

Mathematik I für ET/IT und ITE
WS 25/26

Annett Püttmann

Inhaltsverzeichnis

Kapitel I. Zahlen(bereiche) und mathematische Symbole	1
1. Die natürlichen Zahlen	1
2. Die ganzen Zahlen	3
3. Die rationalen Zahlen	6
4. Die reellen Zahlen	7
5. Die komplexen Zahlen	10
6. Symbole	15
Kapitel II. Formeln, Gleichungen und Ungleichungen	17
1. Summenformel der geometrischen Reihe	17
2. Gaußsche Summenformel	18
3. Dreiecksungleichung	18
4. Arithmetisches und geometrisches Mittel	19
5. Schwarzsche Ungleichung	20
6. Der binomische Satz	21
7. Bernoullische Ungleichung	25
8. Qualitatives Wachstum von Potenzen und Fakultäten	26
Kapitel III. Primzahlen	31
Kapitel IV. Zahldarstellungen bezüglich verschiedener Basen	33
1. Darstellung zur Basis b für natürliche Zahlen	33
2. Darstellung zur Basis b für reelle Zahlen	34
3. Schriftliches Rechnen im Binärsystem	35
Kapitel V. Funktionen – Grundbegriffe	37
1. Graphische Darstellung einfacher Funktionen	38
2. Verknüpfung von Funktionen	39
3. Die Exponentialfunktion	41
4. Die Winkelfunktionen	42
5. Injektivität, Monotonie und Umkehrfunktion	47
6. Polarkoordinaten der komplexen Zahlen	48
Kapitel VI. Stetigkeit	57
1. Folgen	57
2. Folgenstetigkeit	59
3. Eigenschaften stetiger Funktionen	61
4. Verknüpfung stetiger Funktionen	62

5. Nullstellen von Polynomen	64
Kapitel VII. Differenzierbarkeit	71
1. Ableitungsregeln	73
2. l'Hospitalsche Regel	77
3. Monotonie	79
4. Extremwerte	81
5. Differenzierbarkeit der Umkehrfunktion	82
6. Mittelwertsatz	83
Kapitel VIII. Das Riemann-Integral	85
1. Definition des Riemann-Integrals	86
2. Eigenschaften des Integrals	89
3. Uneigentliche Integrale	91
4. Hauptsatz der Differential- und Integralrechnung	92
5. Der Logarithmus und die Exponentialfunktion	96
6. Kettenregel und partielle Integration	100
7. Substitutionsregel	103
8. Partialbruchzerlegung	105
9. Die Gammafunktion	109
Kapitel IX. Ableitungen höherer Ordnung	113
1. Bedeutung der zweiten Ableitung	114
2. Taylorpolynome	117
3. Newtonverfahren	122
Kapitel X. Reihen	125
1. Die geometrische Reihe	125
2. Leibnizkriterium	126
3. Absolute Konvergenz	127
4. Potenzreihen	128
5. Die Hyperbelfunktionen	134
6. Komplexe Potenzreihen	135
Kapitel XI. Lineare Gleichungssysteme	137
1. Die Einsetzmethode	139
2. Das Gaußverfahren	141
3. Rang einer Matrix	148
4. Interpolationspolynome	149
Kapitel XII. Lineare Vektorräume	155
1. Basis, Koordinaten, Dimension	156
2. Norm und Längenmessung	158
3. Skalarprodukt und Winkelmessung	159
4. Untervektorräume	161
5. Lösungsmenge linearer Gleichungssysteme	162
6. Geraden in $\mathbb{R}^2$	163

7. Das Kreuzprodukt auf $\mathbb{R}^3$	165
8. Lösungsmenge einer linearen Differentialgleichung	167
Kapitel XIII. Lineare Abbildungen	171
1. Matrizenmultiplikation	173
2. Die Determinante	174
3. Die Inverse Matrix	178
4. Orthogonale Abbildungen	180
5. Normalformen einer quadratischen Matrix	182
6. Homogene, lineare DGL mit konstanten Koeffizienten	188
Kapitel XIV. Quadratische Formen	193
1. Extremwerte quadratischer Formen auf der Kugel	194
2. Hauptachsentransformation in $\mathbb{R}^2$, Kegelschnitte	196
Kapitel XV. Fourierreihen und Orthonormalsysteme	203
1. Fourierreihen	203
2. Orthonormalsysteme	227

KAPITEL I

Zahlen(bereiche) und mathematische Symbole

1. Die natürlichen Zahlen

1.1. Die Menge der natürlichen Zahlen.

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

Die Menge $\mathbb{N}$ der natürlichen Zahlen ist die kleinste Menge M mit folgenden Eigenschaften:

- i) $0 \in M$
- ii) Es existiert eine bijektive Abbildung $M \rightarrow M \setminus \{0\}$, $m \mapsto m'$.

