die wir hier einüben wollen), mit welchen man gewisse Integrale von Hand gut berechnen kann. Doch leider gibt es zahlreiche Integrale, die trotz Anwendung all dieser Tricks nicht berechenbar sind. Diese besitzen nachweislich keine aus den elementaren Funktionen zusammengesetzte Stammfunktion. In diesen Fällen liefert das Integral

$$F_a(x) = \int_a^x f(t) dt$$

eine neue Funktion. Zum Beispiel ist der *Integralsinus* $\text{Si}(x)$ für $x \geq 0$ definiert durch

$$\text{Si}(x) = \int_0^x \frac{\sin(t)}{t} dt$$

(vgl. Seite 89) und die *Fehlerfunktion* $\text{erf}(x)$ für $x \geq 0$ ist definiert durch

$$\text{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt$$

(vgl. Seite 89 und nächstes Semester).

Stossen Sie also eines Tages auf ein Integral, welches Sie nicht mit den üblichen Tricks lösen können, sollten Sie nicht zu lange zögern, in eine Formelsammlung zu schauen. Finden Sie dort Ihr Integral nicht, versuchen Sie es mit einem CAS (Computeralgebrasystem wie zum Beispiel Maple oder Mathematica) oder fragen Sie direkt eine*n Mathematiker*in (oder besser eine*n theoretische*n Physiker*in).

Schliesslich sagt man:

Ableiten ist ein Handwerk, Integrieren ist eine Kunst.

Partielle Integration

Seien $u = u(x)$ und $v = v(x)$ zwei differenzierbare Funktionen. Aus der Produktregel folgt

$$(uv)' = u'v + uv',$$

das heisst,

$$u'v = (uv)' - uv'.$$

Integrieren ergibt

$$\int u'v dx = uv - \int uv' dx.$$

Satz 5.4

$$\int_a^b u'(x) v(x) dx = u(x) v(x) \Big|_a^b - \int_a^b u(x) v'(x) dx$$

Ist also die zu integrierende Funktion $f(x) = f_1(x)f_2(x)$ ein Produkt von zwei Funktionen f_1, f_2 und das Produkt $F_1(x)f_2'(x)$ oder $f_1'(x)F_2(x)$, wobei F_1, F_2 Stammfunktionen von f_1 , bzw. f_2 sind, einfacher als $f(x)$ zu integrieren, dann ist partielle Integration empfehlenswert.

Beispiele

1. $\int x e^x dx = ?$

Durch Ableiten können wir das Resultat kontrollieren:

$$(x e^x - e^x + c)' =$$

2. $\int x^2 \cos(x) dx = ?$

3. $\int \ln(x) dx = ?$

Hier wenden wir einen Trick an, und zwar wählen wir $u' = 1$ und $v = \ln(x)$. Dann ist $u = x$ und $v' = \frac{1}{x}$. Es folgt

$$\int \ln(x) dx = \int 1 \cdot \ln(x) dx = x \ln(x) - \int x \frac{1}{x} dx = x \ln(x) - x + c .$$

4. $I = \int_0^{\frac{\pi}{2}} \sin^2(x) dx = ?$

Wegen $\sin^2(x) + \cos^2(x) = 1$ gilt $\cos^2(x) = 1 - \sin^2(x)$. Es folgt

Substitution

Sei F eine Stammfunktion der Funktion f . Dann gilt nach der Kettenregel

$$(F(g(x)))' = F'(g(x)) \cdot g'(x) = f(g(x)) \cdot g'(x) .$$

Durch Integration erhalten wir

$$\int f(g(x)) \cdot g'(x) dx = F(g(x)) + c .$$

Wir können auch $u = g(x)$ substituieren. Es gilt dann (formal) $g'(x) = \frac{du}{dx}$, und damit $g'(x)dx = du$.

Satz 5.5 *Es gilt*

$$\int f(g(x)) \cdot g'(x) dx = \int f(u) du ,$$

wobei $u = g(x)$ substituiert wurde.

Ist also die zu integrierende Funktion ein Produkt von zwei Funktionen, wobei die eine Funktion zusammengesetzt und die andere die Ableitung der inneren Funktion ist, dann ist Substitution empfehlenswert.

Beispiel

$$\int 2x \cos(x^2) dx = ?$$

Für das bestimmte Integral gilt mit den Bemerkungen vor Satz 5.5, dass

$$\int_a^b f(g(x)) \cdot g'(x) dx = F(g(x)) \Big|_a^b = F(g(b)) - F(g(a)) = F(u) \Big|_{g(a)}^{g(b)} = \int_{g(a)}^{g(b)} f(u) du.$$

Wir können also $u = g(x)$ substituieren wie vorher, doch müssen wir die Integrationsgrenzen entsprechend anpassen.

Satz 5.6 *Es gilt*

$$\int_a^b f(g(x)) \cdot g'(x) dx = \int_{g(a)}^{g(b)} f(u) du$$

mit $u = g(x)$.

Beispiele

1. $\int_{\frac{\pi}{2}}^{\pi} 2x \cos(x^2) dx = ?$ Wie oben setzen wir $u = g(x) = x^2$, und damit ist $du = 2x dx$.

Dies stimmt überein mit der Berechnung via unbestimmtes Integral vom Beispiel oben.

2. $\int_0^{\frac{\pi}{2}} \frac{\sin x}{\sqrt{1 + \cos x}} dx = ?$

3. $\int_e^{e^3} \frac{1}{x} \ln(\ln x) dx = ?$

4. Manchmal ist eine Substitution $u = g(x)$ hilfreich, obwohl $g'(x)$ nicht explizit im Integranden vorkommt.

$$\int_0^4 e^{\sqrt{x}} dx = ?$$

Partialbruchzerlegung

Ziel der Partialbruchzerlegung ist, eine rationale Funktion auf eine Linearkombination der folgenden Bestandteile zu bringen:

$$x^n, \quad \frac{1}{(x-a)^n}, \quad \frac{2x+p}{(x^2+px+q)^n}, \quad \frac{1}{x^2+1}$$

für $n \geq 1$. Die Stammfunktionen dieser Funktionen sind nämlich bekannt. Es gilt

$$\begin{aligned} \int \frac{1}{(x-a)^n} dx &= \frac{-1}{(n-1)(x-a)^{n-1}} + c \quad (n \geq 2), & \int \frac{1}{x-a} dx &= \ln(|x-a|) + c, \\ \int \frac{2x+p}{x^2+px+q} dx &= \ln(|x^2+px+q|) + c, & \int \frac{1}{x^2+1} dx &= \arctan(x) + c. \end{aligned}$$

Um das Vorgehen zu verstehen, betrachten wir hier zwei typische Beispiele. Weitere Beispiele sind in den Übungs(zusatz)aufgaben zu finden.

Beispiele

$$1. \quad \frac{x+1}{x^2-4} = \frac{x+1}{(x-2)(x+2)}$$

Zuerst faktorisiert man das Nennerpolynom (wenn möglich). Wegen der beiden Linearfaktoren $x-2$ und $x+2$ erwartet man Ausdrücke der Gestalt $\frac{a}{x-2}$ und $\frac{b}{x+2}$ für reelle Zahlen a, b . Wir machen daher den Ansatz

$$\frac{x+1}{x^2-4} = \frac{a}{x-2} + \frac{b}{x+2}.$$

Es gilt also $a = \frac{3}{4}$, $b = \frac{1}{4}$ und wir finden

$$\begin{aligned}\int \frac{x+1}{x^2-4} dx &= \frac{1}{4} \left(\int \frac{3}{x-2} dx + \int \frac{1}{x+2} dx \right) \\ &= \frac{1}{4} (3 \ln |x-2| + \ln |x+2|) + c = \frac{1}{4} \ln(|x-2|^2 |x+2|) + c.\end{aligned}$$

Alternative Methode zur Bestimmung von a und b :

Wir gehen vom gleichen Ansatz wie oben aus. Nun multiplizieren wir diese Gleichung mit $x-2$ und setzen anschliessend $x=2$.

Zur Bestimmung von b multiplizieren wir den Ansatz mit $x+2$ und setzen anschliessend $x=-2$.

2. Ziel ist die Bestimmung einer Stammfunktion von

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

Hier gilt

$$\text{Grad}(\text{Zählerpolynom}) \geq \text{Grad}(\text{Nennerpolynom}).$$

In diesem Fall führen wir zuerst eine Polynomdivision durch. Wir erhalten

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

Für den Bruch machen wir nun wieder einen Ansatz.

Für die Stammfunktion von f finden wir damit

Integration einer Potenzreihe

Wir haben in Satz 4.1 (Seite 53) gesehen, dass eine Potenzreihe gliedweise abgeleitet werden kann. Analog kann eine Potenzreihe $f(x) = \sum_{k=0}^{\infty} a_k x^k$ gliedweise integriert werden. Ist $f(x)$ konvergent auf einem Intervall I , dann ist die auf diese Weise erhaltene Stammfunktion von f ebenfalls konvergent auf I .

Beispiele

1. Mit Hilfe der Potenzreihe von $\sin(t)$ wollen wir den Integralsinus

$$\text{Si}(x) = \int_0^x \frac{\sin(t)}{t} dt$$

durch eine Potenzreihe beschreiben. Wir dividieren die Potenzreihe von $\sin(t)$ gliedweise durch t und erhalten

$$\frac{\sin t}{t} = \frac{1}{t} \left(t - \frac{t^3}{3!} + \frac{t^5}{5!} - \cdots + \cdots \right) = 1 - \frac{t^2}{3!} + \frac{t^4}{5!} - \cdots + \cdots.$$

Diese Reihe ist konvergent für alle $t \in \mathbb{R}$. Nun integrieren wir gliedweise. Dies ergibt

$$\int \frac{\sin t}{t} dt = t - \frac{t^3}{3 \cdot 3!} + \frac{t^5}{5 \cdot 5!} - \cdots + \cdots + c.$$

Damit gilt für $x \geq 0$

$$\text{Si}(x) = \int_0^x \frac{\sin(t)}{t} dt = x - \frac{x^3}{3 \cdot 3!} + \frac{x^5}{5 \cdot 5!} - \cdots + \cdots.$$

2. Analog schreiben wir für die Fehlerfunktion $\text{erf}(x)$ die Funktion e^{-t^2} als Potenzreihe:

$$e^x = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \cdots \quad \implies \quad e^{-t^2} = 1 - t^2 + \frac{t^4}{2!} - \frac{t^6}{3!} + \cdots$$

Gliedweises integrieren ergibt

$$\int e^{-t^2} dt = t - \frac{t^3}{3} + \frac{t^5}{5 \cdot 2!} - \frac{t^7}{7 \cdot 3!} + \cdots - \cdots + c$$

und wir erhalten für $x \geq 0$

$$\text{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt = \frac{2}{\sqrt{\pi}} \left(x - \frac{x^3}{3} + \frac{x^5}{5 \cdot 2!} - \frac{x^7}{7 \cdot 3!} + \cdots \right).$$

Wie berechne ich ein Integral $\int_a^b f(x) dx$?

- *Direkt*, das heisst mit Hilfe einer *Stammfunktion* F von f und Satz 5.3.
- Wenn $f(x) = f_1(x)f_2(x)$ ein Produkt von zwei Funktionen f_1, f_2 mit Stammfunktionen F_1 , bzw. F_2 ist, kann man es versuchen mit
 - *partieller Integration* (Satz 5.4), falls $F_1(x)f_2'(x)$ oder $f_1'(x)F_2(x)$ einfacher als $f(x)$ zu integrieren ist,
 - *Substitution* (Satz 5.6), falls $f_1(x) = g(h(x))$ eine zusammengesetzte Funktion und $f_2(x) = h'(x)$ die Ableitung der inneren Funktion ist.
- Wenn $f(x) = \frac{p(x)}{q(x)}$ eine rationale Funktion ist
 - und $q'(x) = p(x)$, dann ist $F(x) = \ln |q(x)|$ eine *Stammfunktion* von f ,
 - und $q'(x) \neq p(x)$, dann kann man es mit einer *Partialbruchzerlegung* (Seite 87) versuchen.
- Wenn $f(x)$ als Potenzreihe darstellbar ist, kann gliedweise integriert werden (Seite 89).

5.4 Uneigentliche Integrale

Gegeben sei eine auf dem rechts (bzw. links) offenen Intervall $[a, b)$ (resp. $(a, b]$) definierte und stetige Funktion $f(x)$. Wir wollen den Begriff des bestimmten Integrals erweitern, um eine Möglichkeit zu haben,

- Funktionen f , die bei der Annäherung $x \rightarrow b$ (resp. $x \rightarrow a$) nicht beschränkt sind, und
- Funktionen f über unbeschränkte Integrationsintervalle $[a, \infty)$ (resp. $(-\infty, b]$)

zu integrieren. Dazu definieren wir zunächst die folgenden vier Ausdrücke.

Definition

$$f(x) \text{ für } x \rightarrow b \text{ nicht beschränkt: } \int_a^b f(x) dx = \lim_{r \uparrow b} \int_a^r f(x) dx$$

$$f(x) \text{ für } x \rightarrow a \text{ nicht beschränkt: } \int_a^b f(x) dx = \lim_{r \downarrow a} \int_r^b f(x) dx$$

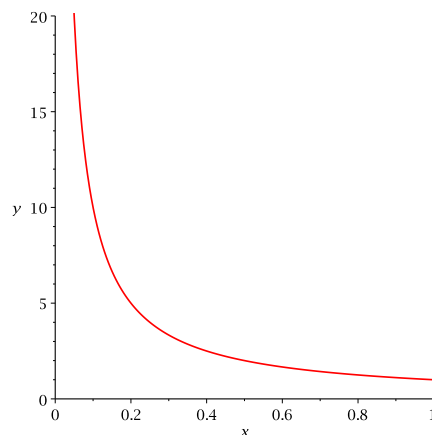
$$\text{unbeschränktes Intervall } [a, \infty): \int_a^\infty f(x) dx = \lim_{r \rightarrow \infty} \int_a^r f(x) dx$$

$$\text{unbeschränktes Intervall } (-\infty, b]: \int_{-\infty}^b f(x) dx = \lim_{r \rightarrow -\infty} \int_r^b f(x) dx$$

Diese Integrale nennt man *uneigentliche Integrale*. Ein uneigentliches Integral *konvergiert* (bzw. *divergiert*), wenn der zugehörige Grenzwert existiert (bzw. nicht existiert).

Beispiele

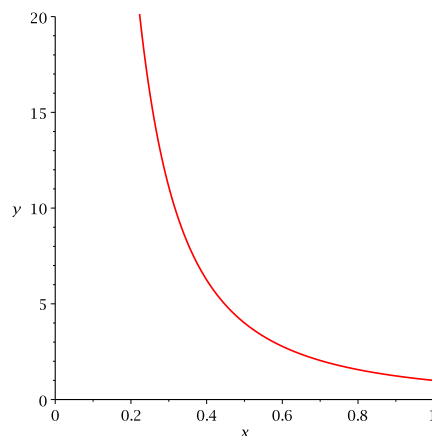
$$1. \int_0^1 \frac{1}{x} dx =$$



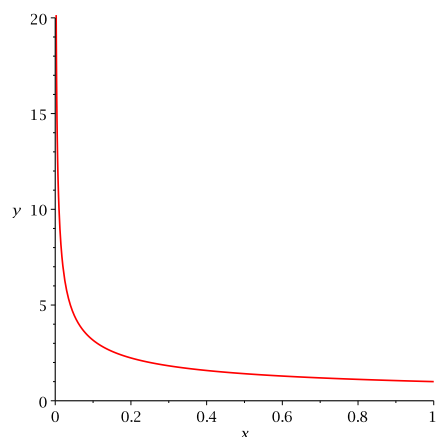
Zunächst ist nicht klar, was dieses Integral überhaupt bedeutet. Geometrisch misst das Integral den Flächeninhalt unter der Kurve der Funktion $f(x) = \frac{1}{x}$, welche allerdings im Nullpunkt eine Polstelle hat. Intuitiv könnten hier zwei Dinge passieren: Entweder der Flächeninhalt ist unendlich gross (da die Funktion unendlich wächst) oder der Flächeninhalt ist endlich (da das unendliche Wachstum der Funktion durch die schnelle Annäherung an die y -Achse kompensiert wird).

$$2. \int_0^1 \frac{1}{x^2} dx = \lim_{r \downarrow 0} \int_r^1 \frac{1}{x^2} dx = \infty, \text{ denn}$$

$$\int_r^1 \frac{1}{x^2} dx = -\frac{1}{x} \Big|_r^1 = -1 + \frac{1}{r} \quad \text{und} \quad \lim_{r \downarrow 0} \frac{1}{r} = \infty$$



3. $\int_0^1 \frac{1}{\sqrt{x}} dx =$

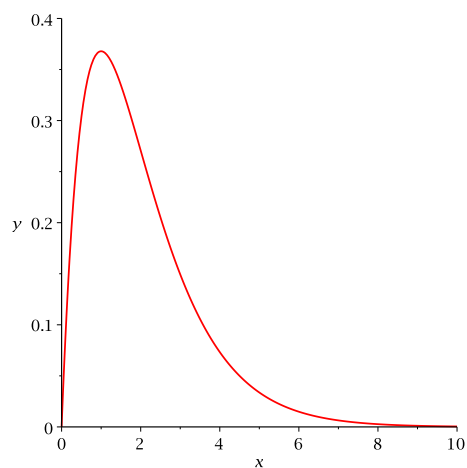


4. $\int_0^\infty x e^{-x} dx = \lim_{r \rightarrow \infty} \int_0^r x e^{-x} dx = 1$, denn

$$\int_0^r x e^{-x} dx = -(x+1)e^{-x} \Big|_0^r = -\frac{r+1}{e^r} + 1$$

wobei hier partiell integriert wurde wie im 1. Beispiel auf Seite 84. Mit der Regel von Bernoulli-de l'Hôpital folgt

$$\lim_{r \rightarrow \infty} -\frac{r+1}{e^r} = \lim_{r \rightarrow \infty} -\frac{1}{e^r} = 0.$$



Intuitiv könnten auch hier zwei Dinge passieren: Entweder der Flächeninhalt ist unendlich gross (da das Intervall unendlich lang ist) oder der Flächeninhalt ist endlich (da die unendliche Intervalllänge durch die schnelle Annäherung an die x -Achse kompensiert wird).

6 Differentialgleichungen

Definition Eine *Differentialgleichung* ist eine Gleichung, in der eine unbekannte Funktion $y = y(x)$ und Ableitungen (die erste oder auch höhere) von y vorkommen. Lösungen einer Differentialgleichung sind Funktionen y , welche die Gleichung erfüllen.

Gewisse Differentialgleichungen sind uns schon in Kapitel 5 begegnet, nämlich in der Form

$$y' = f(x)$$

wobei f eine gegebene reelle Funktion ist. Eine Lösung dieser Gleichung ist eine Stammfunktion $y = y(x) = F(x)$ von $f(x)$. Diese Lösung ist eindeutig bis auf eine Konstante c . Es gibt eine eindeutige Lösung, wenn zusätzlich eine sogenannte *Anfangsbedingung* vorgegeben ist, zum Beispiel $y(0) = 0$.

Differentialgleichungen kommen vor allem in der Physik vor, aber auch das Wachstumsverhalten von Populationen wird oft mit Hilfe von Differentialgleichungen beschrieben.

Beispiel

Fruchtfliegen vermehren sich im Sommer unter idealen Bedingungen besonders schnell. Eine Biologin stellt fest, dass für die Wachstumsrate $y'(t)$ von Fruchtfliegen zum Zeitpunkt t die Beziehung

$$y'(t) = 1,5 y(t)$$

gilt. Dabei bezeichnet $y(t)$ die Anzahl der Fruchtfliegen nach t Tagen. Am ersten Tag (zum Zeitpunkt $t = 0$) zählt sie 20 Fruchtfliegen. Mit wievielen Fruchtfliegen muss sie nach einer Woche rechnen?

Die Differentialgleichung im vorhergehenden Beispiel nennt man Differentialgleichung (kurz DGL) *erster Ordnung*, weil nur y und die erste Ableitung von y vorkommen. Allgemein nennt man eine Differentialgleichung von *n-ter Ordnung*, wenn die n -te Ableitung von y die höchste in der Differentialgleichung vorkommende Ableitung ist.

Differentialgleichungen von erster Ordnung sind zum Beispiel

$$y' = 0, \quad y' = f(x), \quad y' = ay + b, \quad y' = p(x)y + q(x), \quad y' = ay^2 + by + c.$$

Differentialgleichungen von zweiter Ordnung sind zum Beispiel

$$y'' = 0, \quad y'' + ay = 0, \quad y'' + by' + cy = \cos x.$$

6.1 Separierbare Differentialgleichungen

Exponentielles Wachstum

Die Gleichung

$$y' = y \tag{1}$$

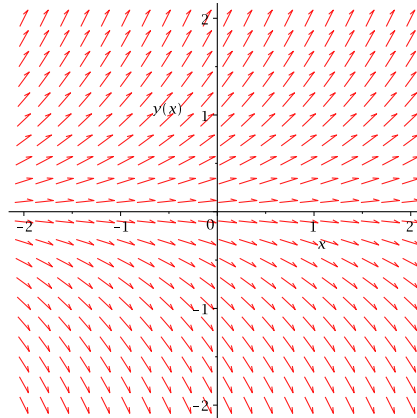
ist eine Differentialgleichung erster Ordnung. Genauer müsste man schreiben $f'(x) = f(x)$ für alle $x \in \mathbb{R}$. Um aber anzudeuten, dass die Funktion f variabel ist, schreiben wir in diesem Kapitel statt f immer y , also $y'(x) = y(x)$ für alle $x \in \mathbb{R}$. Und weil in dieser Gleichung y und nicht x gesucht ist, lassen wir in der Differentialgleichung wie oben das x einfach weg.

Eine Lösung der DGL (1) können wir erraten: Die Funktion

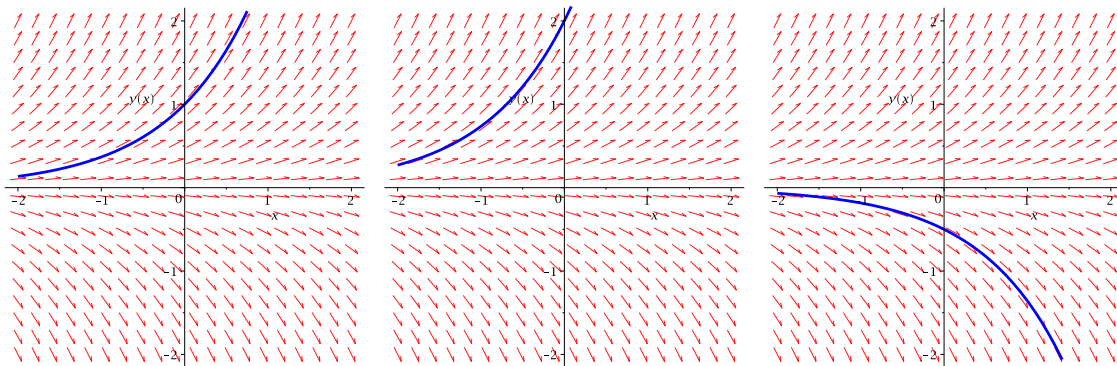
$$y(x) = e^x$$

erfüllt die DGL (1), denn $(e^x)' = e^x$. Ist dies die einzige Lösung?

Betrachten wir die DGL $y' = y$ geometrisch, indem wir in jedem Punkt (x, y) des Koordinatensystems $y' = y$ als Steigung einzeichnen. Das heisst, wir zeichnen in (x_0, y_0) einen Geradenabschnitt der Geraden durch (x_0, y_0) mit der Steigung y_0 ein.



Dies nennt man ein *Richtungsfeld*. Eine Lösung $y(x)$ der DGL $y' = y$ “passt” in dieses Richtungsfeld, denn die rot eingezeichneten Steigungen sind gerade Tangentenabschnitte an den Graphen von $y(x)$ (y' ist ja die Steigung der Tangente an die Funktion $y(x)$).



Im ersten Bild ist die Lösung $y = e^x$ eingezeichnet, denn für diese Lösung gilt die Anfangsbedingung

$$y(0) = e^0 = 1.$$

Im mittleren Bild gilt die Anfangsbedingung $y(0) = 2$, was durch die Funktion $y(x) = 2e^x$ erfüllt ist. Im Bild rechts sehen wir die Lösung $y(x) = -\frac{1}{2}e^x$, da dort die Anfangsbedingung $y(0) = -\frac{1}{2}$ gilt.

Die allgemeine Lösung der Differentialgleichung $y' = y$ ist also gegeben durch

$$y(x) = A e^x,$$

wobei die reelle Zahl A durch die Anfangsbedingung $y(0) = A e^0 = A$ bestimmt ist.

Betrachten wir nun die leicht allgemeinere Differentialgleichung

$$y' = ay \quad (2)$$

für eine reelle Zahl a . Für jede Konstante A (d.h. reelle Zahl A) ist die Funktion

$$y(x) = A e^{ax}$$

eine Lösung dieser DGL, und jede Lösung hat diese Form. Durch Ableiten können wir überprüfen, dass $y(x)$ eine Lösung ist:

Für $a > 0$ ist die Funktion e^{ax} monoton wachsend. Deshalb beschreibt die Differentialgleichung (2) ein exponentielles Wachstum falls $A > 0$ und einen exponentiellen Zerfall falls $A < 0$ ist, wobei die momentane Wachstums-, bzw. Zerfallsgeschwindigkeit y' proportional vom Bestand y abhängt.

Beispiel

Die Wachstumsrate von Fruchtfliegen aus dem Beispiel zu Beginn des Kapitels ist genau von dieser Form, und zwar gilt

$$y'(t) = 1,5 y(t) ,$$

das heisst, $a = 1,5$. Die Lösung ist also von der Form

$$y(t) = A e^{1,5t} .$$

Mit der Anfangsbedingung $y(0) = 20$ können wir A und damit die Lösung $y(t)$ bestimmen:

Nach einer Woche muss die Biologin bei idealen Bedingungen also bereits mit

$$y(7) = 726310$$

Fruchtfliegen rechnen!

Beschränktes Wachstum

Wir verallgemeinern noch einmal und betrachten die Differentialgleichung

$$y' = ay + b \quad (3)$$

für reelle Zahlen $a \neq 0$ und b .

Um eine Lösung dieser Differentialgleichung zu finden, schreiben wir

$$\frac{dy}{dx} = y'(x) = ay + b .$$

Nun bringen wir die Variablen y auf die linke Seite (dabei setzen wir voraus, dass $y \neq -\frac{b}{a}$) und die Variablen x auf die rechte Seite und integrieren beide Seiten.

$$\frac{dy}{ay + b} = dx \quad \implies \quad \int \frac{dy}{ay + b} = \int dx .$$

Wir erhalten

und lösen diese Gleichung nach y auf:

Zu Beginn dieses Beispiels hatten wir vorausgesetzt, dass $y \neq -\frac{b}{a}$. Aber $y = -\frac{b}{a}$ ist auch eine Lösung der DGL (3), denn die Ableitung einer konstanten Funktion ist gleich Null.

Die Differentialgleichung (3) hat also die allgemeine Lösung

$$y(x) = A e^{ax} - \frac{b}{a}$$

für eine beliebige Konstante A .

Bei vielen Wachstumsprozessen in der Natur ist die Zu- oder Abnahme eines Bestandes durch eine natürliche Grenze (man nennt sie *Sättigungs-* oder *Kapazitätsgrenze*) beschränkt. Zum Beispiel kann ein Teich nicht unendlich viele Fische aufnehmen oder eine warme Flüssigkeit kann sich nicht unendlich stark abkühlen. Die Wachstumsgeschwindigkeit ist in diesen Fällen proportional zur Differenz aus Sättigungsgrenze S und Bestand, das heisst, man erhält eine Differentialgleichung der Form

$$y' = a(S - y) = -ay + aS$$

für eine Konstante a . Diese Differentialgleichung ist vom Typ der DGL (3) und hat damit die allgemeine Lösung

$$y(x) = S + A e^{-ax}.$$

Für $a > 0$ ist die Funktion e^{-ax} monoton fallend. Deshalb beschreibt $y(x)$ einen Wachstumsprozess, falls $A < 0$ und einen Zerfallsprozess, falls $A > 0$.

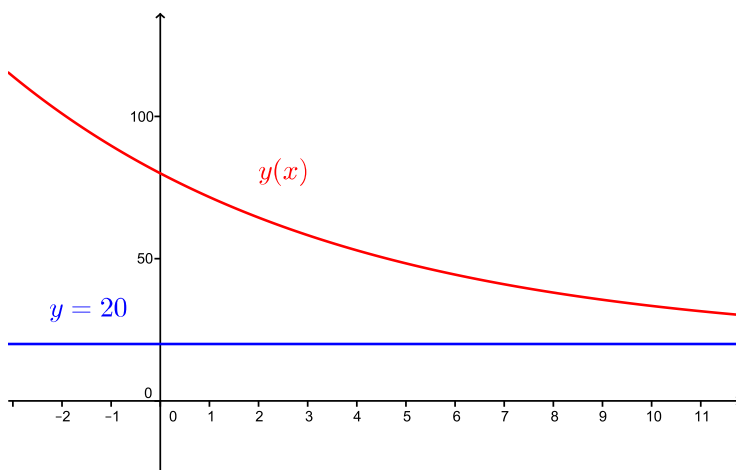
Beispiel

Frisch aufgebrühter Kaffee hat eine Temperatur um die 80° C. Als angenehme Trinktemperatur gilt etwa 45° C. Wir lassen den Kaffee im 20° C warmen Zimmer stehen. Pro Minute kühlt der Kaffee um 15 Prozent der aktuellen Temperaturdifferenz zur Raumtemperatur ab. Nach wieviel Minuten können wir den Kaffee trinken, ohne uns die Zunge zu verbrennen?

Sei $y = y(x)$ die Temperatur des Kaffees in $^{\circ}\text{C}$ nach x Minuten. Dann gilt

Für die eindeutige Lösung erhalten wir

Hier ist der Graph von $y(x) = 20 + 60e^{-0,15 \cdot x}$:



Für die optimale Trinktemperatur von 45°C müssen wir nun die Gleichung $y(x) = 45$ nach x auflösen:

Wir können unseren Kaffee also nach ungefähr 6 Minuten trinken.

Trennung der Variablen

Kehren wir nochmals zur allgemeinen Differentialgleichung (3) zurück. Die Methode, mit der wir diese Differentialgleichung gelöst haben, ist auch auf andere Differentialgleichungen anwendbar, nämlich auf sogenannte separierbare Differentialgleichungen.

Definition Eine Differentialgleichung der Form

$$y' = g(x)h(y)$$

für Funktionen $g(x)$ und $h(y)$ in x , bzw. y heisst *separierbar*.

Eine separierbare Differentialgleichung kann nach der folgenden Methode gelöst werden.

Trennung der Variablen

1. Schreibe

$$\frac{dy}{dx} = g(x)h(y) .$$

2. Trenne die Variablen, das heisst bringe y nach links und x nach rechts:

$$\frac{1}{h(y)} dy = g(x) dx$$

3. Integriere unbestimmt:

$$\int \frac{1}{h(y)} dy = \int g(x) dx$$

4. Löse nach y auf.

5. Jede Nullstelle $y = y_0$ von $h(y)$ ergibt zusätzlich eine konstante Lösung $y(x) = y_0$.
(Diese Lösungen wurden im 2. Schritt bei der Division durch $h(y)$ ausgeschlossen.)

Beispiel

Gesucht ist die Lösung der Differentialgleichung

$$y' = (y - 2)(x + 1)^2$$

mit der Anfangsbedingung $y(-1) = 2,5$. In der obigen Notation ist hier

$$g(x) = (x + 1)^2 \quad \text{und} \quad h(y) = y - 2 .$$

Diese DGL ist also separierbar und wir können sie durch Trennung der Variablen lösen.

Die ersten beiden Schritte ergeben

Im dritten Schritt wird integriert:

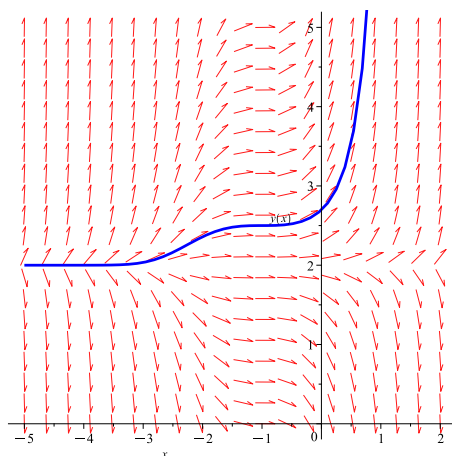
Und schliesslich müssen wir die Gleichung nach y auflösen:

Die Funktion $y(x) = 2$ ist hier wegen der Anfangsbedingung $y(-1) = 2,5$ keine Lösung.

Mit der Anfangsbedingung folgt

Die eindeutige Lösung der Differentialgleichung ist also

$$y(x) = \frac{1}{2} e^{\frac{1}{3}(x+1)^3} + 2.$$



Logistisches Wachstum

Wachstumsprozesse können zu verschiedenen Zeitpunkten unterschiedlich verlaufen. Es kann zum Beispiel vorkommen, dass eine Population zunächst exponentiell wächst, da die Sättigungsgrenze zunächst noch kein Hindernis für die noch kleine Population darstellt. Doch nähert sich die Grösse der Population der Sättigungsgrenze, dann verringert sich das Wachstum und es ist begrenzt. Man nennt diese Kombination von exponentiellem und begrenztem Wachstum *logistisches Wachstum*. Die Wachstumsgeschwindigkeit ist hier proportional zum Produkt vom Bestand und der Differenz aus Sättigungsgrenze und Bestand. Wir erhalten also eine Differentialgleichung der Form

$$y' = ay(S - y). \quad (4)$$

Ist die Population y noch klein, dann ist $S - y \approx S$ und $y' \approx aSy$ beschreibt ein exponentielles Wachstum. Je mehr sich y der Grenze S nähert, desto kleiner wird $S - y$, und die Änderungsrate von y nimmt immer mehr ab.

Beispiel

Wir beobachten das Wachstum einer Pantoffeltierchenpopulation im Labor mit konstanten Umweltbedingungen. Die Population wird durch die Gleichung

$$y' = 1,1 \cdot y \cdot \left(\frac{900 - y}{900} \right)$$

beschrieben. Welche Funktion y beschreibt die Anzahl der Pantoffeltierchen in Abhängigkeit der Zeit, wenn wir zur Zeit $x = 0$ ein einziges Pantoffeltierchen haben?

Diese Differentialgleichung beschreibt ein logistisches Wachstum mit der Sättigungsgrenze $S = 900$ und $a = \frac{1,1}{900}$.

Wir können die Differentialgleichung (4) (und damit auch die DGL des Beispiels) durch Trennung der Variablen lösen.

Das Integral auf der linken Seite berechnen wir mit Hilfe einer Partialbruchzerlegung. Wir machen den Ansatz

Es folgt

Damit erhalten wir

$$\ln \left(\frac{|y - S|}{|y|} \right) = -aSx + c$$

und wir müssen diese Gleichung noch nach y auflösen.

Die allgemeine Lösung der Differentialgleichung (4) ist also

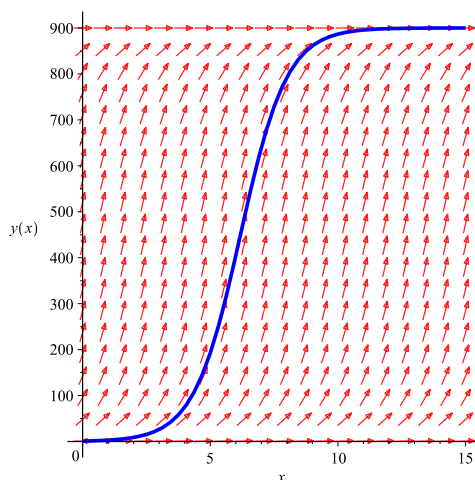
$$y(x) = \frac{S}{1 - Ae^{-aSx}}$$

für eine Konstante $A \neq 0$. Zusätzliche Lösungen sind die konstanten Funktionen $y(x) = 0$ und $y(x) = S$ (da hier $h(y) = ay(S - y) = 0$ für $y_0 = 0$ und $y_0 = S$).

Beispiel

Für das Wachstum der Pantoffeltierchenpopulation gilt $S = 900$ und $a = \frac{1,1}{900}$. Mit der Anfangsbedingung $y(0) = 1$ erhalten wir für A :

$$y(x) = \frac{900}{1 + 899 e^{-1,1 \cdot x}}.$$



Die Differentialgleichung (4) ist ein Spezialfall einer Differentialgleichung der Form

$$y' = ay^2 + by + c \quad (5)$$

für reelle Zahlen $a \neq 0$, b und c . Hat das quadratische Polynom auf der rechten Seite zwei verschiedene reelle Nullstellen y_1 , y_2 , dann können wir die Differentialgleichung (5) analog zur DGL (4) lösen. Die Geraden $y = y_1$ und $y = y_2$ bilden in diesem Fall Asymptoten für die Lösung $y(x)$. Der Fall $y_1 = y_2$ führt durch direkte Integration (ohne Partialbruchzerlegung) auf eine rationale Funktion y . Und hat das quadratische Polynom keine reelle Nullstelle, dann kommt mit Hilfe einer Substitution der Arcustangens ins Spiel (s. Zusatzübung).

6.2 Lineare Differentialgleichungen erster Ordnung

Definition Eine Differentialgleichung erster Ordnung heisst *linear*, wenn sie auf die Form

$$y' = p(x)y + q(x) \quad (\text{I})$$

für Funktionen $p(x)$ und $q(x)$ gebracht werden kann. Die DGL

$$y' = p(x)y \quad (\text{H})$$

heisst die zur Gleichung (I) zugehörige *homogene* Gleichung. Entsprechend nennt man die Gleichung (I) manchmal *inhomogen*.

Da die DGL (I) linear in y' und y ist, gilt das Folgende. Ist y_1 eine Lösung von (I) und y_0 eine Lösung von (H), dann ist die Summe $y_2 = y_1 + y_0$ wieder eine Lösung von (I). Sind umgekehrt y_1 und y_2 zwei Lösungen von (I), dann ist die Differenz $y_0 = y_1 - y_2$ eine Lösung von (H). Damit gilt der folgende Satz.

Satz 6.1 Die allgemeine Lösung von (I) erhält man durch Addition einer partikulären Lösung von (I) und der allgemeinen Lösung von (H).

Die allgemeine Lösung von (H) erhält man durch Trennung der Variablen, da (H) eine separierbare DGL ist. Wie wir eine *partikuläre* (d.h. einzelne) Lösung von (I) finden, schauen wir uns zuerst an einem Beispiel an.

Beispiel

Wir betrachten die DGL $xy' - 2y = x^3$. Dies ist eine lineare DGL erster Ordnung, denn wir können sie umschreiben zu

$$y' = \frac{2}{x}y + x^2. \quad (\text{I})$$

Es ist also $p(x) = \frac{2}{x}$ und $q(x) = x^2$. Die zugehörige homogene Gleichung ist

$$y' = \frac{2}{x}y. \quad (\text{H})$$

Diese homogene Gleichung (H) können wir durch Trennung der Variablen lösen:

$$\frac{dy}{dx} = \frac{2}{x}y \quad \implies \quad \int \frac{dy}{y} = \int \frac{2}{x} dx$$

Integration führt zu

$$\ln |y| = 2 \ln |x| + c = \ln(x^2) + c$$

und auflösen nach y ergibt

$$y = \pm e^{\ln(x^2)+c} = \pm e^c \cdot x^2 = A \cdot x^2 \quad \text{mit } A = \pm e^c \neq 0.$$

Hinzu kommt die konstante Lösung $y = 0$. Die allgemeine Lösung von (H) ist also

$$y_H(x) = Ax^2.$$

für eine beliebige Konstante A .

Um eine partikuläre Lösung der inhomogenen Gleichung (I) zu finden, machen wir nun einen Ansatz, der *Variation der Konstanten* genannt wird und erstmals von JOSEPH-LOUIS LAGRANGE (1736 – 1813) benutzt wurde. Wir nehmen die allgemeine Lösung $y_H(x) = Ax^2$ von (H) und ersetzen die Konstante A durch eine (noch unbekannte) Funktion $a(x)$, das heißt wir setzen

$$y(x) = a(x) \cdot x^2.$$

Diesen Ansatz setzen wir in die Gleichung (I) ein:

Linke Seite:

Rechte Seite:

Da “Linke Seite = Rechte Seite” wegen der Gleichung (I), folgt

Die Konstante c können wir weglassen, da wir nur eine einzelne Lösung von (I) suchen; also $a(x) = x$ genügt. Damit erhalten wir eine partikuläre Lösung der Gleichung (I),

$$y_P(x) = a(x) \cdot x^2 = x^3.$$

Nun addieren wir die partikuläre Lösung von (I) und die allgemeine Lösung von (H),

$$y(x) = y_P(x) + y_H(x) = x^3 + Ax^2.$$

Gemäss Satz 6.1 ist dies die allgemeine Lösung der inhomogenen Gleichung (I).

Dieser Ansatz mit der Variation der Konstanten lässt sich auf eine beliebige lineare DGL erster Ordnung (I) $y' = p(x)y + q(x)$ anwenden. Sei (H) $y' = p(x)y$ wie vorher.

Herleitung und Bestimmung der allgemeinen Lösung von (I)

Schritt 1: Bestimmung der allgemeinen Lösung y_H von (H) durch Trennung der Variablen.

$$\frac{dy}{dx} = y' = p(x) \cdot y \quad \implies \quad \int \frac{dy}{y} = \int p(x) dx.$$

Integration führt zu

$$\ln |y| = P(x) + c$$

für eine Stammfunktion $P(x)$ von $p(x)$. Auflösen nach y ergibt die allgemeine Lösung

$$y_H(x) = A e^{P(x)}$$

für eine beliebige Konstante A .

Schritt 2: Bestimmung einer partikulären Lösung y_P von (I) durch Variation der Konstanten.

Aufgrund der allgemeinen Lösung in Schritt 1 machen wir den Ansatz

$$y_P(x) = a(x) e^{P(x)}$$

und setzen ihn in die Gleichung (I) ein:

Es folgt

$$a'(x) e^{P(x)} = q(x) \quad \implies \quad a'(x) = q(x) e^{-P(x)}.$$

Ist $B(x)$ eine Stammfunktion von $q(x) e^{-P(x)}$, dann erhalten wir eine partikuläre Lösung

$$y_P(x) = B(x) e^{P(x)}.$$

Schritt 3: Gemäss Satz 6.1 ist nun **die allgemeine Lösung von (I)** gegeben durch

$$y(x) = y_P(x) + y_H(x) = B(x) e^{P(x)} + A e^{P(x)} = (B(x) + A) e^{P(x)},$$

wobei $P(x)$ und $B(x)$ Stammfunktionen von $p(x)$, bzw. $q(x) e^{-P(x)}$ sind.

Beispiele

1. Gesucht ist die allgemeine Lösung der inhomogenen linearen DGL

$$y' = y + x.$$

Hier ist also $p(x) = 1$ und $q(x) = x$.

Schritt 1: Eine Stammfunktion von $p(x) = 1$ ist $P(x) = x$. Also ist

$$y_H(x) = A e^{P(x)} = A e^x$$

die allgemeine Lösung der zugehörigen homogenen DGL $y' = y$. Das wissen wir auch vom Beginn dieses Kapitels (Seite 94).

Schritt 2: Für eine partikuläre Lösung brauchen wir eine Stammfunktion $B(x)$ von

$$q(x) e^{-P(x)} = x e^{-x}.$$

Mit partieller Integration wie im 1. Beispiel auf Seite 84 finden wir

$$B(x) = -(1 + x) e^{-x}.$$

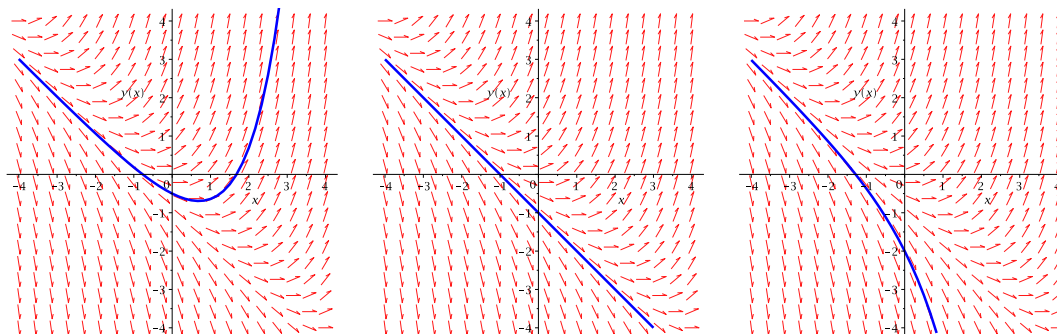
Damit erhalten wir die partikuläre Lösung

Schritt 3: Die allgemeine Lösung der DGL $y' = y + x$ ist demnach

$$y(x) = y_P + y_H = -1 - x + A e^x.$$

Die Bilder zeigen die eindeutigen Lösungen zu verschiedenen Anfangsbedingungen:

- Bild links: $y(0) = -\frac{1}{2} \implies A = \frac{1}{2}$
- Bild Mitte: $y(0) = -1 \implies A = 0$
- Bild rechts: $y(0) = -2 \implies A = -1$



2. Gesucht ist die eindeutige Lösung der inhomogenen linearen DGL

$$y' = \cos(x) y + e^{\sin(x)}$$

mit der Anfangsbedingung $y(\pi) = 0$. Hier ist also $p(x) = \cos(x)$ und $q(x) = e^{\sin(x)}$.

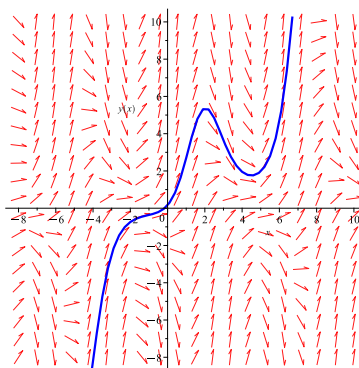
Eine Stammfunktion von $p(x) = \cos(x)$ ist

Für y_p brauchen wir eine Stammfunktion $B(x)$ von

Damit erhalten wir die allgemeine Lösung der DGL

Mit der Anfangsbedingung $y(\pi) = 0$ erhalten wir für die Konstante A die Gleichung

Die eindeutige Lösung ist also gegeben durch $y(x) = (x - \pi) e^{\sin(x)}$.



6.3 Lineare Differentialgleichungen zweiter Ordnung

Wir beginnen mit drei **Beispielen**.

1. Die DGL $y'' = 0$ hat die allgemeine Lösung

$$y(x) = Ax + B$$

für beliebige Konstanten A und B .

2. Die DGL $y'' = y$ hat die allgemeine Lösung

$$y(x) = Ae^x + Be^{-x}$$

für beliebige Konstanten A und B . Tatsächlich sind dies Lösungen, denn

3. Die DGL $y'' = -y$ hat die allgemeine Lösung

$$y(x) = A \cos(x) + B \sin(x)$$

für beliebige Konstanten A und B .

In allen drei Beispielen können zwei Konstanten A und B frei gewählt werden. Um eine eindeutige Lösung zu erhalten, müssen deshalb zwei Anfangsbedingungen vorgegeben werden, zum Beispiel $y(0)$ und $y'(0)$.

Definition Eine Differentialgleichung der Form

$$ay'' + by' + cy = 0 \tag{6}$$

mit $a \neq 0$ nennt man *homogene lineare Differentialgleichung zweiter Ordnung* (mit konstanten Koeffizienten).

Wir bemerken zuerst, dass für eine Lösung y von (6) auch Ay , für jede reelle Zahl A , eine Lösung ist, und sind y_1, y_2 zwei Lösungen von (6), dann ist auch die Summe $y_1 + y_2$ eine Lösung von (6). Die Lösungen von (6) bilden deshalb einen sogenannten Vektorraum, wie wir in Kapite 9 sehen werden.

Um eine Lösung von (6) zu finden, machen wir den Ansatz $y(x) = e^{\lambda x}$ und setzen ihn in die Gleichung (6) ein:

Da $e^{\lambda x} \neq 0$, erhalten wir die folgende Gleichung.

Definition Die Gleichung

$$a\lambda^2 + b\lambda + c = 0$$

nennt man *charakteristische Gleichung* von (6).

Wir erhalten also eine Lösung $y(x) = e^{\lambda x}$ von (6), wenn λ die charakteristische Gleichung erfüllt. Diese Gleichung hat entweder zwei verschiedene reelle Lösungen, eine reelle Lösung oder zwei konjugiert komplexe Lösungen. Wir untersuchen diese drei Fälle nacheinander.

1. Fall: Die Gleichung $a\lambda^2 + b\lambda + c = 0$ hat zwei verschiedene reelle Lösungen.

Dies ist genau dann der Fall, wenn $b^2 - 4ac > 0$. Die beiden reellen Lösungen sind dann gegeben durch

$$\lambda_1 = \frac{-b + \sqrt{b^2 - 4ac}}{2a} \quad \text{und} \quad \lambda_2 = \frac{-b - \sqrt{b^2 - 4ac}}{2a}.$$

Die allgemeine Lösung der DGL (6) ist in diesem Fall

$$y(x) = Ae^{\lambda_1 x} + Be^{\lambda_2 x}.$$

Da λ_1 und λ_2 Lösungen der charakteristischen Gleichung sind, sind $e^{\lambda_1 x}$ und $e^{\lambda_2 x}$ Lösungen der DGL (6). Nach der Bemerkung oben ist dann auch jede Linearkombination eine Lösung.

Beispiel

Gesucht ist die Lösung der DGL

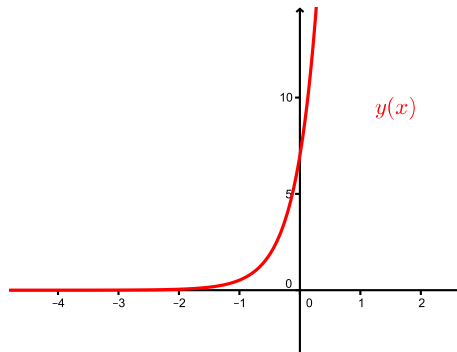
$$2y'' - 10y' + 12y = 0$$

mit den Anfangsbedingungen $y(0) = 7$ und $y'(0) = 19$.

Charakteristische Gleichung:

Mit den Anfangsbedingungen gilt

Die (eindeutige) Lösung ist also $y(x) = 2e^{2x} + 5e^{3x}$.

**2. Fall: Die Gleichung $a\lambda^2 + b\lambda + c = 0$ hat eine reelle Lösung.**

Dies ist genau dann der Fall, wenn $b^2 - 4ac = 0$. Die Lösung ist dann gegeben durch

$$\lambda_0 = \frac{-b}{2a}.$$

Die allgemeine Lösung der DGL (6) ist in diesem Fall

$$y(x) = (Ax + B)e^{\lambda_0 x}.$$

Dass $e^{\lambda_0 x}$ eine Lösung ist, ist klar von der Herleitung. Dass $xe^{\lambda_0 x}$ eine Lösung ist, kann man durch Einsetzen in die Gleichung (6) überprüfen.

Beispiel

Gesucht ist die Lösung der DGL

$$y'' - 10y' + 25y = 0$$

mit den Anfangsbedingungen $y(0) = 3$ und $y'(0) = 13$.

Die charakteristische Gleichung lautet:

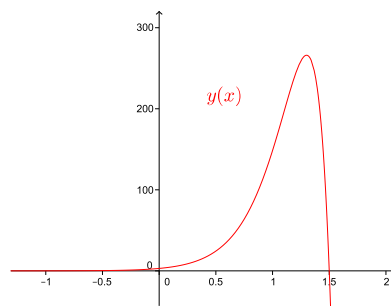
$$0 = \lambda^2 - 10\lambda + 25 = (\lambda - 5)^2 \quad \implies \quad \lambda_0 = 5$$

Die allgemeine Lösung ist also gegeben durch

$$y(x) = (Ax + B)e^{5x}.$$

Von der ersten Anfangsbedingung erhalten wir direkt $3 = y(0) = B$. Für die zweite Anfangsbedingung müssen wir zuerst $y(x)$ ableiten.

Die (eindeutige) Lösung ist also $y(x) = (-2x + 3)e^{5x}$.



3. Fall: Die Gleichung $a\lambda^2 + b\lambda + c = 0$ hat zwei konjugiert komplexe Lösungen.

Dies ist genau dann der Fall, wenn $b^2 - 4ac < 0$. Die beiden konjugiert komplexen Lösungen sind dann gegeben durch

$$\lambda_1 = \alpha + i\omega \quad \text{und} \quad \lambda_2 = \alpha - i\omega, \quad \text{wobei} \quad \alpha = -\frac{b}{2a}, \quad \omega = \frac{\sqrt{4ac - b^2}}{2a}.$$

Die allgemeine Lösung der DGL (6) ist in diesem Fall

$$y(x) = e^{\alpha x} (A \sin(\omega x) + B \cos(\omega x)).$$

Dass $e^{\alpha x} \sin(\omega x)$ und $e^{\alpha x} \cos(\omega x)$ Lösungen von (6) sind, folgt daraus, dass

$$\begin{aligned} e^{\lambda_1 x} &= e^{(\alpha + i\omega)x} = e^{\alpha x} e^{i\omega x} = e^{\alpha x} (\cos(\omega x) + i \sin(\omega x)) \\ e^{\lambda_2 x} &= e^{(\alpha - i\omega)x} = e^{\alpha x} e^{-i\omega x} = e^{\alpha x} (\cos(\omega x) - i \sin(\omega x)) \end{aligned}$$

Lösungen von (6) sind und die DGL linear in y, y' und y'' ist.

Beispiel

Gesucht ist die Lösung der DGL

$$y'' - 4y' + 13y = 0$$

mit den Anfangsbedingungen $y(0) = 3$ und $y'(0) = 9$.

Charakteristische Gleichung:

Die allgemeine Lösung ist also gegeben durch

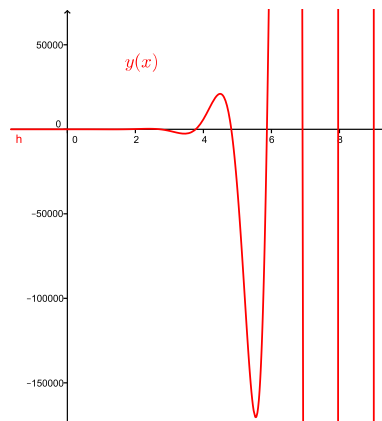
$$y(x) = e^{2x}(A \sin(3x) + B \cos(3x)) .$$

Die erste Anfangsbedingung gibt wieder direkt $3 = y(0) = B$. Für die zweite Anfangsbedingung müssen wir wieder zuerst $y(x)$ ableiten.

Die (eindeutige) Lösung ist also

$$y(x) = e^{2x}(\sin(3x) + 3 \cos(3x)) = \sqrt{10} e^{2x}(\sin(3x + u))$$

mit $u = \arccos\left(\frac{1}{\sqrt{10}}\right) \approx 1,25$ (vgl. Satz 1.3).



Physikalische Anwendung: Das Federpendel

Wir betrachten einen an einer Spiralfeder aufgehängten Körper der Masse m . Die Funktion $y(t)$ soll die Auslenkung des Körpers aus der Ruhelage zum Zeitpunkt t beschreiben.