Für jede natürliche Zahl m heißt m' *Nachfolger* der Zahl m . Wir definieren $1 := 0'$.

BEMERKUNG 1. Man kann eine Menge durch die vollständige Aufzählung ihrer Elemente oder durch die Angabe von Eigenschaften beschreiben oder durch Mengenoperationen (Vereinigung, Durchschnitt oder Differenz anderer Mengen) beschreiben. Die obigen Eigenschaften, die die natürlichen Zahlen definieren, heißen *Peano-Axiome*. $\square$

1.2. Rechenoperationen. Auf $\mathbb{N}$ sind zwei zweistellige Operationen, also Abbildungen $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, definiert:

$$\text{Addition: } (a, b) \mapsto a + b$$

$$\text{Multiplikation: } (a, b) \mapsto ab$$

Diese beiden Rechenoperationen haben folgende grundlegenden Eigenschaften.

SATZ 1 (Rechengesetze). Für beliebige $a, b, c \in \mathbb{N}$ gilt

$a + b = b + a$	<i>Kommutativgesetz der Addition</i>
$ab = ba$	<i>Kommutativgesetz der Multiplikation</i>
$(a + b) + c = a + (b + c)$	<i>Assoziativgesetz der Addition</i>
$(ab)c = a(bc)$	<i>Assoziativgesetz der Multiplikation</i>
$a(b + c) = (ab) + (ac)$	<i>Distributivgesetz</i>
$a + 0 = a$	<i>0 ist neutrales Element der Addition.</i>
$a \cdot 1 = a$	<i>1 := 0' ist neutrales Element der Multiplikation.</i>
$a \cdot 0 = 0$	

BEMERKUNG 2. Die Rechenoperationen Addition und Multiplikation können induktiv oder durch die Mächtigkeit der diskreten Vereinigung und des kartesischen Produktes definiert werden:

$$m + 0 := m, \quad m + n' := (m + n)', \quad m \cdot 0 := 0, \quad m \cdot n' := (m \cdot n) + n \quad \forall m, n \in \mathbb{N}$$

Es seien A, B endliche Mengen mit $A \cap B = \emptyset$.

$$|A| + |B| := |A \cup B|, \quad |A| \cdot |B| := |A \times B| \quad \square$$

1.3. Ordnungsrelation. Man kann natürliche Zahlen miteinander vergleichen. Eine natürliche Zahl a heißt **kleiner als** eine natürliche Zahl b , falls eine natürliche Zahl $c \neq 0$ mit der Eigenschaft $a + c = b$ existiert. Man schreibt dann $a < b$.

Mit Hilfe mathematischer Symbole kann man die Ordnungsrelation „kleiner als“ so definieren:

$$a < b :\Leftrightarrow \exists c \in \mathbb{N} : c \neq 0 \text{ und } a + c = b$$

Die Relation „kleiner als“ ist eine *totale, strenge Ordnungsrelation*. D.h. für beliebige natürliche Zahlen a, b, c gilt

- i) genau eine der Beziehungen $a < b$, $a = b$ oder $b < a$.
- ii) Wenn $a < b$ und $b < c$, dann $a < c$. (Transitivität)

LEMMA 1. Wenn $a, b \in \mathbb{N}$ und $a < b$, dann existiert genau ein $x \in \mathbb{N}$ mit $a + x = b$.

BEWEIS. Da $a < b$, existiert nach Definition der Ordnungsrelation ein $c \in \mathbb{N}$ mit $a + c = b$.

Angenommen die Gleichung $b = a + x$ hat zwei verschiedene Lösungen. Dann muss eine der Lösungen größer sein als die andere. Es existieren also $x_1, x_2 \in \mathbb{N}$ mit $x_1 < x_2$ und $b = a + x_1 = a + x_2$.