Experimente zeigen, dass bei kleinen Auslenkungen aus der Ruhelage die rücktreibende Kraft

proportional zur Auslenkung ist. Der Proportionalitätsfaktor $k > 0$ hängt nur von der Feder ab und heisst *Federkonstante*. Da diese Kraft entgegen der Auslenkung wirkt, können wir sie durch $-ky$ angeben. Die Auslenkung wird zudem durch eine Reibungskraft (Luftwiderstand) gebremst. Diese ist proportional zur Geschwindigkeit $y'(t)$. Den Proportionalitätsfaktor bezeichnen wir mit ρ . Da die Kraft gleich Masse m mal Beschleunigung $y''(t)$ ist, erhalten wir die Differentialgleichung

$$my'' = -ky - \rho y' ,$$

die wir umschreiben zu

$$my'' + \rho y' + ky = 0 .$$

Dies ist eine homogene lineare DGL zweiter Ordnung. Die zugehörige charakteristische Gleichung lautet

$$m\lambda^2 + \rho\lambda + k = 0 .$$

Wir nehmen nun an, dass die Reibung klein ist, das heisst, es gelte $\rho^2 < 4mk$. Damit sind wir im 3. Fall. Für die allgemeine Lösung erhalten wir

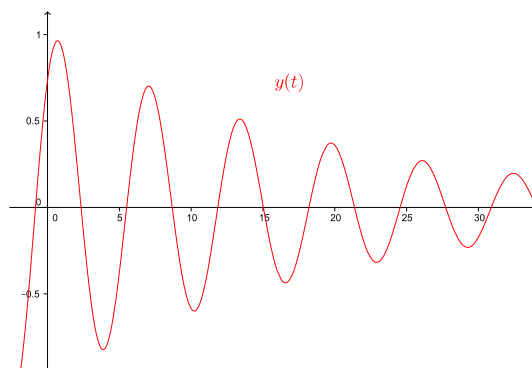
$$y(t) = e^{\alpha t}(A \sin(\omega t) + B \cos(\omega t)) = Ce^{\alpha t} \sin(\omega t + u)$$

mit

$$\alpha = -\frac{\rho}{2m} \quad \text{und} \quad \omega = \sqrt{\frac{k}{m} - \frac{\rho^2}{4m^2}}$$

und reellen Zahlen A, B (bzw. C, u), welche durch die Anfangsbedingungen festgelegt werden. Man erkennt, dass die Schwingung *gedämpft* ist, das heisst, die Amplitude nimmt ab und zwar exponentiell mit $Ce^{\alpha t} = Ce^{-\frac{\rho}{2m}t}$.

Für $k = m = 1$ und $\rho = 0,1$ beispielsweise sieht der Graph von $y(t)$ wie folgt aus:



6.4 Systeme von linearen Differentialgleichungen

Die einfachste Form eines solchen Systems besteht aus zwei Differentialgleichungen

$$\begin{aligned} y_1' &= ay_1 + by_2 \\ y_2' &= cy_1 + dy_2 \end{aligned} \tag{7}$$

mit reellen Zahlen a, b, c, d und hat als Lösung Paare von Funktionen $y_1(x), y_2(x)$. Diese beiden Funktionen sind also untereinander “gekoppelt”.

Systeme von zwei und mehr linearen Differentialgleichungen können mit Hilfe von sogenannten Eigenwerten und Eigenvektoren (der Koeffizientenmatrix) gelöst werden. Diese Begriffe

werden wir jedoch erst im nächsten Semester kennenlernen. Wir werden dann, als Anwendung, ein System von linearen Differentialgleichungen lösen.

Sind die Koeffizienten a, b, c, d des Systems (7) abhängig von x , dann gibt es keine allgemeine Lösungsverfahren.

Räuber-Beute-Beziehungen

Interessant sind speziell Systeme von Differentialgleichungen, die zwei (oder mehrere) Populationen beschreiben, die sich gegenseitig beeinflussen. Betrachten wir zum Beispiel zwei konkurrierende Populationen, Raubtiere und ihre Beutetiere. Abhängig von der Zeit t bezeichne $y_1 = y_1(t)$ die Anzahl der Beutetiere und $y_2 = y_2(t)$ die Anzahl der Räubertiere.

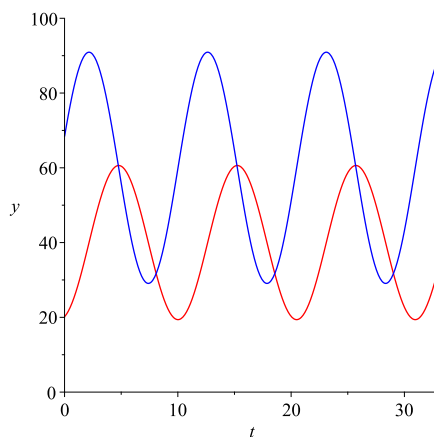
Ohne Räuber würde sich die Beute gemäss $y_1' = ay_1$ vermehren. Sind Räuber anwesend, vermindert sich die Wachstumsrate y_1' um einen Term proportional zur Anzahl der Räuber. Für die Population der Beute ergibt dies eine Differentialgleichung der Form

$$y_1' = (a - by_2) y_1 .$$

Für die Räuber gilt, dass sie ohne Beute gemäss $y_2' = -cy_2$ aussterben würden. Gibt es Beutetiere, dann vermindert sich die Sterberate um einen Term proportional zur Anzahl der Beutetiere. Für die Population der Räuber ergibt dies eine Differentialgleichung der Form

$$y_2' = -(c - dy_1) y_2 .$$

Diese beiden Differentialgleichungen bilden ein System von Differentialgleichungen, welches *Volterra-Lotka Räuber-Beute-Modell* genannt wird. Lösungsfunktionen y_1 und y_2 können nur mit Hilfe von numerischen Verfahren berechnet werden. Eine typische Lösung sieht wie folgt aus, wobei die blaue Kurve die Populationsdichte der **Beute** und die rote Kurve die Populationsdichte der **Räuber** anzeigt.



Zu Beginn wachsen beide Populationen. Ab einem bestimmten Zeitpunkt gibt es zuviele Räuber, so dass der Bestand an Beutetieren stark zurückgeht. Mit der Zeit gibt es dadurch zu wenig Beute für die Räuber, so dass die Räuberpopulation zeitverzögert abnimmt. Dadurch erholt sich jedoch die Beutepopulation wieder. Nach einer gewissen Zeit gibt es wieder genug Beute, so dass auch wieder die Räuberpopulation zunimmt und sich der Zyklus wiederholt.

7 Lineare Gleichungssysteme

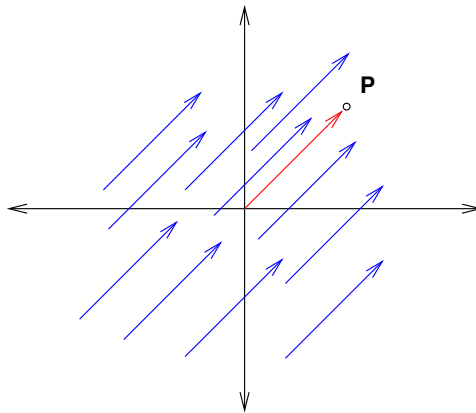
Lineare Gleichungssysteme treten in vielen mathematischen, aber auch naturwissenschaftlichen Problemen auf; zum Beispiel beim Lösen von Differentialgleichungen, bei Optimierungsaufgaben, in der Elektrotechnik und auch in der Chemie. Bei Anwendungen treten meist sehr viele Gleichungen und Unbekannte auf, was effiziente Lösungsmethoden unabdingbar macht. Hilfsmittel dieser Lösungsmethoden sind Vektoren und Matrizen.

7.1 Vektoren in der Ebene und im Raum

In diesem Abschnitt ist das Wichtigste über Vektoren in der Ebene und im Raum zusammengefasst. All dies sollte aus der Schule bekannt sein. Der restliche Stoff dieses Semesters baut auf diesen Grundlagen auf.

In der Ebene und im Raum lassen sich Vektoren geometrisch als gerichtete Strecken oder Pfeile darstellen. Wir beschreiben die Vektoren ausschliesslich durch Länge und Richtung. Deshalb betrachten wir zwei Vektoren als gleich, wenn ihre Richtung und ihre Länge übereinstimmen.

Unter einem *Ortsvektor* $\overrightarrow{OP}$ eines Punktes P verstehen wir den gerichteten Pfeil im Koordinatensystem mit Anfangspunkt im Ursprung O und Endpunkt P . Wir können uns also jeden Vektor als Ortsvektor vorstellen.



Wir schreiben einen Vektor sowohl in der Ebene als auch im Raum als *Spaltenvektor*:

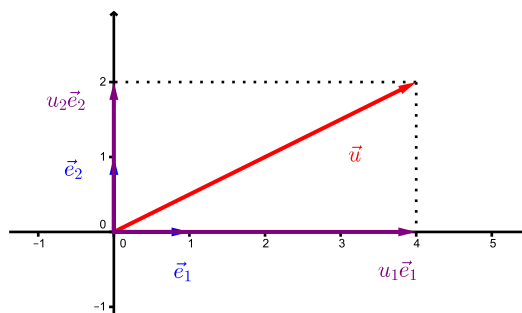
$$\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}, \quad \vec{v} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}.$$

Die reellen Zahlen u_1, u_2 bzw. v_1, v_2, v_3 heissen *Komponenten* des Vektors $\vec{u}$ bzw. $\vec{v}$. Diese Komponenten beziehen sich auf die *Standardbasis* $\vec{e}_1, \vec{e}_2$ der Ebene bzw. $\vec{e}_1, \vec{e}_2, \vec{e}_3$ des Raumes. Das heisst, es gilt

$$\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix} = u_1 \vec{e}_1 + u_2 \vec{e}_2 \quad \text{mit } \vec{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \vec{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

bzw.

$$\vec{v} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = v_1 \vec{e}_1 + v_2 \vec{e}_2 + v_3 \vec{e}_3 \quad \text{mit } \vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}$$

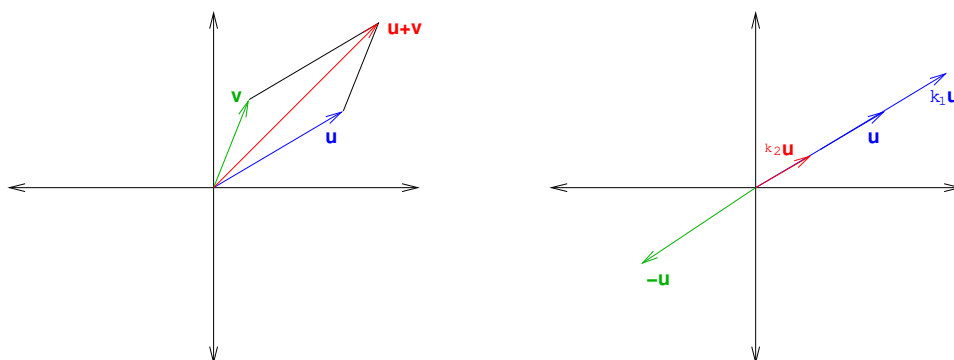


Fassen wir $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}$ als Ortsvektor $\vec{u} = \overrightarrow{OP}$ auf, dann sind die Komponenten u_1, u_2 von $\vec{u}$ gerade die Koordinaten des Punktes P : $P = (u_1, u_2)$. Analog für $\vec{v} = \overrightarrow{OP}$ im Raum.

Rechenregeln

Die Definition von *Summe* $\vec{u} + \vec{v}$, *Differenz* $\vec{u} - \vec{v}$ und *Skalarmultiplikation* $k\vec{u}$ mit einer reellen Zahl k erfolgt komponentenweise:

$$\vec{u} \pm \vec{v} = \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} \pm \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = \begin{pmatrix} u_1 \pm v_1 \\ u_2 \pm v_2 \\ u_3 \pm v_3 \end{pmatrix} \quad \text{und} \quad k\vec{u} = k \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} = \begin{pmatrix} ku_1 \\ ku_2 \\ ku_3 \end{pmatrix}.$$



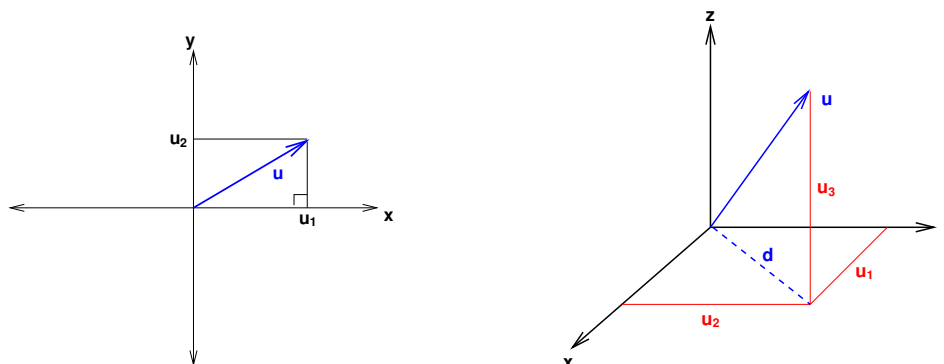
Diese Vektoroperationen gehorchen den folgenden Regeln.

Satz 7.1 Für Vektoren $\vec{u}, \vec{v}$ und $\vec{w}$ in der Ebene oder im Raum und reelle Zahlen k, l gilt:

- (i) $\vec{u} + \vec{v} = \vec{v} + \vec{u}$
- (ii) $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$
- (iii) $\vec{u} + \vec{0} = \vec{0} + \vec{u} = \vec{u}$
- (iv) $\vec{u} + (-\vec{u}) = \vec{0}$
- (v) $k(l\vec{u}) = (kl)\vec{u}$
- (vi) $k(\vec{u} + \vec{v}) = k\vec{u} + k\vec{v}$
- (vii) $(k + l)\vec{u} = k\vec{u} + l\vec{u}$
- (viii) $1\vec{u} = \vec{u}$

Die *Länge* (oder *Norm*) $\|\vec{u}\|$ eines Vektors $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}$ in der Ebene, bzw. $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix}$ im Raum ist gegeben durch

$$\|\vec{u}\| = \sqrt{u_1^2 + u_2^2}, \quad \text{bzw.} \quad \|\vec{u}\| = \sqrt{u_1^2 + u_2^2 + u_3^2}.$$



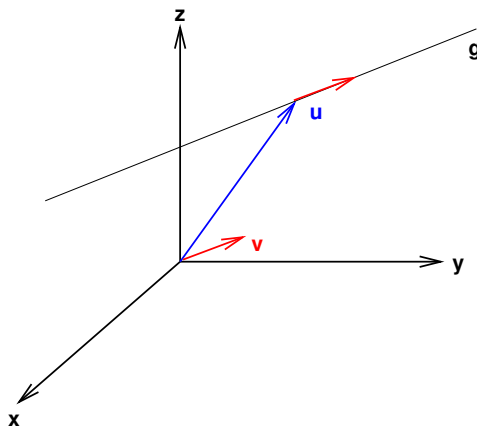
Geraden und Ebenen im Raum

Es gibt im Wesentlichen zwei verschiedene Beschreibungsformen von Geraden und Ebenen, die *Parametergleichung* und die *Koordinatengleichung*.

Parametergleichung einer Geraden. Eine Gerade in der Ebene bzw. im Raum ist gegeben durch

$$\vec{r} = \begin{pmatrix} x \\ y \end{pmatrix} = \vec{u} + t \vec{v} \quad \text{bzw.} \quad \vec{r} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \vec{u} + t \vec{v} \quad \text{mit } t \in \mathbb{R}$$

wobei $\vec{u} = \overrightarrow{OP}$ der Ortsvektor eines (fest gewählten) beliebigen Punktes P auf der Geraden und $\vec{v}$ ein Richtungsvektor längs der Geraden ist. Für jeden Punkt der Geraden gibt es also (genau) ein t in $\mathbb{R}$, so dass der Vektor $\vec{r}$ der Ortsvektor dieses Punktes beschreibt. Man nennt t einen *Parameter*.

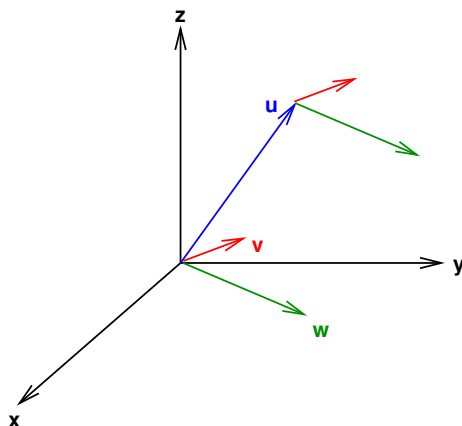


$$\text{Punkt } Q \in g \iff \text{es gibt ein } t \in \mathbb{R} \text{ mit } \vec{r} = \vec{u} + t \vec{v} = \overrightarrow{OQ}$$

Parametergleichung einer Ebene. Eine Ebene im Raum ist gegeben durch

$$\vec{r} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \vec{u} + s\vec{v} + t\vec{w} \quad \text{mit } s, t \in \mathbb{R}$$

wobei $\vec{u} = \overrightarrow{OP}$ der Ortsvektor eines beliebigen Punktes P auf der Ebene und $\vec{v}$ und $\vec{w}$ zwei nicht parallele Richtungsvektoren in der Ebene sind. Hier sind s und t zwei Parameter.



Aus der Schule wissen Sie, dass eine Gleichung der Form $y = mx + q$ eine Gerade in der Ebene beschreibt, wobei m die Steigung und q der y -Achsenabschnitt der Geraden ist. Die folgende leicht allgemeinere Schreibweise dieser Gleichung beschreibt auch die Geraden, die parallel zur y -Achse sind.

Koordinatengleichung einer Geraden in der Ebene. Eine Gerade in der Ebene ist gegeben durch

$$ax + by + c = 0$$

wobei a, b, c reelle Zahlen sind. Dabei gilt:

$$\text{Der Punkt } P = (x_0, y_0) \text{ liegt auf der Geraden} \quad \Longleftrightarrow \quad ax_0 + by_0 + c = 0$$

Koordinatengleichung einer Ebene im Raum. Eine Ebene im Raum ist gegeben durch

$$ax + by + cz + d = 0$$

wobei a, b, c, d reelle Zahlen sind.

Tatsächlich beschreibt diese Gleichung eine Ebene und nicht eine Gerade. Es ist eine Gleichung in den drei Unbekannten x, y, z . Zwei Unbekannte sind also frei wählbar, dann ist die dritte bestimmt. Dies entspricht den zwei Dimensionen der Ebene (bzw. den zwei Parametern der Parametergleichung einer Ebene).

Eine Gerade im Raum ist nicht durch eine einzige Koordinatengleichung beschreibbar. Es sind zwei Koordinatengleichungen dafür nötig. Geometrisch bedeuten die zwei Gleichungen zwei Ebenen. Die Gerade wird also als Schnittgerade zweier Ebenen beschrieben.

Beispiel

Man bestimme die Schnittgerade der beiden Ebenen gegeben durch $2x + 3y - z + 1 = 0$ und $x - y + 2z - 1 = 0$.

7.2 Der n -dimensionale Raum

Im vorhergehenden Abschnitt haben wir speziell die Ebene und den (dreidimensionalen) Raum untersucht. In Räumen mit vier oder mehr Dimensionen kann ganz ähnlich gerechnet werden. Solche Räume kommen ins Spiel, wenn wir lineare Gleichungssysteme lösen wollen. Hat nämlich ein lineares Gleichungssystem vier oder mehr Unbekannte, so liegen seine Lösungen nicht in der Ebene oder im (dreidimensionalen) Raum, sondern in einem Raum von höherer Dimension. Die Rechenarten des vorhergehenden Abschnitts können problemlos auf höherdimensionale Räume erweitert werden.

Sei n eine natürliche Zahl. Die Menge aller geordneten n -Tupel

$$(x_1, x_2, \dots, x_n)$$

mit reellen Zahlen $x_1, x_2, \dots, x_n$ heisst *n -dimensionaler Raum* und wird mit $\mathbb{R}^n$ bezeichnet. Den 2- und 3-dimensionalen Raum kennen wir schon. Der Raum $\mathbb{R}^2$ ist die Ebene und $\mathbb{R}^3$ ist der Raum aus dem vorhergehenden Abschnitt.

Wie in $\mathbb{R}^2$ und $\mathbb{R}^3$ bezeichnen wir mit $P = (x_1, \dots, x_n)$ einen *Punkt* in $\mathbb{R}^n$ und mit

$$\vec{v} = \begin{pmatrix} v_1 \\ \vdots \\ v_n \end{pmatrix}$$

einen *Vektor* in $\mathbb{R}^n$. Eine *Addition* und eine *Skalarmultiplikation* können wir für alle $n \geq 1$

komponentenweise definieren,

$$\vec{u} + \vec{v} = \begin{pmatrix} u_1 \\ \vdots \\ u_n \end{pmatrix} + \begin{pmatrix} v_1 \\ \vdots \\ v_n \end{pmatrix} = \begin{pmatrix} u_1 + v_1 \\ \vdots \\ u_n + v_n \end{pmatrix} \quad \text{und} \quad k\vec{v} = \begin{pmatrix} kv_1 \\ \vdots \\ kv_n \end{pmatrix}$$

für $k \in \mathbb{R}$. Die Subtraktion kann dann als $\vec{u} - \vec{v} = \vec{u} + (-\vec{v})$ definiert werden, wobei $-\vec{v} = (-1)\vec{v}$ der Vektor mit den Komponenten $-v_1, \dots, -v_n$ ist.

Die Vektoren im $\mathbb{R}^n$ gehorchen damit denselben Rechenregeln wie die Vektoren im $\mathbb{R}^2$ und $\mathbb{R}^3$. Es sind die Gesetze (i)–(viii) aus Satz 7.1.

Weiter ist die *Länge* (oder *Norm*) eines Vektors $\vec{u}$ in $\mathbb{R}^n$ definiert durch

$$\|\vec{u}\| = \sqrt{u_1^2 + \dots + u_n^2}.$$

7.3 Lineare Gleichungssysteme und Matrizen

Wir beginnen mit einem Beispiel eines linearen Gleichungssystems aus der Chemie.

Beispiel: Reaktionsgleichung

Wird Kaliumdichromat $\text{K}_2\text{Cr}_2\text{O}_7$ auf über 500°C erhitzt, zerfällt es in Kaliumchromat K_2CrO_4 , Chromoxid Cr_2O_3 und Sauerstoff O_2 . Die Reaktionsgleichung lautet mit unbekannten Moleküllzahlen:



Wie sehen die Koeffizienten x_1, x_2, x_3, x_4 in der Reaktionsgleichung aus, die gewährleisten, dass bei den Reaktanden und Produkten der Gleichung die Anzahlen der jeweiligen Atome gleich sind?

Bringen wir alle Terme auf die linke Seite, erhalten wir die folgenden drei linearen Gleichungen

$$\begin{aligned} 2x_1 - 2x_2 &= 0 \\ 2x_1 - x_2 - 2x_3 &= 0 \\ 7x_1 - 4x_2 - 3x_3 - 2x_4 &= 0 \end{aligned}$$

also ein lineares Gleichungssystem. Wir könnten dieses Gleichungssystem mit Methoden aus der Schule lösen, zum Beispiel durch Einsetzen. Hätte dieses System jedoch mehr Unbekannte und Gleichungen, wäre diese Lösungsmethode sehr aufwendig. Es ist deshalb sinnvoll, eine Lösungsmethode zu kennen, mit der man jedes lineare System effizient lösen kann.

Zuerst halten wir noch ein paar wichtige allgemeine Tatsachen über lineare Gleichungssysteme fest.

Definition Man nennt m lineare Gleichungen

$$\begin{aligned} a_{11}x_1 + \cdots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + \cdots + a_{2n}x_n &= b_2 \\ &\vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n &= b_m \end{aligned} \quad (\text{G})$$

in n Variablen ein *lineares Gleichungssystem*. Die reellen Zahlen $a_{11}, \dots, a_{mn}$ nennt man *Koeffizienten* des Gleichungssystems (oder der Gleichungen). Eine *Lösung* des Systems besteht aus n Zahlen $z_1, \dots, z_n$ mit der Eigenschaft, dass jede Gleichung des Systems durch die Substitution $x_1 = z_1, x_2 = z_2, \dots, x_n = z_n$ erfüllt wird. Die Gesamtheit aller Lösungen heisst *Lösungsmenge* oder *allgemeine Lösung* des Gleichungssystems.

Zum Beispiel bilden die Gleichungen

$$\begin{aligned} x_1 - 3x_2 &= -7 \\ 2x_1 + x_2 &= 7 \end{aligned}$$

ein lineares Gleichungssystem.

Wir verändern die zweite Gleichung und erhalten ein neues Gleichungssystem:

$$\begin{aligned} x_1 - 3x_2 &= -7 \\ 2x_1 - 6x_2 &= 7 \end{aligned}$$

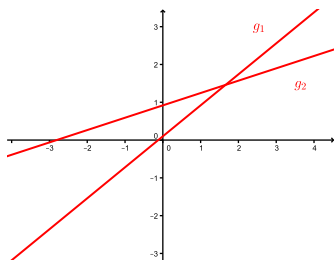
Dieses System hat keine Lösung! Also ist nicht jedes Gleichungssystem lösbar.

Wir verändern die zweite Gleichung noch einmal und erhalten wieder ein neues System:

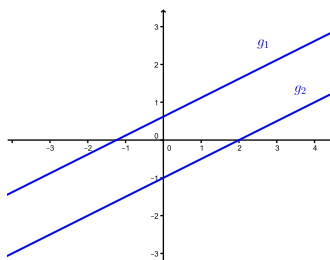
$$\begin{aligned} x_1 - 3x_2 &= -7 \\ 2x_1 - 6x_2 &= -14 \end{aligned}$$

Diese zweite Gleichung ist nun überflüssig, da sie ein Vielfaches der ersten Gleichung ist und deshalb keine neue Bedingung an x_1 und x_2 stellt. Die beiden Gleichungen beschreiben dieselbe Gerade in der Ebene (als Koordinatengleichungen). Alle Punkte (x_1, x_2) auf dieser Geraden sind Lösungen des Gleichungssystems. Dieses System hat also unendlich viele Lösungen.

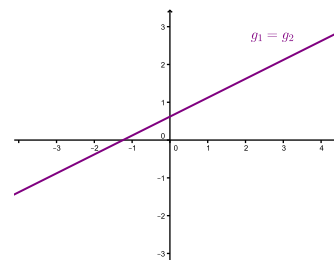
Wir haben nun je ein lineares Gleichungssystem mit genau einer Lösung, mit keiner Lösung und mit unendlich vielen Lösungen gesehen. Tatsächlich können bei einem beliebigen linearen Gleichungssystem stets nur genau diese drei Fälle auftreten.



genau eine Lösung



keine Lösung



unendlich viele Lösungen

Für Systeme in 2 oder 3 Unbekannten kann diese Tatsache geometrisch begründet werden, denn eine Gleichung in 2 Unbekannten beschreibt eine Gerade in der Ebene und eine Gleichung in 3 Unbekannten beschreibt eine Ebene im Raum. Das heisst, das lineare Gleichungssystem hat eine Lösung, wenn die Geraden (bzw. Ebenen) einen gemeinsamen Schnittpunkt haben.

Um lineare Gleichungssysteme effizient lösen zu können, schreibt man sie mit Hilfe von sogenannten Matrizen,

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

Definition Die Matrix A nennt man *Koeffizientenmatrix* des Systems (G) und die Matrix $(A|\vec{b})$ heisst *erweiterte Matrix* des Systems.

Matrizen

Definition Eine (reelle oder komplexe) *Matrix* ist ein rechteckiges Zahlenschema

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

wobei a_{ij} für $i = 1, \dots, m$ und $j = 1, \dots, n$ (reelle oder komplexe) Zahlen sind; man nennt sie *Elemente* oder *Einträge* der Matrix (nächstes Semester werden wir auch Matrizen benutzen, deren Einträge Funktionen sind). Hat die Matrix m *Zeilen* und n *Spalten*, dann bezeichnet man die Matrix als $m \times n$ -*Matrix*.

Ist speziell $n = 1$, so hat die $m \times 1$ -Matrix nur eine Spalte

$$\begin{pmatrix} a_1 \\ \vdots \\ a_m \end{pmatrix},$$

ist also nichts anderes als ein *Spaltenvektor*. Ist $m = 1$, so hat die $1 \times n$ -Matrix nur eine Zeile

$$(a_1 \quad \cdots \quad a_n)$$

und wird auch als *Zeilenvektor* bezeichnet.

Ist $m = n$, so nennt man die $n \times n$ -Matrix *quadratisch der Ordnung n* . Die Elemente $a_{11}, a_{22}, \dots, a_{nn}$ heissen *Diagonalelemente*. Sind alle Elemente ausser den Diagonalelementen einer Matrix gleich Null, dann nennt man die Matrix eine *Diagonalmatrix*.

Eine spezielle Diagonalmatrix ist die *Einheitsmatrix*

$$E = E_n = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}$$

der Ordnung n . Weiter nennt man die Matrix, deren sämtliche Elemente gleich Null sind, *Nullmatrix* und bezeichnet sie mit 0.

Schliesslich nennt man zwei Matrizen *gleich*, wenn sie dieselbe Anzahl Zeilen und Spalten haben (d.h. dieselbe *Grösse* haben) und einander entsprechende Einträge übereinstimmen.

7.4 Der Gauß-Algorithmus

Es gibt verschiedene Methoden, ein lineares Gleichungssystem zu lösen. Der Gauß-Algorithmus (oder das Gauß-Jordan-Verfahren) ist ein Lösungsverfahren, das für beliebige lineare Systeme anwendbar ist. Er ist leicht auf einem Computer programmierbar und vor allem sehr effizient. Weitere Lösungsmethoden sind zum Beispiel die Cramersche Regel und das Lösen mit Hilfe der inversen Matrix, die jedoch nur für gewisse quadratische Koeffizientenmatrizen anwendbar sind, das heisst insbesondere, für lineare Systeme mit gleich vielen Gleichungen wie Unbekannten. Sie sind eher vom theoretischen Standpunkt her interessant.

Gewisse lineare Gleichungssysteme sind ganz einfach zu lösen, zum Beispiel

$$\begin{aligned} x_1 + x_2 - x_3 &= 0 \\ x_2 - x_3 &= 1 \\ x_3 &= 3 \end{aligned} \quad (\text{S1})$$

Durch Rückwärtseinsetzen (d.h. die 3. Gleichung in die 2. Gleichung einsetzen) erhält man

und mit der ersten Gleichung

Noch einfacher zu lösen ist

$$\begin{aligned} x_1 + x_4 &= 4 \\ x_2 - 2x_4 &= 6 \\ x_3 + 3x_4 &= 3 \end{aligned} \quad (\text{S2})$$

Wir setzen $x_4 = t \in \mathbb{R}$ und erhalten damit

Das Ziel des Gauß-Algorithmus ist, ein gegebenes lineares Gleichungssystem in eine der Formen der beiden Beispiele (S1) und (S2) zu bringen. Dadurch kann schliesslich die Lösungsmenge leicht abgelesen werden.

Der Gauß-Algorithmus wird an der erweiterten Matrix des Systems durchgeführt, da dadurch viel Schreibarbeit gespart werden kann. Die erweiterte Matrix des linearen Systems (S1) ist

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

Sie ist in sogenannter Zeilenstufenform. Die erweiterte Matrix des Systems (S2) ist

$$\left(\begin{array}{cccc|c} 1 & 0 & 0 & 1 & 4 \\ 0 & 1 & 0 & -2 & 6 \\ 0 & 0 & 1 & 3 & 3 \end{array} \right)$$

Diese Matrix hat sogenannte reduzierte Zeilenstufenform.

Definition Eine Matrix ist in *Zeilenstufenform*, wenn sie die folgende Gestalt hat:

$$\begin{pmatrix} 1 & * & * & \cdots & * & * \\ & 1 & * & \cdots & * & * \\ & & & & 1 & * \\ & 0 & & & & 0 \end{pmatrix}$$

Dabei gilt:

- Hat die Matrix Zeilen, die nur Nullen enthalten, dann stehen diese in den untersten Zeilen der Matrix.
- Wenn eine Zeile nicht nur aus Nullen besteht, so ist die erste von Null verschiedene Zahl eine Eins (man nennt sie *führende Eins* der Zeile).
- In zwei aufeinanderfolgenden Zeilen, die Einträge $\neq 0$ enthalten, steht die führende Eins der unteren Zeile rechts von der führenden Eins der oberen Zeile.

Zum Beispiel sind die folgenden Matrizen in Zeilenstufenform:

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

Hingegen sind

$$F = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{pmatrix}, \quad G = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 2 & 0 \end{pmatrix}, \quad H = \begin{pmatrix} 1 & 0 & 2 \\ 0 & 1 & 3 \\ 1 & 1 & 5 \end{pmatrix}$$

nicht in Zeilenstufenform.

Definition Eine Matrix ist in *reduzierter Zeilenstufenform*, wenn sie in Zeilenstufenform ist und zusätzlich gilt:

- Eine Spalte, die eine führende Eins enthält, hat keine weiteren Einträge $\neq 0$.

Zum Beispiel sind die folgenden Matrizen in reduzierter Zeilenstufenform:

$$E = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \quad M = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}, \quad N = \begin{pmatrix} 1 & 0 & 2 & 0 \\ 0 & 1 & 2 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Das Ziel des Gauß-Algorithmus ist also, die erweiterte Matrix eines linearen Systems in (reduzierte) Zeilenstufenform zu bringen. Dies kann durch Zeilenumformungen erreicht werden. Zeilenumformungen der erweiterten Matrix eines linearen Systems bedeuten Operationen mit den linearen Gleichungen. Zulässig sind also nur Zeilenumformungen, welche die Lösungsmenge der Gleichungen nicht verändern.

Definition Zeilenumformungen, welche die Lösungsmenge des linearen Gleichungssystems nicht verändern, heissen *elementare Zeilenumformungen*. Es gibt drei verschiedene Typen davon:

1. Vertauschen von zwei Zeilen
2. Multiplikation einer Zeile mit einer Zahl $\neq 0$
3. Addition eines Vielfachen einer Zeile zu einer anderen Zeile

Beispiel

Der **Gauß-Algorithmus** besteht aus den folgenden Schritten:

1. Wir bestimmen die am weitesten links stehende Spalte, die Einträge $\neq 0$ enthält.
2. Ist die oberste Zahl der in Schritt 1 gefundenen Spalte eine Null, dann vertauschen wir die erste Zeile mit einer geeigneten anderen Zeile.
3. Ist a der erste Eintrag der in Schritt 1 gefundenen Spalte, dann dividieren wir die erste Zeile durch a , um die führende Eins zu erzeugen.
- 1.-3. zusammengefasst: Wir erzeugen "oben links" eine (führende) Eins.
4. Wir addieren passende Vielfache der ersten Zeile zu den übrigen Zeilen, um unterhalb der führenden Eins Nullen zu erzeugen.
5. Wir wenden die ersten vier Schritte auf den Teil der Matrix an, den wir durch Streichen der ersten Zeile erhalten, und wiederholen dieses Verfahren, bis die erweiterte Koeffizientenmatrix Zeilenstufenform hat.
6. Mit der letzten nicht verschwindenden Zeile beginnend, addieren wir geeignete Vielfache jeder Zeile zu den darüberliegenden Zeilen, um über den führenden Einsen Nullen zu erzeugen.

Die Schritte 1-5 haben für das lineare System zur Folge, dass sukzessive die Variablen $x_1, x_2, \dots$ eliminiert werden.

Nun ist die erweiterte Matrix in Zeilenstufenform. Man könnte an diesem Punkt das Gleichungssystem durch Rückwärtseinsetzen lösen. Wir tun dies nicht, sondern fahren mit dem Algorithmus (Schritt 6) fort, bis die erweiterte Matrix in reduzierter Zeilenstufenform ist.

$$\begin{array}{rcl} x + 4y & = & -5 \\ y & = & -2 \\ z & = & 2 \end{array}$$

$$\begin{array}{rcl} x & = & 3 \\ y & = & -2 \\ z & = & 2 \end{array}$$

Damit ist das Gleichungssystem gelöst!

2. Wir lösen das Beispiel zu Beginn dieses Abschnitts (zur Reaktionsgleichung) mit Hilfe des Gauß-Algorithmus. Das System war

$$\begin{array}{rcl} 2x_1 - 2x_2 & = & 0 \\ 2x_1 - x_2 - 2x_3 & = & 0 \\ 7x_1 - 4x_2 - 3x_3 - 2x_4 & = & 0. \end{array}$$

Da auf den rechten Seiten der Gleichungen alles Nullen stehen, ist die erweiterte Koeffizientenmatrix von der Form $(A | \vec{0})$. Die letzte Spalte $\vec{0}$ können wir für den Gauß-Algorithmus weglassen, da sie durch die Zeilenumformungen nicht verändert wird.

Wir erhalten also die unendlich vielen Lösungen

$$\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = t \begin{pmatrix} \frac{4}{3} \\ \frac{4}{3} \\ \frac{2}{3} \\ 1 \end{pmatrix}, \quad t \in \mathbb{R}.$$

Zur Vermeidung der Brüche hätte man vorher auch $x_4 = 3t$ setzen können. Damit erhalten wir $x_3 = 2t$, $x_2 = 4t$ und $x_1 = 4t$, das heisst,

$$\vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = t \begin{pmatrix} 4 \\ 4 \\ 2 \\ 3 \end{pmatrix}, \quad t \in \mathbb{R}.$$

Für die Reaktionsgleichung suchen wir eine Lösung mit $x_1, x_2, x_3, x_4 \in \mathbb{N}$. Die kleinste Lösung erhalten wir für $t = 1$,

$$x_1 = 4, \quad x_2 = 4, \quad x_3 = 2, \quad x_4 = 3.$$

Die Reaktionsgleichung lautet damit:



Der Gauß-Algorithmus ist für dieses Beispiel allerdings nicht die effizienteste Lösungsmethode. Die Koeffizientenmatrix A hat “oben rechts” (anstatt “unten links” wie in der Zeilenstufenform) alles Nullen, so dass man direkt $x_1 = t \in \mathbb{R}$ setzen könnte und dann durch sukzessives Einsetzen $x_2 = t$, $x_3 = \frac{1}{2}t$ und $x_4 = \frac{3}{4}t$ erhält.

3. Wir betrachten das System

$$\begin{aligned} x_1 + 3x_2 &= 12 \\ x_1 + x_2 &= 6 \\ 4x_1 - 2x_2 &= 14. \end{aligned}$$

Die dritte Zeile bedeutet (übersetzt in eine lineare Gleichung)

$$0 \cdot x_1 + 0 \cdot x_2 = 8.$$

Das ist ein Widerspruch. Dieses System hat keine Lösung!

7.5 Lösbarkeitskriterien

Gegeben sei eine $m \times n$ -Matrix $A = (a_{ij})$. Wir bringen sie auf Zeilenstufenform

$$A \rightarrow \tilde{A} = \begin{pmatrix} 1 & * & * & \cdots & * & * \\ & 1 & * & \cdots & * & * \\ & & & & 1 & * \\ & 0 & & & & 0 \end{pmatrix}$$

wobei die ersten r Zeilen keine Nullzeilen sind und die letzten $m - r$ Zeilen Nullzeilen sind. Dann nennt man die Zahl r den *Rang* von A ; man schreibt $r = \text{rg}(A)$.

Satz 7.2 *Gegeben ist ein lineares Gleichungssystem mit m Gleichungen und n Unbekannten. Sei A die Koeffizientenmatrix und $(A | \vec{b})$ die erweiterte Matrix.*

(1) *Das System ist lösbar $\iff \text{rg}(A) = \text{rg}(A | \vec{b})$*

(2) *Ist das System lösbar, dann gilt:*

$$\text{Anzahl der freien Parameter} = n - \text{rg}(A)$$

(3) *Das System ist eindeutig lösbar $\iff n = \text{rg}(A) = \text{rg}(A | \vec{b})$*

Beispiele

1. Gegeben sei das lineare Gleichungssystem mit erweiterter Matrix

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

2. Gegeben sei das lineare Gleichungssystem mit erweiterter Matrix

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

Schritte des Gauß-Algorithmus führen auf die erweiterte Matrix

$$(A|\vec{b}) = \left(\begin{array}{ccc|c} -1 & 3 & 2 & 1 \\ 0 & 5 & -1 & 1 \\ 0 & 0 & 0 & 0 \end{array} \right).$$

3. Gegeben sei das lineare Gleichungssystem mit erweiterter Matrix

$$(A|\vec{b}) = \left(\begin{array}{ccc|c} -1 & 3 & 2 & 1 \\ 0 & 5 & -1 & 1 \\ 0 & 0 & 1 & 2 \end{array} \right).$$

Nun bleibt nur noch die Frage offen, ob der Rang einer Matrix unabhängig ist vom Vorgehen, die Matrix auf Zeilenstufenform zu bringen. Tatsächlich ist das der Fall. Das werden wir aber erst im übernächsten Kapitel einsehen können.

7.6 Struktur der Lösungsmenge

Ein lineares Gleichungssystem

$$\begin{aligned} a_{11}x_1 + \cdots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + \cdots + a_{2n}x_n &= b_2 \\ &\vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n &= b_m \end{aligned} \tag{G}$$

heißt *homogen*, falls $b_1 = b_2 = \cdots = b_m = 0$. Andernfalls heißt es *inhomogen*.

Ein homogenes lineares System hat stets die *triviale Lösung* $(x_1, \dots, x_n) = (0, \dots, 0)$.

Satz 7.3 *Ein homogenes lineares Gleichungssystem mit $n > m$ (d.h. mehr Unbekannten als Gleichungen) hat stets unendlich viele Lösungen.*

Dieser Satz ist eine direkte Folgerung von Satz 7.2. Wie oben bemerkt, ist ein homogenes lineares System immer lösbar. Ist A die Koeffizientenmatrix des Systems, dann gilt

Nach Satz 7.2(2) gibt es also mindestens einen freien Parameter, das heißt, unendlich viele Lösungen.

Ist (G) ein lineares Gleichungssystem, so heisst dasjenige Gleichungssystem, welches durch Nullsetzen aller b_i entsteht, das *zugehörige homogene Gleichungssystem* (hG). Wir bezeichnen mit $L(G)$, bzw. $L(hG)$, die Lösungsmenge des Systems (G), bzw. (hG).

Satz 7.4 Wenn das System (G) lösbar ist, gilt

$$L(G) = \vec{x}_P + L(hG)$$

für eine partikuläre Lösung $\vec{x}_P$ von (G).

Dieser Satz kann auf dieselbe Art bewiesen werden wie der Satz 6.1 (Seite 102) über lineare Differentialgleichungen erster Ordnung.

Beispiel

Wir betrachten nochmals das lineare System (S2)

$$\begin{array}{rcrcrcrcrl} x_1 & & & + & x_4 & = & 4 & \\ & x_2 & & - & 2x_4 & = & 6 & \text{(G)} \\ & & x_3 & + & 3x_4 & = & 3 \end{array}$$

von Seite 120. Eine partikuläre (d.h. einzelne) Lösung können wir sofort erraten, nämlich

$$x_4 = 0, x_1 = 4, x_2 = 6, x_3 = 3.$$

Das zugehörige homogene System ist

$$\begin{array}{rcrcrcrcrl} x_1 & & & + & x_4 & = & 0 & \\ & x_2 & & - & 2x_4 & = & 0 & \text{(hG)} \\ & & x_3 & + & 3x_4 & = & 0 \end{array}$$

Die allgemeine Lösung erkennen wir ebenfalls sofort:

$$x_4 = t \in \mathbb{R}, x_1 = -t, x_2 = 2t, x_3 = -3t.$$

Die Lösungen, geschrieben als Vektoren in $\mathbb{R}^4$, sind also

$$\vec{x}_P = \begin{pmatrix} 4 \\ 6 \\ 3 \\ 0 \end{pmatrix}, \quad L(hG) = \left\{ t \begin{pmatrix} -1 \\ 2 \\ -3 \\ 1 \end{pmatrix} \mid t \in \mathbb{R} \right\}.$$

Wir erhalten nun nach Satz 7.4 die Lösungsmenge des Systems (G)

$$L(G) = \vec{x}_P + L(hG) = \left\{ \begin{pmatrix} 4 \\ 6 \\ 3 \\ 0 \end{pmatrix} + t \begin{pmatrix} -1 \\ 2 \\ -3 \\ 1 \end{pmatrix} \mid t \in \mathbb{R} \right\},$$

die natürlich mit der Lösungsmenge, die wir auf Seite 122 gefunden haben, nämlich

$$x_1 = 4 - t, x_2 = 6 + 2t, x_3 = 3 - 3t, x_4 = t \quad \text{für } t \in \mathbb{R},$$

übereinstimmt.

Der Satz 7.4 und das Beispiel zeigen die interessante Struktur der Lösungsmenge $L(G)$ eines inhomogenen linearen Gleichungssystems. Der Gauß-Algorithmus bleibt aber die erste Wahl der Lösungsmethode eines solchen Systems.

8 Rechnen mit Matrizen

Matrizen kann man mit einer reellen Zahl multiplizieren und (mit gewissen Einschränkungen) addieren und multiplizieren. Für Anwendungen sind diese Rechenoperationen sehr nützlich und daher wichtig.

Bezüglich der Addition kann mit Matrizen so gerechnet werden wie mit reellen Zahlen oder Vektoren. Bezüglich der Multiplikation gibt es jedoch einige Rechenregeln der reellen Zahlen, welche für Matrizen nicht mehr gelten.

8.1 Matrixoperationen und ihre Eigenschaften

Addition und skalare Multiplikation

Matrizen können nur dann addiert oder subtrahiert werden, wenn sie dieselbe Grösse (d.h. gleich viele Zeilen und gleich viele Spalten) haben. Sind A und B zwei Matrizen derselben Grösse, so ist ihre *Summe* $A + B$ diejenige Matrix, die durch Addition der einander entsprechenden Elemente entsteht.

Beispiel

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

$$\Rightarrow A + B = \begin{pmatrix} 1+0 & 2+2 & 3+(-4) \\ -1+(-2) & 4+1 & 2+3 \end{pmatrix} = \begin{pmatrix} 1 & 4 & -1 \\ -3 & 5 & 5 \end{pmatrix}$$

Die *Differenz* $A - B$ erhält man durch Subtraktion der Elemente in B von den entsprechenden Elementen in A .

$$\Rightarrow A - B = \begin{pmatrix} 1-0 & 2-2 & 3-(-4) \\ -1-(-2) & 4-1 & 2-3 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 7 \\ 1 & 3 & -1 \end{pmatrix}$$

Sei nun A eine $m \times n$ -Matrix und λ in $\mathbb{R}$. Dann ist die *skalare Multiplikation* λA (d.h. die Multiplikation der Matrix A mit der reellen Zahl λ) die $m \times n$ -Matrix, die durch Multiplikation jedes Elementes von A mit der Zahl λ entsteht. Für $\lambda = 2$ und A wie oben gilt zum Beispiel

$$2A = \begin{pmatrix} 2 \cdot 1 & 2 \cdot 2 & 2 \cdot 3 \\ 2 \cdot (-1) & 2 \cdot 4 & 2 \cdot 2 \end{pmatrix} = \begin{pmatrix} 2 & 4 & 6 \\ -2 & 8 & 4 \end{pmatrix}.$$

Für $\lambda = -1$ schreibt man $-A$ für das Produkt $(-1)A$. Mit A wie oben gilt

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

Die Matrix $-A$ ist das additiv Inverse der Matrix A , das heisst

$$A + (-A) = (-A) + A = 0,$$

wobei mit 0 auf der rechten Seite die Nullmatrix (mit m Zeilen und n Spalten) gemeint ist.

Die wichtigsten Rechenregeln für diese Matrixoperationen sind diejenigen aus Satz 7.1 aus dem 7. Kapitel, wenn Vektoren $\vec{u}, \vec{v}, \vec{w}$ durch Matrizen A, B, C derselben Grösse ersetzt werden.

Matrixmultiplikation

Die Definition der Matrixmultiplikation ist auf den ersten Blick recht unnatürlich, doch interpretiert man eine Matrix als eine sogenannte lineare Abbildung (was wir im nächsten Semester tun werden), dann ist die folgende Definition sinnvoll.

Seien A eine $m \times k$ -Matrix und B eine $k \times n$ -Matrix,

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1k} \\ a_{21} & a_{22} & \cdots & a_{2k} \\ \vdots & \vdots & & \vdots \\ a_{i1} & a_{i2} & \cdots & a_{ik} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mk} \end{pmatrix}, \quad B = \begin{pmatrix} b_{11} & b_{12} & \cdots & b_{1j} & \cdots & b_{1n} \\ b_{21} & b_{22} & \cdots & b_{2j} & \cdots & b_{2n} \\ \vdots & \vdots & & \vdots & & \vdots \\ b_{k1} & b_{k2} & \cdots & b_{kj} & \cdots & b_{kn} \end{pmatrix}.$$

Dann ist das *Produkt* AB eine $m \times n$ -Matrix mit den Einträgen c_{ij} definiert durch

$$c_{ij} = a_{i1}b_{1j} + a_{i2}b_{2j} + \cdots + a_{ik}b_{kj}$$

für $i = 1, \dots, m$ und $j = 1, \dots, n$. Das Element c_{ij} setzt sich also aus den Elementen der i -ten Zeile von A und den Elementen der j -ten Spalte von B zusammen (man merke sich: “Zeile mal Spalte”).

Nach Definition der Matrixmultiplikation kann das Produkt AB also nur dann gebildet werden, wenn die Anzahl der Spalten von A gleich der Anzahl der Zeilen von B ist. Ist dies nicht der Fall, dann ist das Produkt AB nicht definiert.

Beispiel

Eigenschaften der Matrixmultiplikation

1. Die Matrixmultiplikation ist *nicht kommutativ*.

Das heisst, für Matrizen A, B gilt im Allgemeinen

$$AB \neq BA.$$

Zum Beispiel ist AB definiert, BA hingegen nicht. Oder beide Produkte AB und BA sind definiert, jedoch haben sie verschiedene Grössen. Aber auch wenn AB und BA definiert sind und dieselbe Grösse haben, kann man nicht davon ausgehen, dass die beiden Produkte übereinstimmen!

Beispiel

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

$$AB = \begin{pmatrix} 7 & 10 \\ 3 & 4 \end{pmatrix} \neq BA = \begin{pmatrix} 1 & 4 \\ 3 & 10 \end{pmatrix}$$

2. Die Matrixmultiplikation ist *assoziativ* und *distributiv*.

Das heisst, für Matrizen A, B und C entsprechender Grösse gilt:

$$(1) \quad (AB)C = A(BC)$$

$$(2) \quad (A + B)C = AC + BC \quad \text{und} \quad C(A + B) = CA + CB$$

Beispiel

Seien A, B wie vorher und $C = \begin{pmatrix} 2 \\ 3 \end{pmatrix}$. Dann gilt

$$(AB)C = \begin{pmatrix} 7 & 10 \\ 3 & 4 \end{pmatrix} \begin{pmatrix} 2 \\ 3 \end{pmatrix} = \begin{pmatrix} 44 \\ 18 \end{pmatrix}$$

$$A(BC) = \begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 8 \\ 18 \end{pmatrix} = \begin{pmatrix} 44 \\ 18 \end{pmatrix}$$

3. Die *Einheitsmatrix* E übernimmt die Rolle der “Eins”.

Das heisst, es gilt

$$EA = AE = A$$

für jede beliebige Matrix A . Zur Erinnerung: E ist die quadratische Matrix

$$E = \begin{pmatrix} 1 & & 0 \\ & \ddots & \\ 0 & & 1 \end{pmatrix}.$$

Ist A eine $m \times n$ -Matrix, dann steht E beim Produkt EA für eine $m \times m$ -Matrix, beim Produkt AE steht E für eine $n \times n$ -Matrix.

Beispiel

$$EA = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 \\ -1 & 4 & 2 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ -1 & 4 & 2 \end{pmatrix} = A$$

$$AE = \begin{pmatrix} 1 & 2 & 3 \\ -1 & 4 & 2 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 \\ -1 & 4 & 2 \end{pmatrix} = A$$

4. Es gibt *Nullteiler*.

Das heisst, es gibt Matrizen A, B mit

$$AB = 0,$$

aber $A \neq 0$ und $B \neq 0$.

Beispiel

$$A = \begin{pmatrix} 2 & 4 \\ 1 & 2 \end{pmatrix}, B = \begin{pmatrix} -2 & 4 \\ 1 & -2 \end{pmatrix} \implies AB = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} = 0$$

5. Matrizen darf man im Allgemeinen *nicht kürzen*.

Das heisst, für Matrizen A, B, C gilt im Allgemeinen

$$AB = AC \quad \text{und} \quad A \neq 0 \quad \not\Rightarrow \quad B = C$$

Beispiel

$$A = \begin{pmatrix} 2 & 3 \\ 6 & 9 \end{pmatrix}, B = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}, C = \begin{pmatrix} -2 & 1 \\ 3 & 2 \end{pmatrix}$$
$$\implies AB = AC = \begin{pmatrix} 5 & 8 \\ 15 & 24 \end{pmatrix} \quad \text{aber} \quad B \neq C$$

Für reelle Zahlen a, b, c lautet die Kürzungsregel wie folgt:

$$ab = ac \quad \text{und} \quad a \neq 0 \quad \implies \quad b = c$$

Wir müssen also lediglich $a = 0$ ausschliessen. Dabei ist die 0 genau diejenige Zahl, durch die wir nicht dividieren können.

Bei Matrizen genügt es nicht, nur die Matrix $A = 0$ auszuschliessen (wie uns das Beispiel zeigt). Bedeutet dies etwa, dass wir nicht durch jede Matrix $\neq 0$ dividieren können? Wir werden im nächsten Abschnitt sehen, dass genau das der Fall ist.

Die Transponierte

Sei A eine $m \times n$ -Matrix,

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & \cdots & a_{2n} \\ \vdots & \vdots & & & \vdots \\ a_{m1} & a_{m2} & \cdots & \cdots & a_{mn} \end{pmatrix}.$$

Definition Die *Transponierte* A^T von A ist die $n \times m$ -Matrix

$$A^T = \begin{pmatrix} a_{11} & a_{21} & \cdots & a_{m1} \\ a_{12} & a_{22} & \cdots & a_{m2} \\ \vdots & \vdots & & \vdots \\ \vdots & \vdots & & \vdots \\ a_{1n} & a_{2n} & \cdots & a_{mn} \end{pmatrix},$$

das heisst, die Zeilen werden mit den Spalten vertauscht.

Beispiele

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

$$\Rightarrow A^T = \begin{pmatrix} 1 & 2 \\ 0 & 5 \\ 3 & -1 \end{pmatrix}, \quad B^T = \begin{pmatrix} 1 & 2 \\ 3 & 6 \end{pmatrix}, \quad C^T = (1 \ 0 \ 2)$$

Definition Man nennt eine (quadratische) Matrix A *symmetrisch*, wenn gilt

$$A^T = A.$$

Beispiele

Symmetrisch sind

$$A = \begin{pmatrix} 1 & 3 \\ 3 & 6 \end{pmatrix} \quad \text{und} \quad B = \begin{pmatrix} 1 & 0 & 3 \\ 0 & 5 & -1 \\ 3 & -1 & 2 \end{pmatrix}.$$

Satz 8.1 Für Matrizen entsprechender Grössen gilt:

- (1) $(A^T)^T = A$
- (2) $(A + B)^T = A^T + B^T$
- (3) $(AB)^T = B^T A^T$

8.2 Invertierbare Matrizen

Die Division ist als Umkehroperation der Multiplikation definiert. Das heisst, für reelle Zahlen $a \neq 0$ und b gilt $b = \frac{1}{a}$ genau dann, wenn $a \cdot b = 1$.

Übertragen wir dies von den reellen Zahlen $a \neq 0$, b auf quadratische Matrizen $A \neq 0$, B , dann müssen wir feststellen, dass es nicht zu jeder quadratischen Matrix $A \neq 0$ eine quadratische Matrix B mit $AB = E$ gibt (nach der 3. Eigenschaft, Seite 131, übernimmt ja die Einheitsmatrix E die Rolle der 1).

Beispiel

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

Gesucht sind reelle Zahlen a, b, c, d , so dass die Matrix $B = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ die Gleichung $AB = E$ erfüllt.

Für diese Matrix A gibt es also keine 2×2 -Matrix B mit $AB = E$ (analog auch keine Matrix B mit $BA = E$).

Definition Sei A eine quadratische Matrix. Gibt es eine Matrix A^{-1} mit

$$AA^{-1} = A^{-1}A = E$$

so heisst A *invertierbar* und die Matrix A^{-1} nennt man *Inverse* von A (sie ist eindeutig bestimmt durch A).