Da $x_1 < x_2$, existiert ein $c \in \mathbb{N}$, $c > 0$ mit $x_2 = x_1 + c$. Aus dem Assoziativgesetz der Addition folgt $b = a + x_2 = a + (x_1 + c) = (a + x_1) + c = b + c$. Dies bedeutet $b > b$ und ist ein Widerspruch, da $<$ eine strenge Ordnungsrelation ist. Also ist die Annahme, dass es zwei verschiedene Lösungen der Gleichung $a + x = b$ gibt, falsch. $\square$

BEMERKUNG 3. Zu gegebenen $a, b \in \mathbb{N}$ besitzt die Gleichung $a + x = b$ genau dann eine Lösung $x \in \mathbb{N}$, wenn $a \leq b$. Diese Lösung ist dann eindeutig. Mit Hilfe mathematischer Symbole läßt sich diese Tatsache so ausdrücken:

$$a, b \in \mathbb{N}, a \leq b \Rightarrow \exists! x \in \mathbb{N} : a + x = b \quad \square$$

SATZ 2 (Monotonie der Ordnungsrelation bezüglich der Addition).

$$a < b \Leftrightarrow a + c < b + c \quad \forall a, b, c \in \mathbb{N}$$

SATZ 3 (Monotonie der Ordnungsrelation bezüglich der Multiplikation).

$$a < b \Leftrightarrow ac < bc \quad \forall a, b, c \in \mathbb{N}, c > 0$$

1.4. Teilbarkeit. Eine natürliche Zahl $a \neq 0$ heißt **Teiler** einer natürlichen Zahl b , falls eine natürliche Zahl c mit $ac = b$ existiert. Man sagt dann auch a **teilt** b oder b ist durch a **teilbar** und schreibt $a|b$.

$$a|b :\Leftrightarrow \exists c \in \mathbb{N} : ac = b$$

BEMERKUNG 4. Zu gegebenen $a, b \in \mathbb{N}$ mit $a \neq 0$ besitzt die Gleichung $ax = b$ genau dann eine Lösung $x \in \mathbb{N}$, wenn $a|b$. Ähnlich wie in Lemma 1 kann man zeigen, dass diese Lösung eindeutig ist. Mit Hilfe mathematischer Symbole läßt sich diese Tatsache so ausdrücken:

$$a, b \in \mathbb{N}, a \neq 0, a|b \Rightarrow \exists! x \in \mathbb{N} : ax = b \quad \square$$

2. Die ganzen Zahlen

2.1. Die Menge der ganzen Zahlen. Die Menge der ganzen Zahlen entsteht aus $\mathbb{N}$ durch Hinzunahme der (eindeutigen) Lösungen der Gleichungen $n + x = 0$ für alle $n \in \mathbb{N}, n \neq 0$.

$$\mathbb{Z} := \mathbb{N} \cup \{-n : n \in \mathbb{N}, n \neq 0\}$$

Da die Lösung der Gleichung $n + x = 0$ für alle $n \in \mathbb{N}, n \neq 0$, keine natürliche Zahl ist, enthält der Durchschnitt $\mathbb{N} \cap \{-n : n \in \mathbb{N}, n \neq 0\}$ kein Element, also $\mathbb{N} \cap \{-n : n \in \mathbb{N}, n \neq 0\} = \emptyset$. Die natürlichen Zahlen bilden eine Teilmenge der ganzen Zahlen, $\mathbb{N} \subset \mathbb{Z}$.

2.2. Fortsetzung des Begriffs Teilbarkeit, der Rechenoperationen und der Ordnungsrelation auf $\mathbb{Z}$. Die auf $\mathbb{N}$ definierten Rechenoperationen Addition und Multiplikation sind in eindeutiger Weise zu Rechenoperationen auf $\mathbb{Z}$ fortsetzbar, wenn man fordert, dass die Rechengesetze (Satz 1) für beliebige $a, b, c \in \mathbb{Z}$ gelten und $n + (-n) = 0$ für alle $n \in \mathbb{N}$ gilt.

Die Ordnungsrelation „kleiner als“ kann durch die folgende Definition zu einer strengen, totalen Ordnung auf $\mathbb{Z}$ fortgesetzt werden:

$$-n < -m :\Leftrightarrow n > m \quad \forall n, m \in \mathbb{N}$$

Auch der Begriff der Teilbarkeit ist für ganze Zahlen sinnvoll. Für $a, b \in \mathbb{Z}$ mit $a \neq 0$ ist

$$a|b :\Leftrightarrow \exists c \in \mathbb{Z} : ac = b$$

Falls $a, b \in \mathbb{N} \subset \mathbb{Z}$, so stimmt diese Definition mit der in Abschnitt 1 überein.