Beispiel

Invertierbar ist die Matrix

$$A = \begin{pmatrix} 3 & 5 \\ 1 & 2 \end{pmatrix} \quad \text{mit} \quad A^{-1} = \begin{pmatrix} 2 & -5 \\ -1 & 3 \end{pmatrix}.$$

Für die Inverse einer invertierbaren 2×2 -Matrix gibt es eine einfache Formel.

Satz 8.2 Die Matrix

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

ist invertierbar genau dann, wenn $ad - bc \neq 0$. In diesem Fall gilt

$$A^{-1} = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

Beispiel

Sind die Matrizen

$$A = \begin{pmatrix} 4 & 2 \\ 1 & -2 \end{pmatrix} \quad \text{und} \quad B = \begin{pmatrix} 4 & 2 \\ 2 & 1 \end{pmatrix}$$

invertierbar?

Bevor wir weitere Beispiele von Inversen betrachten, kehren wir nochmals zur Eigenschaft zurück, dass Matrizen im Allgemeinen nicht gekürzt werden dürfen. Es gilt nun nämlich das Folgende.

Kürzungsregel Für Matrizen A, B, C gilt:

$$AB = AC \text{ und } A \text{ invertierbar} \implies B = C$$

Diese Kürzungsregel gilt, da die Gleichung $AB = AC$ von links mit A^{-1} multipliziert werden kann:

Weiter gelten die folgenden Eigenschaften für invertierbare Matrizen.

Satz 8.3 Seien A und B zwei invertierbare Matrizen. Dann gilt:

- (1) $(A^{-1})^{-1} = A$
- (2) $(AB)^{-1} = B^{-1}A^{-1}$
- (3) Auch A^T ist invertierbar und $(A^T)^{-1} = (A^{-1})^T$.

Warum wird bei der Eigenschaft (2) die Reihenfolge der Matrizen vertauscht? Nun, sei C die Inverse von AB . Dann gilt $E = C(AB)$. Multiplizieren wir diese Gleichung von rechts mit B^{-1} , dann erhalten wir wegen $BB^{-1} = E$

$$B^{-1} = EB^{-1} = C(AB)B^{-1} = CA(BB^{-1}) = CAE = CA.$$

Nun multiplizieren wir diese Gleichung von rechts mit A^{-1} und finden

$$B^{-1}A^{-1} = CAA^{-1} = CE = C.$$

Bestimmung der Inversen

Um die Inverse einer Matrix von beliebiger Grösse zu bestimmen, kann der Gaußsche Algorithmus (in leicht veränderter Form) benutzt werden. Wie dies funktioniert, wird am folgenden Beispiel erklärt.

Beispiel

Gegeben ist die 3×3 -Matrix

$$A = \begin{pmatrix} 4 & 0 & 5 \\ 0 & 1 & -6 \\ 3 & 0 & 4 \end{pmatrix}.$$

Gesucht ist

$$A^{-1} = \begin{pmatrix} x_1 & u_1 & v_1 \\ x_2 & u_2 & v_2 \\ x_3 & u_3 & v_3 \end{pmatrix}$$

mit $AA^{-1} = E$ (falls A invertierbar ist). Wenn wir die Matrixmultiplikation AA^{-1} ausführen, erhalten wir aus der Matrixgleichung $AA^{-1} = E$ das folgende lineare Gleichungssystem:

$$\begin{array}{rclcl} 4x_1 & & + & 5x_3 & = & 1 \\ & x_2 & - & 6x_3 & = & 0 \\ 3x_1 & & + & 4x_3 & = & 0 \\ \\ 4u_1 & & + & 5u_3 & = & 0 \\ & u_2 & - & 6u_3 & = & 1 \\ 3u_1 & & + & 4u_3 & = & 0 \\ \\ 4v_1 & & + & 5v_3 & = & 0 \\ & v_2 & - & 6v_3 & = & 0 \\ 3v_1 & & + & 4v_3 & = & 1 \end{array}$$

Dies sind 9 Gleichungen in 9 Unbekannten. Doch fassen wir die drei Gleichungen 1–3, 4–6 und 7–9 je als ein Gleichungssystem auf, dann haben diese drei Gleichungssysteme dieselbe Koeffizientenmatrix und nur die Zahlen auf der rechten Seite sind unterschiedlich. Die Schritte im Gauß-Algorithmus zur Lösung dieser drei Systeme sind deshalb für jedes System dieselben. Also führen wir den Gauß-Algorithmus für diese drei Systeme gleichzeitig aus.

Dazu schreiben wir die Koeffizientenmatrix der drei Systeme hin und fügen die Zahlen der rechten Seiten der Gleichungssysteme als Spalten hinzu:

$$\left(\begin{array}{ccc|ccc} 4 & 0 & 5 & 1 & 0 & 0 \\ 0 & 1 & -6 & 0 & 1 & 0 \\ 3 & 0 & 4 & 0 & 0 & 1 \end{array} \right) = (A|E)$$

Nun führen wir den Gauß-Algorithmus durch, und zwar bis wir die Matrix A auf die reduzierte Zeilenstufenform gebracht haben, welches die Einheitsmatrix E ist, falls A invertierbar ist. Wir starten also mit $(A|E)$, führen elementare Zeilenumformungen durch (eine Zeile besteht hier aus 6 Einträgen), bis wir die Form $(E|B)$ erreichen. Die Matrix B ist dann die gesuchte Inverse $A^{-1} = B$!

$$(A|E) = \left(\begin{array}{ccc|ccc} 4 & 0 & 5 & 1 & 0 & 0 \\ 0 & 1 & -6 & 0 & 1 & 0 \\ 3 & 0 & 4 & 0 & 0 & 1 \end{array} \right)$$

Der erste Schritt vom Gauß-Algorithmus wäre $z_1' = \frac{1}{4}z_1$. Um Brüche zu vermeiden, kann man auch wie folgt vorgehen:

Wir erhalten also die Inverse

$$A^{-1} = \begin{pmatrix} 4 & 0 & -5 \\ -18 & 1 & 24 \\ -3 & 0 & 4 \end{pmatrix}.$$

Bei diesem Verfahren muss man von einer gegebenen Matrix A nicht im Voraus wissen, ob sie invertierbar ist oder nicht. Ist A invertierbar, dann führt das eben beschriebene Vorgehen automatisch zur Inversen. Ist A jedoch nicht invertierbar, dann ist die reduzierte Zeilenstufenform von A nicht die Einheitsmatrix E ; das heisst, es ist nicht möglich, durch Zeilenumformungen zur Matrix E zu gelangen.

Beispiel

Gegeben ist die Matrix

$$A = \begin{pmatrix} 4 & 2 \\ 2 & 1 \end{pmatrix}.$$

Dann gilt

$$(A|E) = \left(\begin{array}{cc|cc} 4 & 2 & 1 & 0 \\ 2 & 1 & 0 & 1 \end{array} \right) \xrightarrow{z'_1 = \frac{1}{4}z_1} \left(\begin{array}{cc|cc} 1 & \frac{1}{2} & \frac{1}{4} & 0 \\ 2 & 1 & 0 & 1 \end{array} \right) \xrightarrow{z'_2 = z_2 - 2z_1} \left(\begin{array}{cc|cc} 1 & \frac{1}{2} & \frac{1}{4} & 0 \\ 0 & 0 & -\frac{1}{2} & 1 \end{array} \right)$$

Die Zeilenstufenform der Matrix A hat eine Nullzeile. Die reduzierte Zeilenstufenform von A kann also nicht E sein. Das heisst, dass A nicht invertierbar ist.

Man kann dies auch mit Hilfe des Ranges ausdrücken.

Satz 8.4 Sei A eine $n \times n$ -Matrix. Dann gilt

$$A \text{ ist invertierbar} \iff \text{rg}(A) = n$$

8.3 Potenzen einer Matrix

Für eine quadratische Matrix A definiert man

$$A^0 = E \quad \text{und} \quad A^n = \underbrace{AA \cdots A}_{n \text{ Faktoren}} \quad \text{für } n \geq 1.$$

Ist A ausserdem invertierbar, so ist

$$A^{-n} = (A^{-1})^n = \underbrace{A^{-1}A^{-1} \cdots A^{-1}}_{n \text{ Faktoren}}.$$

Es gelten die üblichen Potenzgesetze, das heisst für ganze Zahlen r und s gilt

$$A^r A^s = A^{r+s} \quad \text{und} \quad (A^r)^s = A^{rs}.$$

Für eine Diagonalmatrix

$$D = \begin{pmatrix} d_1 & & 0 \\ & \ddots & \\ 0 & & d_n \end{pmatrix}$$

gilt

$$D^r = \begin{pmatrix} d_1^r & & 0 \\ & \ddots & \\ 0 & & d_n^r \end{pmatrix} \quad \text{für } r \text{ in } \mathbb{Z}.$$

Insbesondere ist

$$D^{-1} = \begin{pmatrix} \frac{1}{d_1} & & 0 \\ & \ddots & \\ 0 & & \frac{1}{d_n} \end{pmatrix},$$

falls $d_1 \cdots d_n \neq 0$. Mit Diagonalmatrizen lässt es sich also sehr leicht rechnen.

Anwendung: Markov-Ketten

Die Berechnung einer hohen Potenz einer Matrix ist in vielen Anwendungen notwendig. Matrixpotenzen kommen bei iterativen Prozessen ins Spiel.

Beispiel

Nehmen wir an, wir haben zwei Säcke, X und Y , mit Sand. Sei x_0 , bzw. y_0 die Masse des Sandes in Sack X , bzw. Y zu Beginn. Nun verändern wir die Massen in gewissen (regelmässigen) Zeitabständen, und zwar schütten wir die Hälfte des Sandes von Sack X in den Sack Y und $\frac{2}{5}$ des Sandes von Sack Y in den Sack X .

Sei x_k , bzw. y_k die Masse des Sandes in Sack X , bzw. Y nach k Zeitetappen. Dann gilt

Sei $\vec{v}_k$ der Spaltenvektor in $\mathbb{R}^2$ mit den Einträgen x_k und y_k . Dann können wir schreiben

$$\vec{v}_1 = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} & \frac{2}{5} \\ \frac{1}{2} & \frac{3}{5} \end{pmatrix} \begin{pmatrix} x_0 \\ y_0 \end{pmatrix} = A \vec{v}_0 \quad \text{mit der Matrix } A = \begin{pmatrix} \frac{1}{2} & \frac{2}{5} \\ \frac{1}{2} & \frac{3}{5} \end{pmatrix}.$$

Weiter gilt $\vec{v}_2 = A\vec{v}_1 = A^2\vec{v}_0$ und allgemein $\vec{v}_k = A^k\vec{v}_0$ für $k \geq 1$.

Diesen Umverteilungsprozess können wir auch mit Wahrscheinlichkeiten beschreiben. Nehmen wir an, dass sich zu Beginn ein rotes Sandkorn in Sack X befindet. Nach einer Zeitetappe ist die Wahrscheinlichkeit, dass sich dieses rote Sandkorn in Sack Y befindet, gleich $\frac{1}{2}$ (und ebenfalls gleich $\frac{1}{2}$, dass es in Sack X bleibt). Mit Hilfe der Matrix A von oben können wir dies durch $(\frac{1}{2}, \frac{1}{2})^T = A(1, 0)^T$ berechnen. Allgemeiner, ist x_k , bzw. y_k die Wahrscheinlichkeit, dass sich das rote Sandkorn nach k Zeitetappen in Sack X , bzw. Y befindet, dann gilt $\vec{v}_k = (x_k, y_k)^T = A^k \vec{v}_0$ wie oben. Wenn diese Wahrscheinlichkeiten x_k, y_k nur von x_{k-1} und y_{k-1} abhängen (wie wir das angenommen haben), spricht man von einer Markov-Kette.

Allgemeiner versteht man unter einer *Markov-Kette* die folgende Situation. Gegeben sei ein System mit n verschiedenen Zuständen. Sei p_{ji} die Wahrscheinlichkeit, dass das System vom Zustand i in den Zustand j wechselt und sei $A = (p_{ij})$ die entsprechende $n \times n$ -Matrix. Weiter sei $\vec{v}_k = (x_1^{(k)}, \dots, x_n^{(k)})^T$ der Zustandsvektor nach k Zeitetappen, das heisst, $x_i^{(k)}$ ist die Wahrscheinlichkeit, dass das System nach k Zeitetappen im Zustand i ist. Dann gilt

$$\vec{v}_k = A^k \vec{v}_0.$$

Die Matrix A ist eine sogenannte *stochastische Matrix*. Es gilt $0 \leq p_{ij} \leq 1$ und die Summe aller Einträge in einer Spalte ist gleich 1.

Beispiel

Wir beobachten einen Wolf, der sich abwechselnd in der Nähe von Basel und in der Nähe von Liestal aufhält. Unsere Beobachtung zeigt, dass wenn der Wolf an einem Tag in Basel ist, er am folgenden Tag stets in Liestal herumstreicht. Wenn er in Liestal ist, dann ist er am folgenden Tag mit einer Wahrscheinlichkeit von $\frac{1}{4}$ in Basel. Wenn wir den Wolf heute in Basel sehen, mit welcher Wahrscheinlichkeit ist er drei Tage später in Liestal?

Wir müssen also $\vec{v}_3 = A^3 \vec{v}_0$ für $\vec{v}_0 = (1, 0)^T$ berechnen. Wir sehen sofort, dass $\vec{v}_1 = (0, 1)^T$ (da der Wolf nach einem Tag in Basel stets in Liestal ist) und dass $\vec{v}_2 = (\frac{1}{4}, \frac{3}{4})^T$. Für $\vec{v}_3$ erhalten wir

$$\vec{v}_3 = A \vec{v}_2 = \begin{pmatrix} 0 & \frac{1}{4} \\ 1 & \frac{3}{4} \end{pmatrix} \begin{pmatrix} \frac{1}{4} \\ \frac{3}{4} \end{pmatrix} = \begin{pmatrix} \frac{3}{16} \\ \frac{13}{16} \end{pmatrix} = \begin{pmatrix} 0,1875 \\ 0,8125 \end{pmatrix} = \begin{pmatrix} x_3 \\ y_3 \end{pmatrix}.$$

Der Wolf ist also drei Tage später mit einer Wahrscheinlichkeit von $y_3 = 0,8125$ in Liestal.

Wie gross sind wohl die Wahrscheinlichkeiten für beispielsweise $k = 30$ (nach einem Monat) oder $k = 365$ (nach einem Jahr)? Dazu muss man A^{30} , bzw. A^{365} berechnen, was am schnellsten durch sogenanntes Diagonalisieren der Matrix A geht. Wir werden nächstes Semester lernen, wie man quadratische Matrizen diagonalisiert (falls möglich) und wie man dadurch Matrixpotenzen schnell berechnen kann.

8.4 Determinanten

Die *Determinante* ordnet jeder $n \times n$ -Matrix A eine bestimmte reelle Zahl zu. Man bezeichnet sie mit

$$\det(A).$$

Betrachten wir als erstes den Spezialfall $n = 1$: Eine 1×1 -Matrix A besteht nur aus einem Eintrag a , das heisst, $A = (a)$ und es gilt $\det(A) = a$.

$n = 2$

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} \quad \det(A) = a_{11}a_{22} - a_{12}a_{21}$$

Diesen Ausdruck haben wir schon bei der Inversen von A angetroffen (Satz 8.2). Das heisst, die Matrix A ist invertierbar, genau dann wenn $\det(A) \neq 0$.

$$n = 3$$

$$A = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix}$$

$$\begin{aligned} \det(A) &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} \\ &\quad - a_{13}a_{22}a_{31} - a_{11}a_{23}a_{32} - a_{12}a_{21}a_{33} \end{aligned}$$

Beispiel

$$A = \begin{pmatrix} 1 & -1 & 2 \\ 3 & 1 & 0 \\ -2 & 0 & 2 \end{pmatrix}$$

$$n \geq 2$$

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

Für allgemeines $n \geq 2$ kann $\det(A)$ rekursiv definiert werden.

Definition Sei A_{ij} diejenige $(n-1) \times (n-1)$ -Matrix, die man aus A durch Streichen der i -ten Zeile und j -ten Spalte erhält.

Entwicklung nach der ersten Zeile:

$$\det(A) = a_{11} \det(A_{11}) - a_{12} \det(A_{12}) + \dots + (-1)^{n+1} a_{1n} \det(A_{1n})$$

Die Determinante kann nach einer beliebigen Zeile oder Spalte entwickelt werden.

Entwicklung nach der i -ten Zeile:

$$\det(A) = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det(A_{ij})$$

Entwicklung nach der j -ten Spalte:

$$\det(A) = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det(A_{ij})$$

Die reelle Zahl $(-1)^{i+j} \det(A_{ij})$ heisst *Kofaktor* von a_{ij} . Das Vorzeichen $(-1)^{i+j}$ des Kofaktors kann man sich mit Hilfe des folgenden Schemas merken:

$$\begin{pmatrix} + & - & + & - & \cdots \\ - & + & - & + & \cdots \\ + & - & + & & \\ \vdots & \vdots & & \ddots & \end{pmatrix}$$

Beispiele

1. Entwickeln wir eine 3×3 -Matrix nach der ersten Zeile, so erhalten wir die obige Definition:

$$\begin{aligned} \det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} &= a_{11} \det \begin{pmatrix} a_{22} & a_{23} \\ a_{32} & a_{33} \end{pmatrix} - a_{12} \det \begin{pmatrix} a_{21} & a_{23} \\ a_{31} & a_{33} \end{pmatrix} + a_{13} \det \begin{pmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{pmatrix} \\ &= a_{11}a_{22}a_{33} - a_{11}a_{23}a_{32} - a_{12}a_{21}a_{33} + a_{12}a_{23}a_{31} \\ &\quad + a_{13}a_{21}a_{32} - a_{13}a_{22}a_{31} \end{aligned}$$

2. Sei

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

Da die letzte Spalte zwei Nullen enthält, entwickeln wir nach dieser Spalte und erhalten

3. Sei

$$B = \begin{pmatrix} 5 & 6 & 4 & -7 \\ 0 & -9 & 3 & 4 \\ 0 & 0 & 9 & 1 \\ 0 & 0 & 0 & -5 \end{pmatrix}$$

Entwicklung nach der ersten Spalte ergibt

Satz 8.5 Seien A und B zwei $n \times n$ -Matrizen. Dann gilt

- (a) $\det(AB) = \det(A)\det(B)$
- (b) A invertierbar $\iff \det(A) \neq 0$
- (c) $\det(A^{-1}) = \frac{1}{\det(A)}$, falls A invertierbar ist
- (d) $\det(A^T) = \det(A)$

Beim Berechnen von Determinanten sind weiter die folgenden Regeln nützlich:

1. Vertauscht man in einer Matrix zwei Zeilen (oder zwei Spalten), so ändert die Determinante das Vorzeichen.
2. Sind zwei Zeilen (oder zwei Spalten) einer Matrix gleich, so ist die Determinante gleich 0.
3. Multipliziert man eine Zeile (oder Spalte) einer Matrix mit einer reellen Zahl λ , so multipliziert sich auch die Determinante mit λ .
4. Addiert man in einer Matrix ein Vielfaches einer Zeile (oder Spalte) zu einer anderen, so ändert sich die Determinante nicht.
5. Es gilt

$$\det \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ 0 & a_{22} & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \cdots & 0 & a_{nn} \end{pmatrix} = \det \begin{pmatrix} a_{11} & 0 & \cdots & 0 \\ a_{21} & a_{22} & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ a_{n1} & \cdots & \cdots & a_{nn} \end{pmatrix} = a_{11}a_{22} \cdots a_{nn}.$$

Die Regeln 1, 3 und 4 zeigen uns, wie sich die Determinante einer Matrix bei einer elementaren Zeilenumformung verändert. Zur Berechnung der Determinante können wir also elementare Zeilenumformungen durchführen (wie beim Gauß-Algorithmus), um eine obere oder untere Dreiecksmatrix zu erhalten. Mit der Regel 5 ist dann die Berechnung der Determinante einfach.

Beispiel

$$\det \begin{pmatrix} 1 & 1 & 0 & 3 \\ 2 & 0 & 1 & -1 \\ 0 & 2 & -1 & 4 \\ -1 & -1 & 2 & 1 \end{pmatrix} =$$

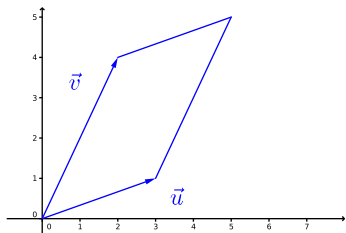
Die Determinante hat eine *geometrische Bedeutung*.

Satz 8.6 Sei A eine $n \times n$ -Matrix. Dann ist $|\det(A)|$ gleich dem Volumen des von den Spaltenvektoren von A aufgespannten Parallelepipeds.

Beispiel

Welchen Flächeninhalt hat das Parallelogramm in $\mathbb{R}^2$ aufgespannt von den Vektoren $\vec{u} = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$

und $\vec{v} = \begin{pmatrix} 2 \\ 4 \end{pmatrix}$?



Schliesslich bleibt noch zu erwähnen, dass die Determinante einer $n \times n$ -Matrix mit Hilfe von Permutationen, das heisst, umkehrbaren Abbildungen $\sigma : \{1, 2, \dots, n\} \rightarrow \{1, 2, \dots, n\}$, definiert werden kann. Die Determinante ist damit eine Summe von $n!$ Summanden (man summiert über alle Permutationen σ). Diese Definition hat aber nur einen theoretischen Nutzen. Deshalb verzichten wir hier auf die genaue Definition.

8.5 Zwei weitere Lösungsmethoden für lineare Gleichungssysteme

Wir betrachten noch einmal ein lineares Gleichungssystem:

$$\begin{aligned} a_{11}x_1 + \cdots + a_{1n}x_n &= b_1 \\ a_{21}x_1 + \cdots + a_{2n}x_n &= b_2 \\ &\vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n &= b_m \end{aligned} \quad (\text{G})$$

Seien

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}, \quad \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \text{ in } \mathbb{R}^n, \quad \vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_m \end{pmatrix} \text{ in } \mathbb{R}^m.$$

Dann können wir das System (G) schreiben als

$$A\vec{x} = \vec{b}.$$

Lösung mit Hilfe der Inversen

Hat das lineare System gleich viele Gleichungen wie Unbekannte, sagen wir n , und ist die zugehörige Koeffizientenmatrix A invertierbar, so gilt $\text{rg}(A) = \text{rg}(A|\vec{b}) = n$. Es gibt also genau eine Lösung. Diese kann mit Hilfe der Inversen von A direkt angegeben werden.

Multiplizieren wir nämlich die Gleichung $A\vec{x} = \vec{b}$ von links mit A^{-1} , dann ist

$$\vec{x} = A^{-1}\vec{b}$$

die eindeutige Lösung des Systems.

Beispiel

Gegeben sei das lineare System

$$\begin{aligned} 3x + y &= 1 \\ 5x + 2y &= 1 \end{aligned}$$

Cramersche Regel

Auch die Cramersche Regel, benannt nach dem Schweizer Mathematiker GABRIEL CRAMER (1704 – 1752), ist nur anwendbar für lineare Systeme mit invertierbarer Koeffizientenmatrix.

Ist $A\vec{x} = \vec{b}$ das lineare System, dann ist die Lösung $\vec{x} = (x_1, \dots, x_n)^T$ gegeben durch

$$x_1 = \frac{\det(A_1)}{\det(A)}, \quad x_2 = \frac{\det(A_2)}{\det(A)}, \quad \dots, \quad x_n = \frac{\det(A_n)}{\det(A)}$$

wobei die Matrix A_j dadurch entsteht, dass die j -te Spalte von A durch den Spaltenvektor $\vec{b}$ ersetzt wird.

Beispiel

Gegeben sei das lineare System

$$\begin{aligned} 3x + y &= 1 \\ 5x + 2y &= 1 \end{aligned}$$

Wie im vorhergehenden Beispiel ist

$$A = \begin{pmatrix} 3 & 1 \\ 5 & 2 \end{pmatrix} \quad \text{mit} \quad \det(A) = 1 \quad \text{und} \quad \vec{b} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

9 Vektorräume

Die Menge $\mathbb{R}^n$ der Vektoren ist nicht die einzige Menge in der Mathematik, deren Elemente man addieren und mit einer reellen Zahl multiplizieren kann. Zur einheitlichen Betrachtung solcher Mengen wurde der Begriff des (abstrakten) Vektorraums eingeführt. Wir werden sehen, dass die Lösungsmenge von jedem homogenen linearen Gleichungssystem ein solcher Vektorraum ist und wir werden dadurch die auftretenden Parameter besser verstehen können. Weiter können wir endlich genau erklären, was Dimension bedeutet.

9.1 Definition und Beispiele

Die Menge $\mathbb{R}^n$ besteht aus (Spalten-)Vektoren $\vec{x}$ mit Komponenten $x_1, \dots, x_n \in \mathbb{R}$. Wir haben in Kapitel 7 gesehen, dass man zwei Vektoren in $\mathbb{R}^n$ addieren und mit einer reellen Zahl multiplizieren kann. Dabei gelten die Rechenregeln von Satz 7.1.

Nun nennt man jede Menge, die genau diese Eigenschaften hat, einen (reellen) Vektorraum (da sich diese Menge eben genau so wie die Menge $\mathbb{R}^n$ der Vektoren verhält). Die Elemente dieser Menge müssen jedoch keine Vektoren sein!

Definition Ein (reeller) Vektorraum ist eine Menge V mit einer Addition und einer Skalarmultiplikation, so dass für alle $\mathbf{u}, \mathbf{v} \in V$, $k \in \mathbb{R}$ auch

$$\mathbf{u} + \mathbf{v} \in V, \quad k\mathbf{v} \in V$$

gilt und alle Eigenschaften aus Satz 7.1 erfüllt sind.

Aus der Bedingung $k \in \mathbb{R}, \mathbf{v} \in V \implies k\mathbf{v} \in V$ folgt insbesondere für $k = 0$, dass

$$\mathbf{0} \in V.$$

Das heisst, ein Vektorraum enthält immer ein *Nullelement*.

Beispiele

1. Die Menge $\mathbb{R}^n$ ist für jedes $n \geq 1$ ein Vektorraum.
2. Die Menge aller $n \times n$ -Matrizen ist ein Vektorraum.

Wir haben in Kapitel 8 gesehen, dass man zwei $n \times n$ -Matrizen addieren und eine Matrix mit einer reellen Zahl multiplizieren kann, wobei die Rechenregeln von Satz 7.1 gelten.

3. Die Menge aller Polynome vom Grad ≤ 2 ,

$$\{ ax^2 + bx + c \mid a, b, c \in \mathbb{R} \},$$

ist ein Vektorraum. Er ist im Wesentlichen derselbe wie $\mathbb{R}^3$.

Analog ist die Menge aller Polynome vom Grad $\leq n$ (oder auch ohne Beschränkung des Grades) ein Vektorraum.

4. Die Menge aller reellen Funktionen $\{ f \mid f : [0, 1] \rightarrow \mathbb{R} \}$ ist ein Vektorraum.

Man definiert Addition und Skalarmultiplikation durch

$$(f + g)(x) = f(x) + g(x) \quad \text{und} \quad (kf)(x) = k f(x) \quad \text{für } x \in [0, 1].$$

Das Nullelement ist dabei die Funktion f mit $f(x) = 0$ für alle $x \in [0, 1]$.

Der Raum $\mathbb{R}^n$ ist also für jede natürliche Zahl n ein Vektorraum. Man kann sich nun fragen, ob es Teilmengen U von $\mathbb{R}^n$ (oder allgemein eines Vektorraums V) gibt, die bezüglich der Addition und Skalarmultiplikation in $\mathbb{R}^n$ (bzw. V) einen Vektorraum bilden.

Glücklicherweise kann man recht schnell überprüfen, ob eine gegebene Teilmenge U eines Vektorraums V selbst ein Vektorraum ist. Die (nichtleere) Menge U ist nämlich genau dann ein Vektorraum in V , wenn gilt

$$\mathbf{u}, \mathbf{v} \in U, k \in \mathbb{R} \implies \mathbf{u} + \mathbf{v} \in U, k\mathbf{u} \in U.$$

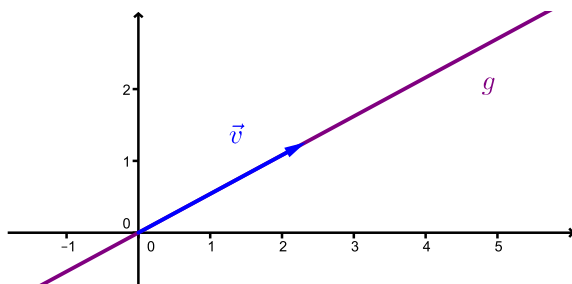
Man muss also nur überprüfen, ob die Teilmenge U *abgeschlossen* ist bezüglich der Addition und der Skalarmultiplikation. Als Teilmenge des Vektorraums V gelten die Eigenschaften von Satz 7.1 automatisch! Wählt man in der Bedingung oben $k = 0$, so sieht man, dass U das Nullelement $\mathbf{0}$ des Vektorraums V enthalten muss.

Jeder Vektorraum $V \neq \{\mathbf{0}\}$ enthält mindestens zwei Teilmengen, die selbst Vektorräume sind, nämlich den ganzen Raum V und den *Nullvektorraum* $\{\mathbf{0}\}$.

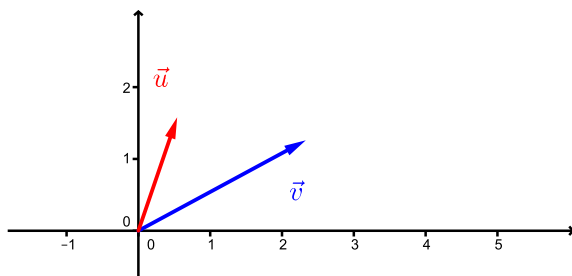
Vektorräume in $\mathbb{R}^2$

Wie eben bemerkt sind $\mathbb{R}^2$ und $\{\vec{0}\}$ Vektorräume.

Ein einzelner Vektor $\vec{v} \neq \vec{0}$ (zusammen mit dem Nullvektor $\vec{0}$) ist kein Vektorraum, da mit $\vec{v}$ auch alle Vielfachen $k\vec{v}$, für $k \in \mathbb{R}$, im Vektorraum liegen müssen. Ist also eine Gerade durch den Ursprung ein Vektorraum?



Auch zwei einzelne (nicht auf einer Geraden liegende) Vektoren $\neq \vec{0}$ (zusammen mit dem Nullvektor $\vec{0}$) bilden keinen Vektorraum. Zunächst müssen die Vielfachen $k\vec{u}$ und $l\vec{v}$, für $k, l \in \mathbb{R}$, im Vektorraum liegen. Aber auch die Summe $\vec{u} + \vec{v}$ und dann $\vec{u} + (\vec{u} + \vec{v})$, usw. muss im Vektorraum liegen. Damit erhält man ganz $\mathbb{R}^2$.



Die Vektorräume in $\mathbb{R}^2$ sind also

- $\{\vec{0}\}$
- Geraden durch den Ursprung
- $\mathbb{R}^2$

Vektorräume in $\mathbb{R}^3$

Die Vektorräume in $\mathbb{R}^3$ sind (gemäss ähnlichen Überlegungen wie in $\mathbb{R}^2$)

- $\{\vec{0}\}$
- Geraden durch den Ursprung
- Ebenen durch den Ursprung
- $\mathbb{R}^3$

Analog zu reellen Vektorräumen kann man *komplexe* Vektorräume definieren. In der Definition des Vektorraums bedeutet die Skalarmultiplikation dann Multiplikation mit einer komplexen Zahl k (anstelle einer reellen Zahl k). Das typische Beispiel eines komplexen Vektorraums ist die Menge $\mathbb{C}^n$ der Vektoren mit Komponenten $z_1, \dots, z_n$ in $\mathbb{C}$. Komplexe Vektorräume treten zum Beispiel in der Quantenmechanik auf oder als Lösungsmengen von homogenen linearen Gleichungssystemen mit komplexen Zahlen als Koeffizienten.

9.2 Linearkombinationen

Im Folgenden beschränken wir uns auf Vektorräume, die Teilmengen von $\mathbb{R}^n$ sind (zur Veranschaulichung kann dabei stets $n = 2$ oder 3 gewählt werden). Jedoch können alle hier gemachten Aussagen von $\mathbb{R}^n$ auf einen beliebigen reellen (oder komplexen) Vektorraum übertragen werden.

Definition Gegeben sind Vektoren $\vec{v}_1, \dots, \vec{v}_r$ in $\mathbb{R}^n$. Ein Vektor $\vec{w}$ in $\mathbb{R}^n$ ist eine *Linearkombination* der Vektoren $\vec{v}_1, \dots, \vec{v}_r$, wenn es reelle Zahlen $c_1, \dots, c_r$ gibt, so dass $\vec{w}$ geschrieben werden kann als

$$\vec{w} = c_1 \vec{v}_1 + \dots + c_r \vec{v}_r.$$

Die Zahlen $c_1, \dots, c_r$ nennt man *Koeffizienten*.

Beispiel

Ist der Vektor $\vec{w} = \begin{pmatrix} 5 \\ 6 \end{pmatrix}$ eine Linearkombination von $\vec{v}_1 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ und $\vec{v}_2 = \begin{pmatrix} 2 \\ 0 \end{pmatrix}$?

Nun sei V die Menge aller Linearkombinationen der Vektoren $\vec{v}_1, \dots, \vec{v}_r$ in $\mathbb{R}^n$. Dann ist V ein Vektorraum in $\mathbb{R}^n$. Man schreibt

$$V = \langle \vec{v}_1, \dots, \vec{v}_r \rangle$$

(oder $V = \text{Lin}(\vec{v}_1, \dots, \vec{v}_r)$ oder $V = \text{span}(\vec{v}_1, \dots, \vec{v}_r)$). Man sagt auch, dass $\vec{v}_1, \dots, \vec{v}_r$ den Vektorraum V *aufspannen* oder *erzeugen*.

Beispiele

Gegeben sind die Vektoren $\vec{v}_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}$, $\vec{v}_3 = \begin{pmatrix} 2 \\ 0 \\ 0 \end{pmatrix}$ und $\vec{v}_4 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$.

1. Sei $V_1 = \langle \vec{v}_1 \rangle$ der von $\vec{v}_1$ aufgespannte Vektorraum in $\mathbb{R}^3$.

Dann besteht V_1 aus allen Vielfachen von $\vec{v}_1$,

$$V_1 = \{ c \vec{v}_1 \mid c \in \mathbb{R} \}.$$

Also ist V_1 die Gerade durch den Ursprung mit Richtungsvektor $\vec{v}_1$.

2. Sei $V_2 = \langle \vec{v}_1, \vec{v}_2 \rangle$ der von den Vektoren $\vec{v}_1$ und $\vec{v}_2$ aufgespannte Vektorraum.

Das heisst,

$$V_2 = \{ c_1 \vec{v}_1 + c_2 \vec{v}_2 \mid c_1, c_2 \in \mathbb{R} \}.$$

Insbesondere sind

$$\vec{e}_1 = \frac{1}{2} \vec{v}_1 + \frac{1}{2} \vec{v}_2 \quad \text{und} \quad \vec{e}_3 = \frac{1}{2} \vec{v}_1 - \frac{1}{2} \vec{v}_2 \quad \text{in } V_2.$$

Der Vektorraum V_2 ist also die xz -Ebene.

3. Sei $V_3 = \langle \vec{v}_1, \vec{v}_2, \vec{v}_3 \rangle$.

Der Vektor $\vec{v}_3$ ist eine Linearkombination der Vektoren $\vec{v}_1$ und $\vec{v}_2$,

$$\vec{v}_3 = \vec{v}_1 + \vec{v}_2$$

Also ist $\vec{v}_3 \in V_2$ und damit ist $V_3 = V_2 = \langle \vec{v}_1, \vec{v}_2 \rangle$.

4. Sei $V_4 = \langle \vec{v}_1, \vec{v}_2, \vec{v}_4 \rangle$.

Nun ist $\vec{v}_4$ keine Linearkombination der Vektoren $\vec{v}_1$ und $\vec{v}_2$. Da $\vec{v}_4 = \vec{e}_2$, sehen wir sofort, dass $V_4 = \mathbb{R}^3$.

Man möchte nun einen Vektorraum $V = \langle \vec{v}_1, \dots, \vec{v}_r \rangle$ in $\mathbb{R}^n$ mit möglichst wenigen Vektoren erzeugen. Wir haben im 3. Beispiel gesehen, dass man einen Vektor $\vec{v}_i$ weglassen kann, wenn man ihn als Linearkombination der anderen Vektoren schreiben kann.

Definition Sei $r \geq 2$. Die Vektoren $\vec{v}_1, \dots, \vec{v}_r$ in $\mathbb{R}^n$ heissen *linear abhängig*, wenn sich einer der r Vektoren als Linearkombination der anderen $r - 1$ Vektoren schreiben lässt. Ist dies nicht möglich, dann nennt man die Vektoren $\vec{v}_1, \dots, \vec{v}_r$ *linear unabhängig*.

Beispiel

Sind die Vektoren $\vec{v}_1 = \begin{pmatrix} 2 \\ -1 \\ 0 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 1 \\ 2 \\ 5 \end{pmatrix}$, $\vec{v}_3 = \begin{pmatrix} 7 \\ -1 \\ 5 \end{pmatrix}$ linear abhängig?

Wie überprüft man aber nun, dass gewisse Vektoren linear unabhängig sind? Dazu ist oft die folgende Bemerkung hilfreich.

Bemerkung Die Vektoren $\vec{v}_1, \dots, \vec{v}_r$ in $\mathbb{R}^n$ sind linear unabhängig genau dann, wenn die Gleichung

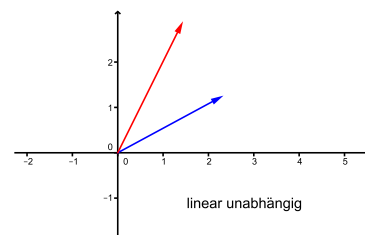
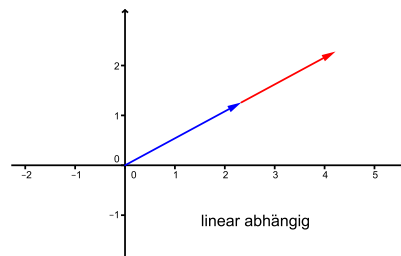
$$c_1 \vec{v}_1 + \dots + c_r \vec{v}_r = \vec{0}$$

nur die (triviale) Lösung $c_1 = c_2 = \dots = c_r = 0$ hat.

Beschränken wir uns jedoch auf Vektoren in $\mathbb{R}^2$ und $\mathbb{R}^3$, dann kann die Frage nach der linearen (Un-)Abhängigkeit mit Hilfe von geometrischen Betrachtungen beantwortet werden.

Vektoren in $\mathbb{R}^2$

- 2 Vektoren sind linear abhängig
 - $\iff$ einer ist ein Vielfaches des anderen
 - $\iff$ sie liegen auf der gleichen Geraden durch den Ursprung



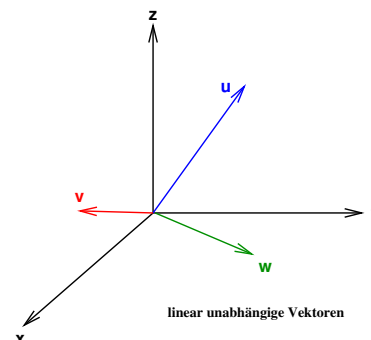
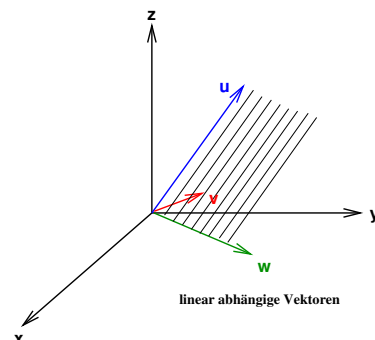
- 3 oder mehr Vektoren sind stets linear abhängig.
Machen wir nämlich den Ansatz

$$c_1 \vec{v}_1 + c_2 \vec{v}_2 + c_3 \vec{v}_3 = \vec{0}$$

für die drei Vektoren $\vec{v}_1, \vec{v}_2, \vec{v}_3$ in $\mathbb{R}^2$, dann ist dies ein homogenes lineares Gleichungssystem mit 2 Gleichungen in den 3 Unbekannten c_1, c_2, c_3 . Nach Satz 7.3 hat dieses unendlich viele Lösungen, insbesondere eine nichttriviale Lösung.

Vektoren in $\mathbb{R}^3$

- 2 Vektoren sind linear abhängig
 - $\iff$ einer ist ein Vielfaches des anderen
 - $\iff$ sie liegen auf der gleichen Geraden durch den Ursprung
- 3 Vektoren sind linear abhängig
 - $\iff$ sie liegen auf der gleichen Ebene durch den Ursprung



- 4 und mehr Vektoren sind stets linear abhängig

Die letzte Aussage bei Vektoren in $\mathbb{R}^2$ und $\mathbb{R}^3$ gilt allgemeiner.

Satz 9.1 Seien $\vec{v}_1, \dots, \vec{v}_r$ Vektoren in $\mathbb{R}^n$. Ist $r > n$, dann sind $\vec{v}_1, \dots, \vec{v}_r$ linear abhängig.

Begründen kann man dies wie in $\mathbb{R}^2$ mit Hilfe von Satz 7.3 über homogene lineare Gleichungssysteme (man hat n Gleichungen und $r > n$ Unbekannte).

Wie können wir konkret überprüfen, ob drei Vektoren in $\mathbb{R}^3$ linear unabhängig sind oder nicht? Nun, drei Vektoren $\vec{v}_1, \vec{v}_2, \vec{v}_3$ in $\mathbb{R}^3$ spannen genau dann ein Parallelepiped auf, wenn sie linear unabhängig sind (d.h. nicht in einer Ebene liegen). In diesem Fall (und nur in diesem) ist das Volumen dieses Parallelepipeds eine Zahl ungleich Null. Nach Satz 8.6 ist dieses Volumen gleich $|\det(A)|$, wobei A die 3×3 -Matrix mit den Spalten $\vec{v}_1, \vec{v}_2, \vec{v}_3$ ist. Es gilt also:

$$\vec{v}_1, \vec{v}_2, \vec{v}_3 \text{ linear unabhängig} \iff \det(A) \neq 0$$

Auch dies gilt allgemeiner.

Satz 9.2 Seien $\vec{v}_1, \dots, \vec{v}_n$ Vektoren in $\mathbb{R}^n$ und A die $n \times n$ -Matrix mit $\vec{v}_1, \dots, \vec{v}_n$ als Spalten. Dann gilt:

$$\vec{v}_1, \dots, \vec{v}_n \text{ linear unabhängig} \iff \det(A) \neq 0$$

Beispiele

1. Sind die Vektoren $\vec{v}_1 = \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 2 \\ 0 \\ 3 \end{pmatrix}$ und $\vec{v}_3 = \begin{pmatrix} 0 \\ -1 \\ 2 \end{pmatrix}$ linear unabhängig?

2. Auf Seite 148 haben wir gesehen, dass $\vec{v}_1 = \begin{pmatrix} 2 \\ -1 \\ 0 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 1 \\ 2 \\ 5 \end{pmatrix}$, $\vec{v}_3 = \begin{pmatrix} 7 \\ -1 \\ 5 \end{pmatrix}$ linear abhängig sind. Tatsächlich gilt $\det(\vec{v}_1 \vec{v}_2 \vec{v}_3) = 0$:

9.3 Basis und Dimension

Im letzten Abschnitt haben wir versucht, einen Vektorraum mit so wenigen Vektoren wie möglich zu erzeugen. Wird ein Vektorraum $V = \langle \vec{v}_1, \dots, \vec{v}_r \rangle$ in $\mathbb{R}^n$ mit linear abhängigen Vektoren $\vec{v}_1, \dots, \vec{v}_r$ erzeugt, so kann mindestens ein Vektor (einer, der sich als Linearkombination der anderen schreiben lässt) weggelassen werden, ohne dass sich der Vektorraum V verkleinert. Optimal ist also, einen Vektorraum mit linear unabhängigen Vektoren zu erzeugen. Man nennt diese Vektoren eine *Basis* des Vektorraums.

Definition Vektoren $\vec{v}_1, \dots, \vec{v}_n$ bilden eine *Basis* eines Vektorraums V , wenn sie die folgenden zwei Bedingungen erfüllen:

- (1) $V = \langle \vec{v}_1, \dots, \vec{v}_n \rangle$, das heisst, $\vec{v}_1, \dots, \vec{v}_n$ erzeugen V .
- (2) $\vec{v}_1, \dots, \vec{v}_n$ sind linear unabhängig.

Man nennt $\vec{v}_1, \dots, \vec{v}_n$ *Basisvektoren* von V .

Basen sind wegen der folgenden Tatsache sehr wichtig und nützlich.

Satz 9.3 *Bilden $\vec{v}_1, \dots, \vec{v}_n$ eine Basis des Vektorraums V , so kann jeder Vektor $\vec{v}$ in V eindeutig geschrieben werden als*

$$\vec{v} = c_1 \vec{v}_1 + \dots + c_n \vec{v}_n$$

mit reellen Zahlen $c_1, \dots, c_n$.

Basen von $\mathbb{R}^2$

In Kapitel 7 haben wir bemerkt, dass die Schreibweise $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ bedeutet

$$\vec{v} = v_1 \begin{pmatrix} 1 \\ 0 \end{pmatrix} + v_2 \begin{pmatrix} 0 \\ 1 \end{pmatrix} = v_1 \vec{e}_1 + v_2 \vec{e}_2.$$

Die beiden Vektoren $\vec{e}_1$ und $\vec{e}_2$ bilden eine Basis von $\mathbb{R}^2$, die sogenannte *Standardbasis*.

Die Vektoren $\vec{e}_1$ und $\vec{e}_2$ spannen den ganzen Raum $\mathbb{R}^2$ auf, denn wie gerade gezeigt ist jedes $\vec{v}$ in $\mathbb{R}^2$ als Linearkombination von $\vec{e}_1$ und $\vec{e}_2$ schreibbar. Zudem sind $\vec{e}_1$ und $\vec{e}_2$ linear unabhängig.

Ebensogut könnte man jedoch die Vektoren

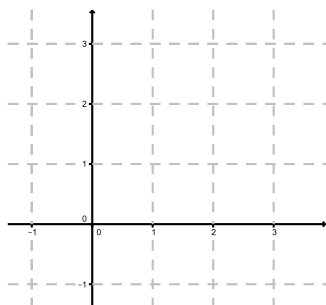
$$\vec{u}_1 = \begin{pmatrix} 3 \\ 1 \end{pmatrix} \quad \text{und} \quad \vec{u}_2 = \begin{pmatrix} 1 \\ -1 \end{pmatrix}$$

als Basis des $\mathbb{R}^2$ wählen. Denn jeder Vektor $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ von $\mathbb{R}^2$ lässt sich als Linearkombination von $\vec{u}_1$ und $\vec{u}_2$ schreiben:

$$\vec{v} = \frac{1}{4}(v_1 + v_2) \vec{u}_1 + \frac{1}{4}(v_1 - 3v_2) \vec{u}_2$$

Das heisst, die Vektoren $\vec{u}_1$ und $\vec{u}_2$ erzeugen $\mathbb{R}^2$. Zudem sind $\vec{u}_1$ und $\vec{u}_2$ linear unabhängig.

Die Wahl einer Basis von $\mathbb{R}^2$ (bzw. $\mathbb{R}^n$) ist nichts anderes als die Wahl eines Koordinatensystems für $\mathbb{R}^2$ (bzw. $\mathbb{R}^n$). Die Richtungen der Basisvektoren definieren die positiven Koordinatenachsen und ihre Längen legen die Masseneinheiten fest.



Können drei Vektoren in $\mathbb{R}^2$ eine Basis für $\mathbb{R}^2$ bilden? Nein, denn drei Vektoren in $\mathbb{R}^2$ sind stets linear abhängig (Satz 9.1). Kann ein einziger Vektor in $\mathbb{R}^2$ eine Basis für $\mathbb{R}^2$ sein? Nein, denn ein einzelner Vektor spannt nur eine Gerade auf, der $\mathbb{R}^2$ ist jedoch eine Ebene. Es folgt, dass eine Basis für $\mathbb{R}^2$ immer aus zwei Vektoren besteht.

Satz 9.4 *Hat ein Vektorraum endlich viele Basisvektoren, so ist die Anzahl der Basisvektoren für alle Basen gleich.*

Definition Die Anzahl der Basisvektoren eines Vektorraums V heisst *Dimension* von V . Man schreibt $\dim(V)$.

Für $V = \{\vec{0}\}$ setzt man $\dim(\{\vec{0}\}) = 0$.

Eine Gerade durch den Ursprung ist ein Vektorraum erzeugt durch *einen* Vektor $\neq \vec{0}$, also ist seine Dimension gleich 1.

Der Vektorraum $\mathbb{R}^2$ hat (wie oben bemerkt) stets zwei Basisvektoren, also ist seine Dimension gleich 2. Tatsächlich bilden *zwei* Vektoren in $\mathbb{R}^2$ eine Basis genau dann, wenn sie linear unabhängig sind.

Satz 9.5 *Gegeben seien n Vektoren $\vec{v}_1, \dots, \vec{v}_n$ eines Vektorraums V , wobei $n = \dim(V)$. Dann gilt:*

$$\vec{v}_1, \dots, \vec{v}_n \text{ bilden eine Basis von } V \iff \vec{v}_1, \dots, \vec{v}_n \text{ sind linear unabhängig}$$

Beispiel

Bilden die Vektoren $\vec{v}_1 = \begin{pmatrix} 2 \\ 4 \end{pmatrix}$ und $\vec{v}_2 = \begin{pmatrix} 3 \\ 5 \end{pmatrix}$ eine Basis von $\mathbb{R}^2$?

Basen von $\mathbb{R}^3$

Schon in Kapitel 7 haben wir festgehalten, dass die Schreibweise $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}$ bedeutet

$$\vec{v} = v_1 \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} + v_2 \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} + v_3 \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = v_1 \vec{e}_1 + v_2 \vec{e}_2 + v_3 \vec{e}_3.$$

Die Vektoren $\vec{e}_1, \vec{e}_2$ und $\vec{e}_3$ bilden die sogenannte *Standardbasis* von $\mathbb{R}^3$. Die Dimension von $\mathbb{R}^3$ ist also 3.

Wie für $\mathbb{R}^2$ können wir auch andere Basisvektoren für $\mathbb{R}^3$ wählen. Wegen Satz 9.5 bilden beliebige drei linear unabhängige Vektoren von $\mathbb{R}^3$ eine Basis von $\mathbb{R}^3$; zum Beispiel die drei Vektoren $\vec{v}_1, \vec{v}_2, \vec{v}_3$ vom 1. Beispiel auf Seite 150.

Beispiel

Welche Dimension hat der Vektorraum $V = \langle \vec{v}_1, \vec{v}_2, \vec{v}_3 \rangle$ von $\mathbb{R}^3$ mit $\vec{v}_1 = \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 1 \\ -3 \end{pmatrix}$ und

$$\vec{v}_3 = \begin{pmatrix} -5 \\ 8 \\ -3 \end{pmatrix} ?$$

Die Vektoren $\vec{v}_1, \vec{v}_2, \vec{v}_3$ sind also linear abhängig und bilden keine Basis von V . Die Dimension von V ist demnach kleiner als 3. Ist die Dimension 2, das heisst, gibt es zwei linear unabhängige Vektoren?

Standardbasis von $\mathbb{R}^n$

Analog zu $\mathbb{R}^2$ und $\mathbb{R}^3$ hat $\mathbb{R}^n$ für beliebige n in $\mathbb{N}$ die *Standardbasis*

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

Es gilt also $\dim(\mathbb{R}^n) = n$.

Die Vektorräume $\mathbb{R}^n$ und alle darin enthaltenen Vektorräume sind alles Beispiele von *endlich-dimensionalen* Vektorräumen, das heisst von Vektorräumen mit endlich vielen Basisvektoren. Hat eine Basis eines Vektorraums V unendlich viele Vektoren, so nennt man V *unendlich-dimensional*. Zum Beispiel ist der Vektorraum aller reellen Funktionen unendlich-dimensional.

1. Anwendung: Der Rang einer Matrix

Mit Hilfe der neuen Kenntnisse über Vektorräume können wir den Rang einer Matrix ohne Zuhilfenahme einer Zeilenstufenform definieren.

Definition Sei A eine $m \times n$ -Matrix. Der *Zeilenraum* von A ist der von den m Zeilen von A aufgespannte Vektorraum in $\mathbb{R}^n$. Der *Spaltenraum* von A ist der von den n Spalten von A aufgespannte Vektorraum in $\mathbb{R}^m$.

Die Zeilen von A fassen wir also als Vektoren in $\mathbb{R}^n$ auf, die Spalten als Vektoren von $\mathbb{R}^m$.

Nun gilt der folgende Zusammenhang mit dem Rang $\text{rg}(A)$ der Matrix A :

$$\begin{aligned} \text{rg}(A) &= \text{maximale Anzahl linear unabhängiger Zeilenvektoren} \\ &= \text{Dimension des Zeilenraums} \end{aligned}$$

Bringen wir nämlich die Matrix A auf Zeilenstufenform, dann ersetzen wir bei einer elementaren Zeilenumformung eine Zeile durch eine Linearkombination dieser Zeile mit einer anderen; das heisst, der Zeilenraum wird dabei nicht verändert. Die Nichtnullzeilen in der Zeilenstufenform sind schliesslich linear unabhängig (wegen den Nullen “unten links”).

Beispiel

Sei

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

Erstaunlicherweise spielt es überhaupt keine Rolle, ob wir die Zeilen oder die Spalten der Matrix A betrachten, um den Rang zu bestimmen.

Satz 9.6 $\operatorname{rg}(A) = \operatorname{Dimension}(\text{Zeilenraum}) = \operatorname{Dimension}(\text{Spaltenraum})$

Der Rang einer Matrix ist also tatsächlich unabhängig vom Vorgehen, die Matrix auf Zeilenstufenform zu bringen. Dies war in Kapitel 7 (Seite 127) noch unklar.

An dieser Stelle fassen wir einmal zusammen, was der Rang einer $n \times n$ -Matrix alles aussagt.

Satz 9.7 *Sei A eine $n \times n$ -Matrix. Dann gilt:*

$$\begin{aligned} \operatorname{rg}(A) = n &\iff A \text{ ist invertierbar} \\ &\iff \det(A) \neq 0 \\ &\iff \text{die Spaltenvektoren sind linear unabhängig} \\ &\iff \text{die Spaltenvektoren sind eine Basis von } \mathbb{R}^n \\ &\iff \text{die Zeilenvektoren sind eine Basis von } \mathbb{R}^n \\ &\iff \text{die Zeilenvektoren sind linear unabhängig} \end{aligned}$$

Weiter sind alle Aussagen von Satz 9.7 gleichbedeutend mit der Aussage, dass das lineare Gleichungssystem $A\vec{x} = \vec{b}$ für jedes $\vec{b} \in \mathbb{R}^n$ genau eine Lösung hat.

2. Anwendung: Parameter der Lösungsmenge eines linearen Gleichungssystems

Sei A eine $m \times n$ -Matrix. Wir betrachten das homogene lineare Gleichungssystem

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

für $\vec{x}$ in $\mathbb{R}^n$.

Seien $\vec{x}_1$ und $\vec{x}_2$ zwei Lösungen dieses Systems. Dann ist auch die Summe dieser beiden Lösungen eine Lösung, denn

$$A(\vec{x}_1 + \vec{x}_2) = A\vec{x}_1 + A\vec{x}_2 = \vec{0} + \vec{0} = \vec{0}.$$

Weiter sind die Vielfachen dieser Lösungen auch Lösungen, denn für $t \in \mathbb{R}$ gilt

$$A(t\vec{x}_1) = tA\vec{x}_1 = t \cdot \vec{0} = \vec{0}.$$

Dies bedeutet: Der Lösungsraum ist ein Vektorraum in $\mathbb{R}^n$!

Satz 9.8 Sei A eine $m \times n$ -Matrix. Dann ist die Lösungsmenge von $A\vec{x} = \vec{0}$ ein Vektorraum in $\mathbb{R}^n$ der Dimension $k = n - \text{rg}(A)$.

Das heisst, es gibt k linear unabhängige Vektoren $\vec{x}_1, \dots, \vec{x}_k$, so dass jede Lösung $\vec{x}$ eindeutig geschrieben werden kann als

$$\vec{x} = t_1\vec{x}_1 + \dots + t_k\vec{x}_k$$

für $t_1, \dots, t_k$ in $\mathbb{R}$.

Die reellen Zahlen $t_1, \dots, t_k$ sind die *Parameter* der Lösung. Die Vektoren $\vec{x}_1, \dots, \vec{x}_k$ sind eine Basis des Lösungsraums.

Geometrisch gesehen kommen als Lösungsräume eines homogenen linearen Gleichungssystems mit 3 Unbekannten also nur der Nullvektorraum $\{\vec{0}\}$, eine Gerade durch den Ursprung, eine Ebene durch den Ursprung oder der ganze Raum $\mathbb{R}^3$ in Frage.

Beispiel

Wir betrachten das lineare Gleichungssystem

$$\begin{aligned} x + 2y &= 0 \\ y - z &= 0 \end{aligned}$$

Dies ist ein homogenes System $A\vec{x} = \vec{0}$ mit der Koeffizientenmatrix

$$A = \begin{pmatrix} 1 & 2 & 0 \\ 0 & 1 & -1 \end{pmatrix},$$

die wir schon im Beispiel auf Seite 154 untersucht haben. Nach Satz 9.8 ist die Lösungsmenge ein Vektorraum der Dimension

Es ist also eine Gerade durch den Ursprung. Sie ist gegeben durch

Von Satz 7.4 wissen wir, dass die Lösungsmenge $L(G)$ eines allgemeinen linearen Systems

$$A\vec{x} = \vec{b}$$

von der Form

$$L(G) = \vec{x}_P + L(hG)$$

ist, wobei $L(hG)$ die Lösungsmenge des zugehörigen homogenen Systems ist. Der Raum $L(G)$ ist also ein um den Vektor $\vec{x}_P$ verschobener Vektorraum.

9.4 Orthogonale Vektoren

Seien $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ in $\mathbb{R}^2$, bzw. $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}$ in $\mathbb{R}^3$.

Definition Das *Skalarprodukt* von $\vec{u}$ und $\vec{v}$ ist definiert durch

$$\vec{u} \cdot \vec{v} = u_1 v_1 + u_2 v_2, \quad \text{bzw.} \quad \vec{u} \cdot \vec{v} = u_1 v_1 + u_2 v_2 + u_3 v_3.$$

Man nennt dieses Produkt *Skalarprodukt*, weil das Ergebnis eine reelle Zahl, das heisst ein *Skalar* ist!

Sei φ der (kleinere) Zwischenwinkel von $\vec{u}$ und $\vec{v}$ (d.h. $0 \leq \varphi \leq \pi$). Dann gilt

$$\vec{u} \cdot \vec{v} = \|\vec{u}\| \|\vec{v}\| \cos \varphi.$$

Insbesondere stehen zwei Vektoren $\vec{u}$ und $\vec{v}$ senkrecht aufeinander (d.h. der Zwischenwinkel ist ein rechter Winkel), wenn $\vec{u} \cdot \vec{v} = 0$. Die Vektoren heissen in diesem Fall *orthogonal*. Man setzt fest, dass der Nullvektor senkrecht zu jedem Vektor steht.

Satz 9.9 Es gilt:

$$\vec{u} \text{ und } \vec{v} \text{ sind orthogonal} \iff \vec{u} \cdot \vec{v} = 0$$

Beispiel

Sind die Vektoren $\vec{u} = \begin{pmatrix} 1 \\ -3 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 6 \\ 2 \end{pmatrix}$ orthogonal?

Dieses Skalarprodukt kann man auf (Spalten-)Vektoren in $\mathbb{R}^n$ erweitern.

Definition Sind $\vec{u}$ und $\vec{v}$ Vektoren in $\mathbb{R}^n$ mit Komponenten $u_1, \dots, u_n$, bzw. $v_1, \dots, v_n$, dann ist das *Skalarprodukt* von $\vec{u}$ und $\vec{v}$ definiert durch

$$\vec{u} \cdot \vec{v} = \vec{u}^T \vec{v} = u_1 v_1 + \dots + u_n v_n.$$

Man nennt die Vektoren $\vec{u}$ und $\vec{v}$ *orthogonal*, wenn $\vec{u} \cdot \vec{v} = 0$.

Das Skalarprodukt hat die folgenden Eigenschaften.

Satz 9.10 Für Vektoren $\vec{u}, \vec{v}, \vec{w}$ in $\mathbb{R}^n$ und $k \in \mathbb{R}$ gilt:

- (i) $\vec{u} \cdot \vec{v} = \vec{v} \cdot \vec{u}$
- (ii) $\vec{u} \cdot (\vec{v} + \vec{w}) = \vec{u} \cdot \vec{v} + \vec{u} \cdot \vec{w}$
- (iii) $k(\vec{u} \cdot \vec{v}) = (k\vec{u}) \cdot \vec{v} = \vec{u} \cdot (k\vec{v})$
- (iv) $\vec{v} \cdot \vec{v} = \|\vec{v}\|^2 \geq 0$ und $\vec{v} \cdot \vec{v} = 0 \iff \vec{v} = \vec{0}$

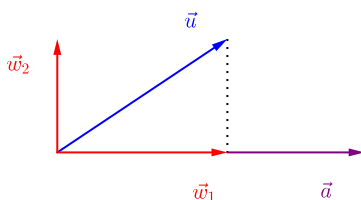
An dieser Stelle sei kurz an das Vektorprodukt von zwei Vektoren in $\mathbb{R}^3$ erinnert.

Definition Das *Vektorprodukt* $\vec{u} \times \vec{v}$ von zwei Vektoren $\vec{u}$ und $\vec{v}$ in $\mathbb{R}^3$ ist definiert durch

$$\vec{u} \times \vec{v} = \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} \times \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = \begin{pmatrix} u_2 v_3 - u_3 v_2 \\ u_3 v_1 - u_1 v_3 \\ u_1 v_2 - u_2 v_1 \end{pmatrix}.$$

Das Ergebnis ist also wieder ein *Vektor* in $\mathbb{R}^3$, wobei der Vektor $\vec{u} \times \vec{v}$ senkrecht auf $\vec{u}$ und auf $\vec{v}$ steht. Das Vektorprodukt ist jedoch nur für Vektoren in $\mathbb{R}^3$ definiert! Es wird auch Kreuzprodukt genannt.

Oft ist es praktisch, einen Vektor $\vec{u}$ in einen zu einem vorgegebenen Vektor $\vec{a} \neq \vec{0}$ parallelen Vektor $\vec{w}_1$ und einen dazu senkrechten Summanden $\vec{w}_2$ zu zerlegen.



Definition Der Vektor $\vec{w}_1$ wie oben beschrieben heisst *Orthogonalprojektion von $\vec{u}$ auf $\vec{a}$* und wird mit $\text{proj}_{\vec{a}}(\vec{u})$ bezeichnet.

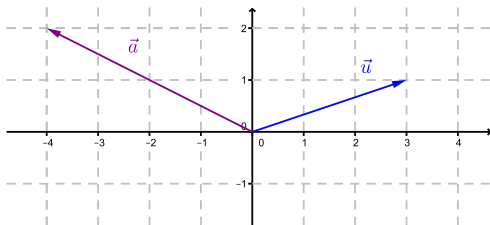
Der zweite Vektor ist dann gegeben durch $\vec{w}_2 = \vec{u} - \vec{w}_1$.

Satz 9.11 Für Vektoren $\vec{u}$ und $\vec{a} \neq \vec{0}$ in $\mathbb{R}^n$ gilt

$$\text{proj}_{\vec{a}}(\vec{u}) = \frac{\vec{u} \cdot \vec{a}}{\|\vec{a}\|^2} \vec{a}.$$

Beispiel

Seien $\vec{u} = \begin{pmatrix} 3 \\ 1 \end{pmatrix}$ und $\vec{a} = \begin{pmatrix} -4 \\ 2 \end{pmatrix}$.



Wir haben gesehen, dass es viele verschiedene Basen für den Vektorraum $\mathbb{R}^n$ gibt. Die Standardbasis $\vec{e}_1, \dots, \vec{e}_n$ zeichnet sich dadurch aus, dass die Vektoren alle die Länge 1 haben und je zwei Vektoren orthogonal zueinander sind. Manchmal ist es nützlich, eine andere Basis mit diesen beiden Eigenschaften zu verwenden.

Sei V ein Vektorraum in $\mathbb{R}^m$.

Definition Man nennt eine Basis $\vec{v}_1, \dots, \vec{v}_n$ eine *Orthonormalbasis* von V , wenn gilt

- (1) $\vec{v}_1, \dots, \vec{v}_n$ sind paarweise orthogonal und
- (2) $\vec{v}_1, \dots, \vec{v}_n$ haben die Länge 1.

Beispiel

Eine Orthonormalbasis von $\mathbb{R}^2$ bilden die beiden Vektoren $\vec{v}_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ und $\vec{v}_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix}$.

Orthonormalbasen haben die folgende wichtige Eigenschaft (welche aus Satz 9.10 folgt).

Satz 9.12 Sei $\vec{v}_1, \dots, \vec{v}_n$ eine Orthonormalbasis von V und $\vec{v} \in V$ beliebig. Dann gilt

$$\vec{v} = c_1 \vec{v}_1 + \dots + c_n \vec{v}_n \quad \text{mit } c_i = \vec{v} \cdot \vec{v}_i.$$

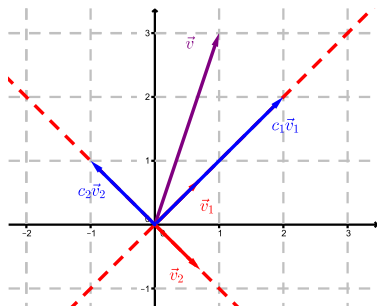
Die einzelnen Summanden $c_i \vec{v}_i$ sind dabei die Orthogonalprojektionen von $\vec{v}$ auf die Basisvektoren $\vec{v}_1, \dots, \vec{v}_n$! Denn wegen $\|\vec{v}_i\| = 1$ gilt

$$\text{proj}_{\vec{v}_i}(\vec{v}) = \frac{\vec{v} \cdot \vec{v}_i}{\|\vec{v}_i\|^2} \vec{v}_i = (\vec{v} \cdot \vec{v}_i) \vec{v}_i = c_i \vec{v}_i.$$

Beispiel

Sei $\vec{v}_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$, $\vec{v}_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix}$ die Orthonormalbasis von $\mathbb{R}^2$ von vorher. Sei $\vec{v} = \begin{pmatrix} 1 \\ 3 \end{pmatrix}$.

Zu bestimmen sind c_1, c_2 mit $\vec{v} = c_1 \vec{v}_1 + c_2 \vec{v}_2$.



Ausgehend von einer beliebigen Basis von V kann mit Hilfe des sogenannten *Gram-Schmidt-schen Orthogonalisierungsverfahrens* eine Orthonormalbasis konstruiert werden. Auf dieses Verfahren wollen wir hier aber nicht eingehen.

Definition Eine $n \times n$ -Matrix A heisst *orthogonal*, wenn gilt

$$A^T A = A A^T = E.$$

Orthogonal ist zum Beispiel die Matrix

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

Satz 9.13 Sei A eine orthogonale $n \times n$ -Matrix. Dann gilt:

- (1) Die Spalten (bzw. Zeilen) bilden eine Orthonormalbasis von $\mathbb{R}^n$.
- (2) A ist invertierbar und $A^{-1} = A^T$.
- (3) $\det(A) = \pm 1$.

Orthogonale Matrizen werden nächstes Semester eine wichtige Rolle spielen.

9.5 Näherungslösungen für unlösbare lineare Gleichungssysteme

Bei Messungen oder bei der Auswertung von statistischen Daten treten oft lineare Gleichungssysteme auf, die überbestimmt sind. Das heisst, das lineare System ist nicht lösbar. Man sucht deshalb nach geeigneten Näherungslösungen.

Definition Sei A eine $m \times n$ -Matrix und $\vec{b} \in \mathbb{R}^m$, so dass das lineare System $A\vec{x} = \vec{b}$ keine Lösung hat. Wir nennen einen Vektor $\vec{x}^*$ eine *Näherungslösung*, wenn

$$\|\vec{b} - A\vec{x}^*\|$$

minimal ist. Man nennt dies *lineares Ausgleichsproblem* oder *Methode der kleinsten Quadrate*.

Die zweite Bezeichnung geht darauf zurück, dass die Summe der Fehlerquadrate

$$\|\vec{b} - A\vec{x}^*\|^2 = \|\vec{r}\|^2 = r_1^2 + \cdots + r_m^2 \quad \text{für } \vec{b} - A\vec{x}^* = \vec{r} = \begin{pmatrix} r_1 \\ \vdots \\ r_m \end{pmatrix}$$

minimiert wird.

Satz 9.14 Jede Lösung $\vec{x}^*$ des (stets lösbaren) linearen Gleichungssystems

$$A^T A \vec{x} = A^T \vec{b}$$

ist eine Näherungslösung für das (unlösbare) System $A\vec{x} = \vec{b}$.

Beispiel

Gesucht ist die Gerade in der Ebene, welche die vier Punkte $(0, 1)$, $(1, 3)$, $(2, 4)$, $(3, 4)$ am besten approximiert.

Schreiben wir $mx + q = y$ für die Gerade und setzen die vier Punkte ein, so erhalten wir das

lineare System

$$\begin{aligned} q &= 1 \\ m + q &= 3 \\ 2m + q &= 4 \\ 3m + q &= 4 \end{aligned}$$

in den Unbekannten m und q , welches offensichtlich unlösbar ist. Wir haben hier

$$A = \begin{pmatrix} 0 & 1 \\ 1 & 1 \\ 2 & 1 \\ 3 & 1 \end{pmatrix} \quad \text{und} \quad \vec{b} = \begin{pmatrix} 1 \\ 3 \\ 4 \\ 4 \end{pmatrix}$$

und berechnen

$$A^T A = \begin{pmatrix} 14 & 6 \\ 6 & 4 \end{pmatrix} \quad \text{und} \quad A^T \vec{b} = \begin{pmatrix} 23 \\ 12 \end{pmatrix}.$$

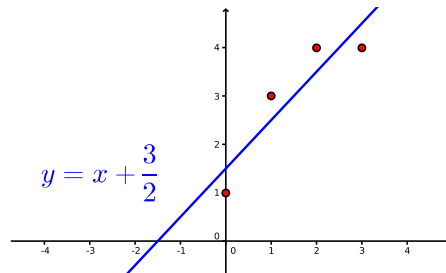
Wir erhalten also eine Näherungslösung $\vec{x}^* = \begin{pmatrix} m \\ q \end{pmatrix}$ als Lösung des Systems

$$\begin{pmatrix} 14 & 6 \\ 6 & 4 \end{pmatrix} \begin{pmatrix} m \\ q \end{pmatrix} = \begin{pmatrix} 23 \\ 12 \end{pmatrix}.$$

Mit Hilfe der Inversen von $A^T A$ erhalten wir

$$\begin{pmatrix} m \\ q \end{pmatrix} = \begin{pmatrix} 14 & 6 \\ 6 & 4 \end{pmatrix}^{-1} \begin{pmatrix} 23 \\ 12 \end{pmatrix} = \frac{1}{20} \begin{pmatrix} 4 & -6 \\ -6 & 14 \end{pmatrix} \begin{pmatrix} 23 \\ 12 \end{pmatrix} = \begin{pmatrix} 1 \\ \frac{3}{2} \end{pmatrix},$$

das heisst $m = 1$ und $q = \frac{3}{2}$. Die gesuchte Näherungsgerade ist also



Wir können zumindest geometrisch nachvollziehen, was hinter Satz 9.14 steckt. Ob ein lineares Gleichungssystem $A\vec{x} = \vec{b}$ eine Lösung hat oder nicht, hängt von den Spaltenvektoren $\vec{a}_1, \dots, \vec{a}_n$ (in $\mathbb{R}^m$) von A und dem Vektor $\vec{b}$ (in $\mathbb{R}^m$) ab. Sind

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} \quad \text{und} \quad \vec{x} = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix},$$

dann gilt

$$A\vec{x} = \begin{pmatrix} a_{11}x_1 + a_{12}x_2 + \cdots + a_{1n}x_n \\ a_{21}x_1 + a_{22}x_2 + \cdots + a_{2n}x_n \\ \vdots \\ a_{m1}x_1 + a_{m2}x_2 + \cdots + a_{mn}x_n \end{pmatrix} = x_1 \begin{pmatrix} a_{11} \\ a_{21} \\ \vdots \\ a_{m1} \end{pmatrix} + x_2 \begin{pmatrix} a_{12} \\ a_{22} \\ \vdots \\ a_{m2} \end{pmatrix} + \cdots + x_n \begin{pmatrix} a_{1n} \\ a_{2n} \\ \vdots \\ a_{mn} \end{pmatrix}.$$

Das heisst, das System $A\vec{x} = \vec{b}$ kann auch als

$$x_1\vec{a}_1 + \cdots + x_n\vec{a}_n = \vec{b}$$

geschrieben werden.

Das System $A\vec{x} = \vec{b}$ hat also genau dann eine Lösung $x_1, \dots, x_n$, falls sich $\vec{b}$ als Linearkombination der Spaltenvektoren $\vec{a}_1, \dots, \vec{a}_n$ von A schreiben lässt. Das ist genau dann der Fall, falls $\vec{b}$ im Spaltenraum von A (d.h. im Vektorraum aufgespannt von $\vec{a}_1, \dots, \vec{a}_n$) liegt. (Weiter gibt es genau eine Lösung, falls $\vec{a}_1, \dots, \vec{a}_n$ linear unabhängig sind, das heisst, eine Basis des Spaltenraums von A sind.)

Hat also das System $A\vec{x} = \vec{b}$ keine Lösung, dann liegt $\vec{b}$ nicht im Spaltenraum von A . Eine Näherungslösung $\vec{x}^*$ hat dann die Eigenschaft, dass $A\vec{x}^*$ die Orthogonalprojektion von $\vec{b}$ auf den Spaltenraum von A ist. Diese Orthogonalprojektion ist dabei eindeutig, die Näherungslösung $\vec{x}^*$ jedoch nicht unbedingt.

10 Fourierreihen

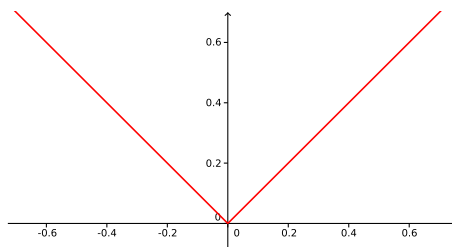
In einigen Bereichen der Physik und der Chemie (zum Beispiel bei der Spektralanalyse zur Bestimmung von chemischen Elementen) ist es aufschlussreich, eine reelle Funktion als unendliche Summe von trigonometrischen Funktionen darstellen zu können. Man nennt diese Darstellung Fourierreihe der Funktion.

Bei der Bild- und Audiokompression beispielsweise genügt es, eine Funktion durch eine endliche Summe von trigonometrischen Funktionen näherungsweise zu beschreiben. Diese endliche Summe nennt man Fourierpolynom der Funktion.

Wir wollen im Folgenden untersuchen, wie man ein Fourierpolynom und die Fourierreihe einer Funktion bestimmt und mit welchem Vektorraum die Fourierreihe zusammenhängt.

10.1 Fourierpolynome

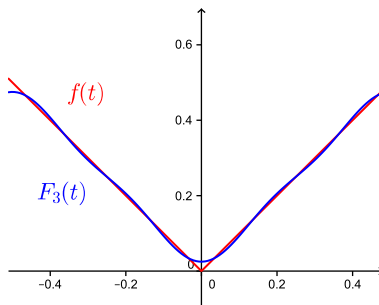
Nehmen wir an, wir haben ein (akustisches) Signal, das wir auf dem Computer abspeichern möchten. Das Signal könnte wie folgt aussehen:



Dieses Signal wird in konstanten Zeitabständen, zum Beispiel alle Hundertstelsekunden, gemessen. Es liegen daher nur die Funktionswerte zu den Zeitpunkten $-0.50, -0.49, \dots, 0.49, 0.50$ vor. Die 100 gemessenen Funktionswerte könnten nun abgespeichert werden, doch bei einer grösseren Anzahl von Messungen würde dies zu Speicherproblemen führen. Wie kann also das Signal mit möglichst geringem Speicheraufwand aber zugleich möglichst geringem Informationsverlust abgespeichert werden?

Tatsächlich ist es möglich, nahezu jede beliebige Funktion durch ein sogenanntes *trigonometrisches Polynom* anzunähern. Nehmen wir an, das Signal sei durch die Funktion $f(t) = |t|$ gegeben und vergleichen wir dies mit dem trigonometrischen Polynom

$$F_3(t) = \frac{1}{4} - \frac{2}{\pi^2} \cos(2\pi t) - \frac{2}{9\pi^2} \cos(6\pi t).$$



Wir sehen, dass die Näherung auf dem gesamten Intervall $[-\frac{1}{2}, \frac{1}{2}]$ recht gut ist. Wären wir mit dieser Näherung zufrieden (und würden wir gewisse Informationsverluste in Kauf nehmen), dann müssten wir nur die drei Koeffizienten $\frac{1}{4}$, $-\frac{2}{\pi^2}$ und $-\frac{2}{9\pi^2}$ speichern!

Betrachten wir nun allgemein reelle Funktionen auf dem Intervall $[-\frac{T}{2}, \frac{T}{2}]$ für $T > 0$.

Definition Ein *trigonometrisches Polynom* oder *Fourierpolynom* ist ein Ausdruck der Form

$$\begin{aligned} F_n(t) &= \frac{a_0}{2} + a_1 \cos(\omega t) + b_1 \sin(\omega t) + \cdots + a_n \cos(n\omega t) + b_n \sin(n\omega t) \\ &= \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos(k\omega t) + b_k \sin(k\omega t)) \end{aligned}$$

wobei $\omega = \frac{2\pi}{T}$. Die Koeffizienten a_k und b_k heissen *Fourierkoeffizienten* von F_n .

Im Beispiel oben ist $F_3(t)$ ein Fourierpolynom vom Grad 3 mit den Koeffizienten

$$a_0 = \frac{1}{2}, a_1 = -\frac{2}{\pi^2}, a_2 = 0, a_3 = -\frac{2}{9\pi^2} \quad \text{und} \quad b_1 = b_2 = b_3 = 0.$$

Anstelle von Potenzen t^k enthält das Fourierpolynom die Funktionen $\cos(k\omega t)$ und $\sin(k\omega t)$, welche periodisch mit der Periode $T = \frac{2\pi}{\omega}$ sind. Es genügt daher, F_n auf einem Intervall der Länge T zu untersuchen, zum Beispiel auf $[-\frac{T}{2}, \frac{T}{2}]$.

Wie müssen nun die Koeffizienten a_k und b_k gewählt werden, damit $F_n(t)$ eine möglichst gute Näherung einer gegebenen Funktion $f(t)$ auf dem ganzen Intervall $[-\frac{T}{2}, \frac{T}{2}]$ ist?

Ähnlich wie im Abschnitt 9.5 über Näherungslösungen von linearen Systemen soll die "quadratische Abweichung" $|f(t) - F_n(t)|^2$ minimal sein, und zwar über das ganze Intervall $[-\frac{T}{2}, \frac{T}{2}]$ betrachtet. Das heisst, der Wert

$$\int_{-\frac{T}{2}}^{\frac{T}{2}} (f(t) - F_n(t))^2 dt$$

soll minimal sein.

Satz 10.1 Sei $f(t)$ eine (stückweise) stetige und beschränkte Funktion und $n \in \mathbb{N}$. Dann sind die Koeffizienten des bestapproximierenden trigonometrischen Polynoms vom Grad n gegeben durch

$$\begin{aligned} a_k &= \frac{2}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} \cos(k\omega t) f(t) dt \\ b_k &= \frac{2}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} \sin(k\omega t) f(t) dt \end{aligned}$$

für $k = 0, \dots, n$.

Für $k = 0$ gilt speziell

$$b_0 = \frac{2}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} 0 dt = 0 \quad \text{und} \quad a_0 = \frac{2}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} f(t) dt.$$

Die Hälfte der Fourierkoeffizienten ist jeweils 0, falls die Funktion f gerade (das heisst $f(-t) = f(t)$) oder ungerade (d.h. $f(-t) = -f(t)$) ist.

Satz 10.2 *Es gilt:*

$$\begin{aligned} f(t) \text{ gerade} &\implies b_k = 0 \text{ für alle } k = 0, \dots, n \\ f(t) \text{ ungerade} &\implies a_k = 0 \text{ für alle } k = 0, \dots, n \end{aligned}$$

Ist nämlich $f(t)$ gerade, dann ist $\sin(k\omega t)f(t)$ eine ungerade Funktion (da $\sin(t)$ ungerade ist) und damit ist das Integral über das Intervall $[-\frac{T}{2}, \frac{T}{2}]$ gleich 0. Ist $f(t)$ ungerade, dann ist die Funktion $\cos(k\omega t)f(t)$ ungerade, da $\cos(t)$ gerade ist.

Beispiel

Wir wollen überprüfen, ob das Fourierpolynom $F_3(t)$ zur Funktion $f(t) = |t|$ vom Beispiel von Seite 162 mit den Aussagen von Satz 10.1 übereinstimmt.

Wir betrachten hier also das Intervall $[-\frac{1}{2}, \frac{1}{2}]$, das heisst $T = 1$ und damit $\omega = \frac{2\pi}{T} = 2\pi$. Das Fourierpolynom vom Grad 3 hat nun die Form

$$\begin{aligned} F_3(t) = & \frac{a_0}{2} + a_1 \cos(2\pi t) + b_1 \sin(2\pi t) + a_2 \cos(4\pi t) + b_2 \sin(4\pi t) \\ & + a_3 \cos(6\pi t) + b_3 \sin(6\pi t), \end{aligned}$$

wobei die Koeffizienten a_k, b_k durch Satz 10.1 gegeben sind. Da die Funktion $f(t) = |t|$ gerade ist, gilt (mit Satz 10.2) für die Koeffizienten

Der Koeffizient a_0 ist einfach zu bestimmen:

Den Koeffizienten a_1 berechnen wir mit Hilfe von partieller Integration:

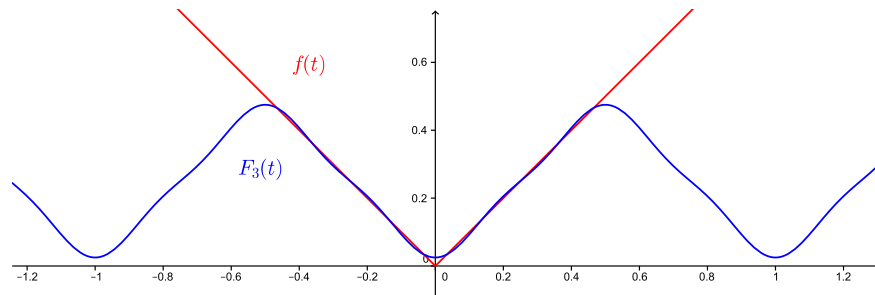
Analog (mit partieller Integration) finden wir die Koeffizienten

$$a_2 = 4 \int_0^{\frac{1}{2}} t \cos(4\pi t) dt = 0 \quad \text{und} \quad a_3 = 4 \int_0^{\frac{1}{2}} t \cos(6\pi t) dt = -\frac{2}{9\pi^2}.$$

Das Fourierpolynom vom Grad 3 lautet also (wie auf Seite 162)

$$F_3(t) = \frac{1}{4} - \frac{2}{\pi^2} \cos(2\pi t) - \frac{2}{9\pi^2} \cos(6\pi t).$$

Die Funktion $f(t) = |t|$ ist natürlich für alle $t \in \mathbb{R}$ definiert, doch die Approximation von f durch das Fourierpolynom F_3 ist nur auf dem Intervall $[-\frac{1}{2}, \frac{1}{2}]$ gut:



Ist die Funktion f periodisch mit der Periode T , das heisst gilt

$$f(t + T) = f(t) \quad \text{für alle } t \in \mathbb{R},$$

dann ist das Fourierpolynom auf ganz $\mathbb{R}$ eine gute Approximation. In diesem Fall kann anstelle des Intervalls $[-\frac{T}{2}, \frac{T}{2}]$ auch $[0, T]$ oder jedes andere Intervall der Länge T verwendet werden.

In der Praxis ist meist nicht die Funktionsgleichung von $f(t)$ gegeben, sondern nur die Funktionswerte $f(t_k)$ zu bestimmten Zeitpunkten t_k . Die Integrale zur Berechnung der Fourierkoeffizienten werden dann durch Summen approximiert.

10.2 Fourierreihen

Das Fourierpolynom approximiert die Funktion umso besser, je grösser der Grad n des Polynoms ist.

Satz 10.3 Sei f eine (stückweise) stetige, beschränkte Funktion mit zugehörigen Fourierpolynomen F_n . Dann kann die Approximation $f(t) \approx F_n(t)$ beliebig genau gemacht werden, indem der Grad n gross genug gewählt wird.

Ist f differenzierbar an einer Stelle t , dann gilt

$$f(t) = \frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos(k\omega t) + b_k \sin(k\omega t)).$$

Die Reihe auf der rechten Seite nennt man Fourierreihe $\mathcal{F}(t)$ von $f(t)$.

Beispiel

Wir betrachten die Funktion

$$f(t) = \begin{cases} 1 & \text{für } t \geq 0 \\ 0 & \text{für } t < 0 \end{cases}$$

im Intervall $[-\pi, \pi]$.

Das Intervall hat nun die Länge $T = 2\pi$, das heisst $\omega = \frac{2\pi}{T} = 1$. Die Fourierreihe hat also die Form

$$\mathcal{F}(t) = \frac{a_0}{2} + a_1 \cos(t) + b_1 \sin(t) + a_2 \cos(2t) + b_2 \sin(2t) + \dots$$

Für die Fourierkoeffizienten a_k gilt

$$a_0 = 1 \quad \text{und} \quad a_k = \frac{1}{\pi} \int_0^\pi \cos(kt) dt = 0 \quad \text{für } k \geq 1.$$

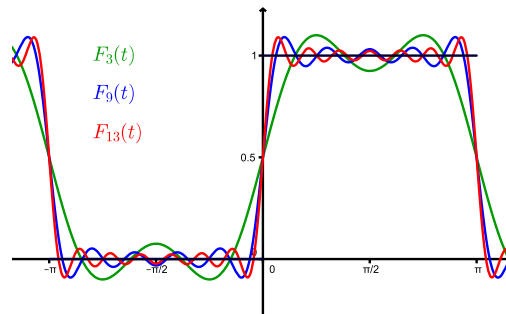
Und für die Koeffizienten b_k erhalten wir

$$b_k = \frac{1}{\pi} \int_0^\pi \sin(kt) dt = \frac{1 - (-1)^k}{k\pi} \quad \text{für } k \geq 1.$$

Die Fourierreihe lautet demnach

$$\begin{aligned} \mathcal{F}(t) &= \frac{1}{2} + \sum_{k=1}^{\infty} \frac{1 - (-1)^k}{k\pi} \sin(kt) \\ &= \frac{1}{2} + \frac{2}{\pi} \sin(t) + \frac{2}{3\pi} \sin(3t) + \frac{2}{5\pi} \sin(5t) + \dots \end{aligned}$$

Nach Satz 10.3 gilt $\mathcal{F}(t) = f(t)$ für jedes $t \neq 0$ in $[-\pi, \pi]$ (denn f ist differenzierbar in $t \neq 0$).



10.3 Eine Orthonormalbasis bestehend aus Funktionen

In diesem Abschnitt untersuchen wir die Fourierreihe einer Funktion mit Hilfe der Theorie über Orthonormalbasen von Vektorräumen.

Wir haben in Kapitel 9 gesehen, dass die Menge der reellen Funktionen von $[0, 1]$ nach $\mathbb{R}$ einen reellen Vektorraum bilden. Ebenso bildet die Menge der im Intervall $[-\pi, \pi]$ stetigen Funktionen einen (reellen) Vektorraum. Wir bezeichnen ihn mit $\mathcal{C} = \mathcal{C}[-\pi, \pi]$.

Auf dem Vektorraum $\mathbb{R}^n$ haben wir in Abschnitt 9.4 ein Skalarprodukt definiert. Die charakteristischen Eigenschaften des Skalarprodukts sind in Satz 9.10 aufgelistet. Wir können nun auch auf dem Vektorraum $\mathcal{C}$ der stetigen Funktionen ein Skalarprodukt definieren.

Definition Das *Skalarprodukt* von zwei Funktionen f, g in $\mathcal{C}$ ist definiert durch

$$\langle f, g \rangle = \int_{-\pi}^{\pi} f(x)g(x)dx.$$

Anstelle von $\langle f, g \rangle$ könnte man (analog zu Vektoren) auch $f \cdot g$ schreiben, doch die Bezeichnung $\langle f, g \rangle$ ist üblicher.

Der Name Skalarprodukt ist gerechtfertigt, da genau die zu Satz 9.10 analogen Eigenschaften gelten!

Für Vektoren $\vec{v} \in \mathbb{R}^n$ gibt es den Zusammenhang $\|\vec{v}\|^2 = \vec{v} \cdot \vec{v}$ zwischen der Länge eines Vektors und dem Skalarprodukt des Vektors mit sich selbst. Für Funktionen f in $\mathcal{C}$ definieren wir die *Länge* oder *Norm* von f durch

$$\|f\| = \sqrt{\langle f, f \rangle} = \sqrt{\int_{-\pi}^{\pi} f(x)^2 dx}.$$

Und analog zu Vektoren nennen wir zwei Funktionen $f, g \in \mathcal{C}$ *orthogonal*, wenn gilt

$$\langle f, g \rangle = 0.$$

Satz 10.4 *Die Funktionen*

$$\frac{1}{\sqrt{\pi}}, \frac{1}{\sqrt{\pi}} \cos(nx), \frac{1}{\sqrt{\pi}} \sin(nx) \quad \text{für } n \in \mathbb{N}$$

bilden eine Orthonormalbasis für den Vektorraum $\mathcal{C} = \mathcal{C}[-\pi, \pi]$.

Konkret bedeutet der Satz, dass gilt

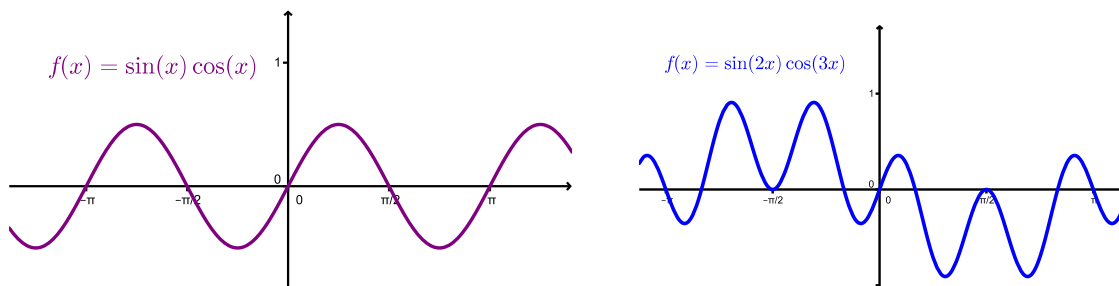
$$\int_{-\pi}^{\pi} \cos(nx) \cos(mx) dx = \begin{cases} \pi & \text{für } m = n \\ 0 & \text{für } m \neq n \end{cases},$$

$$\int_{-\pi}^{\pi} \sin(nx) \sin(mx) dx = \begin{cases} \pi & \text{für } m = n \\ 0 & \text{für } m \neq n \end{cases},$$

und

$$\int_{-\pi}^{\pi} \sin(nx) \cos(mx) dx = 0.$$

Die ersten beiden Integrale können mit Hilfe der partiellen Integration berechnet werden. Das dritte Integral ist klar, da der Integrand eine ungerade Funktion ist. Oder bildlich für zwei Beispiele:



Dass sich jede stetige Funktion f als Fourierreihe darstellen lässt, kann man also so interpretieren, dass sich jede solche Funktion f als unendliche Linearkombination der Basisfunktionen aus Satz 10.4 schreiben lässt (der Vektorraum $\mathcal{C}$ ist ja unendlich-dimensional). Die Fourierkoeffizienten sind nun einfach die Koeffizienten dieser Linearkombination. Doch da wir eine

Orthonormalbasis haben, sind (gemäss Satz 9.12) die Fourierkoeffizienten a_n , b_n nichts anderes als die Orthogonalprojektionen von f auf die Basisfunktionen:

$$\begin{aligned}\langle f(x), \frac{1}{\sqrt{\pi}} \cos(nx) \rangle \frac{1}{\sqrt{\pi}} \cos(nx) &= \left(\frac{1}{\pi} \int_{-\pi}^{\pi} \cos(nx) f(x) dx \right) \cos(nx) = a_n \cos(nx) \\ \langle f(x), \frac{1}{\sqrt{\pi}} \sin(nx) \rangle \frac{1}{\sqrt{\pi}} \sin(nx) &= \left(\frac{1}{\pi} \int_{-\pi}^{\pi} \sin(nx) f(x) dx \right) \sin(nx) = b_n \sin(nx)\end{aligned}$$

Zu erwähnen bleibt, dass sich die Schreibweise der Fourierreihe und der Nachweis von Satz 10.4 vereinfachen, wenn man komplexe Zahlen zu Hilfe nimmt. Wegen der Eulerschen Identität gilt

$$\mathcal{F}(t) = \sum_{k=-\infty}^{\infty} c_k e^{ik\omega t}$$

mit den komplexen Koeffizienten

$$c_k = \frac{1}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} e^{-ik\omega t} f(t) dt .$$

11 Boolesche Algebra

Die Boolesche Algebra ist eine “Algebra der Logik”, die George Boole (1815 – 1864) als erster entwickelt hat. Sie ist die Grundlage für den Entwurf von elektronischen Schaltungen und von Computerprogrammen. Die Boolesche Algebra kennt nur die beiden Zustände “wahr” und “falsch”, die in einem Schaltkreis den grundlegenden Zuständen “Strom fließt” und “Strom fließt nicht” entsprechen. Diese beiden Zustände werden im Folgenden durch die Zahlen 1 und 0 beschrieben.

11.1 Grundlegende Operationen und Gesetze

Die Boolesche Algebra geht von der Menge $\{0, 1\}$ aus.

Definition Auf der Menge $\{0, 1\}$ sind die folgenden drei Operationen definiert.

- (1) Die *Konjunktion* $\wedge$ (*Und-Verknüpfung*) ist eine Verknüpfung, die von zwei Argumenten abhängt. Sie ist genau dann 1, wenn das erste *und* das zweite Argument 1 ist, und in jedem anderen Fall 0. Man liest $a \wedge b$ als “*a und b*”.
- (2) Die *Disjunktion* $\vee$ (*Oder-Verknüpfung*) ist eine weitere von zwei Argumenten abhängige Verknüpfung. Sie ist genau dann 1, wenn das erste *oder* das zweite Argument 1 ist, und sonst 0. Man liest $a \vee b$ als “*a oder b*”.
- (3) Die *Negation* $\neg$ (*Nicht-Operator*) hängt nur von einem Argument ab. Sie ist 0, wenn das Argument 1 ist, und 1, wenn das Argument 0 ist. Man liest $\neg a$ als “*nicht a*”.

Diese drei Verknüpfungen können übersichtlich mit Verknüpfungstafeln dargestellt werden:

$\wedge$	0	1
0	0	0
1	0	1

$\vee$	0	1
0	0	1
1	1	1

x	$\neg x$
0	1
1	0

Diese drei Operationen können nun mehrfach hintereinander ausgeführt werden, um weitere boolesche Ausdrücke zu erhalten. Dabei haben die Operationen unterschiedliche Priorität: $\neg$ kommt vor $\wedge$, und $\wedge$ kommt vor $\vee$. Für andere Prioritäten muss man Klammern setzen.

Beispiel: $\neg 0 \vee 1 \wedge 0 =$

Die folgenden Rechengesetze für $\wedge$, $\vee$ und $\neg$ erinnern an die Rechengesetze für reelle Zahlen, wenn wir $\wedge$ als Multiplikation und $\vee$ als Addition interpretieren.

Satz 11.1 Für alle x, y, z in $\{0, 1\}$ gelten die folgenden Gesetze.

- (a) *Kommutativgesetz:* $x \wedge y = y \wedge x$ und $x \vee y = y \vee x$
- (b) *Assoziativgesetz:* $x \wedge (y \wedge z) = (x \wedge y) \wedge z$ und $x \vee (y \vee z) = (x \vee y) \vee z$
- (c) *Distributivgesetz:* $x \vee (y \wedge z) = (x \vee y) \wedge (x \vee z)$ und $x \wedge (y \vee z) = (x \wedge y) \vee (x \wedge z)$
- (d) *Existenz neutraler Elemente:* $1 \wedge x = x$ und $0 \vee x = x$
- (e) *Existenz des Komplements:* $x \wedge \neg x = 0$ und $x \vee \neg x = 1$

In Satz 11.1 geht jeweils der zweite Teil des Gesetzes aus dem ersten Teil hervor, indem man $\wedge$ und $\vee$ sowie 1 und 0 vertauscht.

Satz 11.2 Jede Aussage, die aus Satz 11.1 folgt, bleibt gültig, wenn die Operationen $\wedge$ und $\vee$ sowie die Zahlen 1 und 0 überall gleichzeitig vertauscht werden.

Diese Eigenschaft der Booleschen Algebra heisst *Dualität*. Eine Aussage, die durch Vertauschen von $\wedge$ und $\vee$ sowie 1 und 0 aus einer anderen hervorgeht, heisst zu dieser *dual*.

Mit Hilfe dieser Eigenschaft können weitere Gesetze hergeleitet werden.

Satz 11.3 Für alle x, y, z in $\{0, 1\}$ gelten die folgenden Gesetze.

- (a) *Absorptionsgesetze:* $x \wedge (x \vee y) = x$ und $x \vee (x \wedge y) = x$
- (b) *Idempotenzgesetze:* $x \vee x = x$ und $x \wedge x = x$
- (c) *Involutionsgesetz:* $\neg(\neg x) = x$
- (d) *Gesetze von de Morgan:* $\neg(x \wedge y) = \neg x \vee \neg y$ und $\neg(x \vee y) = \neg x \wedge \neg y$

Die Gesetze der Sätze 11.1–11.3 kann man beweisen, indem man alle möglichen Werte für x und y einsetzt und überprüft, ob jeweils die linke und die rechte Seite einer Gleichung übereinstimmen. Bei Satz 11.3 geht es teilweise eleganter, indem man schon bekannte Gesetze anwendet. Die Absorptionsgesetze kann man beispielsweise wie folgt nachweisen:

Wegen der Dualität folgt das zweite Absorptionsgesetz.

11.2 Boolesche Funktionen und ihre Normalformen

Definition Eine n -stellige boolesche Funktion ist eine Abbildung

$$f : \{0, 1\}^n \longrightarrow \{0, 1\}.$$

Jedem n -Tupel $(x_1, \dots, x_n)$ mit $x_i \in \{0, 1\}$ wird eindeutig eine Zahl

$$f(x_1, \dots, x_n) \in \{0, 1\}$$

zugeordnet.

Beispiel

Für $n = 1$ gibt es genau 4 verschiedene boolesche Funktionen: Die Nullfunktion $f(x) = 0$, die Einsfunktion $f(x) = 1$, die Identität $f(x) = x$ und die Negation $f(x) = \neg x$.

Satz 11.4 Es gibt genau $2^{(2^n)}$ verschiedene n -stellige boolesche Funktionen.

Es gibt also $2^{(2^2)} = 16$ 2-stellige boolesche Funktionen. In der folgenden Wertetabelle sind alle Funktionen aufgelistet.

x	y	f_1	f_2	f_3	f_4	f_5	f_6	f_7	f_8	f_9	f_{10}	f_{11}	f_{12}	f_{13}	f_{14}	f_{15}	f_{16}
0	0	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
0	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
1	0	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
1	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1

Einige dieser Funktionen sind uns schon bekannt:

Drei weitere Funktionen sind wichtig und als einfacher boolescher Ausdruck beschreibbar.

Die Funktion

$$f_9(x, y) = \neg(x \vee y)$$

ist die negierte Oder-Verknüpfung. Nach dem englischen “not or” bezeichnet man sie als *NOR*-Verknüpfung.

Analog ist die Funktion

$$f_{15}(x, y) = \neg(x \wedge y)$$

die sogenannte *NAND*-Verknüpfung (nach “not and”).

Die Funktion

$$f_7(x, y) = (x \vee y) \wedge \neg(x \wedge y)$$

ergibt genau dann 1, wenn $x \neq y$ ist, das heisst, wenn *entweder* x *oder* y gleich 1 ist. Nach dem englischen “exclusive or” bezeichnet man sie als *XOR*-Verknüpfung.

Bei der letzten Funktion sieht man, dass es gar nicht so einfach ist, von der Wertetabelle einer Funktion auf einen booleschen Ausdruck zu schliessen. Dabei ist ein boolescher Ausdruck für eine Funktion nicht eindeutig. Wir wollen deshalb im Folgenden untersuchen, wie man erstens überhaupt einen booleschen Ausdruck für eine Funktion findet und zweitens, wie man den gefundenen Ausdruck vereinfachen kann.

Normalformen

Ausgehend von der Wertetabelle einer booleschen Funktion kann mit Hilfe eines Verfahrens der boolesche Ausdruck der Funktion in sogenannter disjunktiver, bzw. konjunktiver Normalform aufgestellt werden.

Definition Eine *Vollkonjunktion* ist ein boolescher Ausdruck, in dem alle Variablen genau einmal vorkommen und durch $\wedge$ (konjunktiv) verbunden sind. Dabei dürfen die Variablen auch negiert auftreten.

Ein Ausdruck liegt in der *disjunktiven Normalform* vor, wenn er aus Vollkonjunktionen besteht, die durch $\vee$ (disjunktiv) verknüpft sind.

Beispiel

Der Boolesche Ausdruck

$$(x \wedge \neg y \wedge \neg z) \vee (\neg x \wedge y \wedge \neg z) \vee (x \wedge \neg y \wedge z)$$

liegt in disjunktiver Normalform vor.

Das folgende Beispiel soll zeigen, wie man von der Wertetabelle einer booleschen Funktion ihren booleschen Ausdruck in disjunktiver Normalform erhält.

Beispiel

Wir betrachten die 3-stellige boolesche Funktion f , die durch die folgende Wertetabelle gegeben ist.

Zeile	x	y	z	$f(x, y, z)$
1	0	0	0	0
2	0	0	1	0
3	0	1	0	1
4	0	1	1	1
5	1	0	0	0
6	1	0	1	0
7	1	1	0	0
8	1	1	1	1

1. *Schritt:* Wir suchen diejenigen Zeilen, die den Funktionswert 1 haben.

2. *Schritt:* Für jede dieser Zeilen stellen wir die Vollkonjunktion auf, die für die Variablen dieser Zeile den Wert 1 liefert.

3. *Schritt:* Diese Vollkonjunktionen werden durch $\vee$ verknüpft.

Dieser Ausdruck ist nun genau dann 1, wenn eine der drei Vollkonjunktionen gleich 1 ist, was genau für die Zeilen 3, 4 und 8 der Fall ist. Also haben wir die disjunktive Normalform von f gefunden.

Die disjunktive Normalform besteht demnach aus genau so vielen Vollkonjunktionen, wie in der Wertetabelle der Funktionswert 1 vorkommt. Die disjunktive Normalform ist also ideal, wenn der Funktionswert 1 nicht oft vorkommt. Ansonsten stellt man besser die konjunktive Normalform auf.

Definition Eine *Volldisjunktion* ist ein boolescher Ausdruck, in dem alle Variablen genau einmal vorkommen und durch $\vee$ (disjunktiv) verbunden sind. Dabei dürfen die Variablen auch negiert auftreten.

Ein Ausdruck liegt in der *konjunktiven Normalform* vor, wenn er aus Volldisjunktionen besteht, die durch $\wedge$ (konjunktiv) verknüpft sind.

Beispiel

Der Boolesche Ausdruck

$$(\neg x \vee y \vee \neg z) \wedge (\neg x \vee \neg y \vee z) \wedge (x \vee \neg y \vee z)$$

liegt in konjunktiver Normalform vor.

Das Verfahren, mit dem man von der Wertetabelle einer booleschen Funktion zu ihrem Ausdruck in konjunktiver Normalform gelangt, ist dual zum vorherigen Verfahren.

Vereinfachen von booleschen Ausdrücken

Die disjunktive, bzw. konjunktive Normalform einer booleschen Funktion ist oft ein unnötig langer Ausdruck. Als nächstes sollte man also diesen booleschen Ausdruck vereinfachen (wenn möglich). Das kann man systematisch tun, und zwar zum Beispiel mit dem Verfahren von Karnaugh und Veitch. Wir wollen hier nur kurz andeuten, wie das funktioniert.

Das *Verfahren von Karnaugh und Veitch* geht von der disjunktiven Normalform aus. Die Grundidee des Verfahrens ist, den Ausdruck systematisch so umzuformen, dass Terme der Form $x \vee \neg x$ entstehen. Nach Satz 11.1(e) haben diese Terme stets den Wert 1 und können in einer Konjunktion weggelassen werden.

Beispiel

Wir vereinfachen die 3-stellige boolesche Funktion vom Beispiel auf der Seite 172.

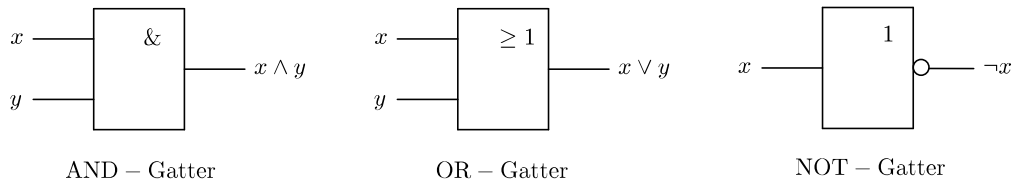
$$f(x, y, z) = (\neg x \wedge y \wedge \neg z) \vee (\neg x \wedge y \wedge z) \vee (x \wedge y \wedge z)$$

Der Vorteil des Verfahrens von Karnaugh und Veitch ist nun, dass diese Umformungen nicht auf gut Glück von Hand durchgeführt werden müssen, sondern graphisch am sogenannten *KV-Diagramm* abgelesen werden können. Darauf wollen wir aber nicht näher eingehen.

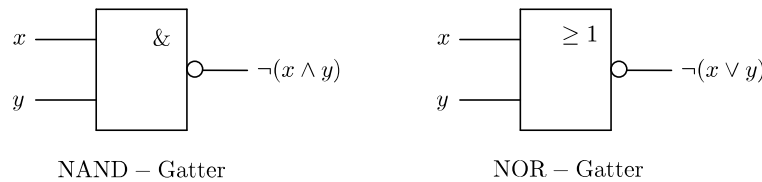
11.3 Logische Schaltungen

Jeder Computer ist aus logischen Schaltungen aufgebaut. Dabei ist eine logische Schaltung nichts anderes als eine physikalische Realisierung einer booleschen Funktion. Die beiden Zustände 0 und 1 der Booleschen Algebra werden durch unterschiedliche elektrische Spannungen realisiert. Meist entspricht der Zustand 0 der Spannung 0 (oder einer minimalen Spannung $U_{\min}$) und der Zustand 1 einer maximalen Spannung $U_{\max}$. Dabei sind gewisse Toleranzbereiche um diese Spannungen erlaubt.

Die grundlegenden booleschen Operationen $\wedge$, $\vee$ und $\neg$ werden durch elektronische Bauteile umgesetzt, die man *Gatter* nennt. Solche Gatter kann man mit einfachen Schaltern und Relais verwirklichen. In Schaltplänen werden Gatter durch ihre jeweiligen Schaltsymbole dargestellt. Die Schaltsymbole der drei Grundgatter AND, OR und NOT sind die Folgenden.



Diese drei Grundgatter kann man nun hintereinanderschalten. Dabei werden vor- oder nachgeschaltete NOT-Gatter am Eingang bzw. Ausgang vereinfacht als Kreis symbolisiert. Das NAND- und das NOR-Gatter sind demnach wie folgt.



Der folgende Satz sagt aus, dass entweder NAND- oder NOR-Gatter genügen, um jede beliebige boolesche Funktion zu verwirklichen.

Satz 11.5 *Die drei booleschen Grundoperationen $\wedge$, $\vee$ und $\neg$ können als Hintereinanderausführung von ausschliesslich NAND-Funktionen oder ausschliesslich NOR-Funktionen geschrieben werden.*

Der Beweis dieses Satzes nutzt die Gesetze von de Morgan (Satz 11.3(d)). Wir zeigen exemplarisch den ersten Teil des Satzes, wobei wir $\text{NAND}(x, y) = \neg(x \wedge y)$ schreiben.

$$\begin{aligned} x \wedge y &= (x \wedge y) \vee 0 = \neg(\neg(x \wedge y) \wedge 1) = \text{NAND}(\text{NAND}(x, y), 1) \\ x \vee y &= (x \wedge 1) \vee (y \wedge 1) = \neg(\neg(x \wedge 1) \wedge \neg(y \wedge 1)) = \text{NAND}(\text{NAND}(x, 1), \text{NAND}(y, 1)) \\ \neg x &= \neg(x \wedge 1) = \text{NAND}(x, 1) \end{aligned}$$

Wenn wir nun eine beliebige boolesche Funktion mit Hilfe der drei Grundgatter und der NAND- und NOR-Gatter realisieren wollen, können wir wie folgt vorgehen:

1. Die Wertetabelle der booleschen Funktion aufstellen.
2. Die disjunktive Normalform der Funktion ablesen.
3. Die disjunktive Normalform vereinfachen.
4. Der vereinfachte Ausdruck mit einer Gatterschaltung realisieren.

Beispiel

Wir wollen die einfachste Form einer Rechenschaltung realisieren, nämlich die Addition von zwei einstelligigen Binärzahlen. Wenn beide Bits gleich 1 sind, entsteht eine zweistellige Binärzahl, das heisst, ein Übertrag in die nächsthöhere Binärstelle:

$$0 + 0 = 0, 0 + 1 = 1, 1 + 0 = 1, 1 + 1 = 10$$

Für jede der beiden Binärstellen benötigt die Schaltung einen Ausgang. Wir bezeichnen diese beiden Ausgänge mit s (für Summe) und u (für Übertrag). Die Wertetabelle sieht damit wie folgt aus.

x	y	u	s
0	0	0	0
0	1	0	1
1	0	0	1
1	1	1	0

Für beide Ausgänge können wir die disjunktive Normalform aus der Tabelle ablesen:

Beide Ausdrücke können nicht weiter vereinfacht werden. Wir können also direkt die Schaltung angeben:

Diese Schaltung nennt man *Halbaddierer*. Um mehrstellige Binärzahlen addieren zu können, reichen Halbaddierer nicht mehr aus. Denn zur Summe zweier Bits muss dann im Allgemeinen noch der Übertrag aus der vorherigen Stelle addiert werden. Insgesamt müssen also an jeder Stelle drei Bits addiert werden. Die Schaltung, die drei Bits addiert, nennt man *Volladdierer*. Ein Volladdierer kann aus zwei Halbaddierern und einem OR-Gatter zusammengesetzt werden (deshalb der Name Halbaddierer). Allgemein kann man durch Zusammenschalten von $n - 1$ Volladdierern und einem Halbaddierer zwei n -stellige Binärzahlen addieren. Solche Addierwerke bilden die Grundlage der Computertechnik.