2.3. Vor- und Nachteile des Übergangs von $\mathbb{N}$ zu $\mathbb{Z}$. Die Rechenoperationen auf $\mathbb{Z}$ genügen den Rechengesetzen, d.h. Satz 1 gilt für alle $a, b, c \in \mathbb{Z}$. Die Ordnung der ganzen Zahlen ist verträglich mit der Addition, d.h. Satz 2 gilt für alle $a, b, c \in \mathbb{Z}$. Die Gleichung $ax = b$ besitzt für gegebene $a, b \in \mathbb{Z}, a \neq 0$, genau dann eine eindeutige Lösung, wenn $a|b$ (siehe Bemerkung 4).

Für beliebige ganze Zahlen $a, b \in \mathbb{Z}$ existiert immer eine (eindeutige) Lösung der Gleichung $a + x = b$, nämlich $x = b - a$.

Die Ordnungsrelation „kleiner als“ ist nicht mehr monoton bezüglich der Multiplikation. Statt Satz 3 gilt für ganze Zahlen

SATZ 4. Es seien $a, b, c \in \mathbb{Z}$.

Wenn $c > 0$, dann gilt

$$a < b \Leftrightarrow ac < bc,$$

Wenn $c < 0$, dann gilt

$$a < b \Leftrightarrow ac > bc,$$

2.4. Der Betrag. Der **Betrag** $|a|$ einer ganzen Zahl a ist definiert als

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

Der Betrag misst den Abstand der Zahl zu 0 und vergisst das Vorzeichen. Für alle $a \in \mathbb{Z}$ gilt $|a| \geq 0$ und $|a| = 0 \Leftrightarrow a = 0$.

BEMERKUNG 5. Für $a, b \in \mathbb{Z}$ ist $|a - b|$ der Betrag der Differenz der Zahlen a und b , also die Entfernung (der Abstand) von a zu b . $\square$

SATZ 5. Für alle $a, b \in \mathbb{Z}$ gilt

$$a^2 < b^2 \Leftrightarrow |a| < |b|.$$

BEWEIS.

$$a^2 < b^2 \Leftrightarrow 0 < b^2 - a^2 \Leftrightarrow 0 < (b - a)(b + a)$$

Die beiden Faktoren $b - a$ und $b + a$ sind ganze Zahlen. Ihr Produkt ist genau dann größer als 0, wenn entweder beide Faktoren größer als Null sind oder beide Faktoren kleiner als Null sind.

Falls beide Faktoren größer als 0 sind, erhält man

$$b - a > 0 \text{ und } b + a > 0 \Leftrightarrow b > a \text{ und } b > -a \Leftrightarrow b > |a| \geq 0.$$

Falls beide Faktoren kleiner als 0 sind, erhält man

$$b - a < 0 \text{ und } b + a < 0 \Leftrightarrow -b > a \text{ und } -b > -a \Leftrightarrow -b > |a| \geq 0.$$

Zusammengefasst ergibt sich

$$a^2 < b^2 \Leftrightarrow b > |a| \text{ oder } -b > |a| \Leftrightarrow |b| > |a|.$$

$\square$

2.5. Lösung von Ungleichungen.

BEISPIEL 1. Gesucht sind alle $x \in \mathbb{Z}$ mit der Eigenschaft $|x - 5| < 3$, also die Lösungsmenge

$$\mathbb{L} := \{x \in \mathbb{Z} : |x - 5| < 3\}.$$

Um die Ungleichung zu lösen, muss man die Betragszeichen auflösen. Dazu ist eine Fallunterscheidung notwendig.

$x \geq 5$: $|x - 5| = x - 5$ und die zu lösende Ungleichung ist äquivalent zu $x - 5 < 3$, also $x < 8$. $\mathbb{L}_1 = \{x \in \mathbb{Z} : x \geq 5 \wedge x < 8\}$

$x < 5$: $|x - 5| = -(x - 5)$ und die zu lösende Ungleichung ist äquivalent zu $-x + 5 < 3$, also $2 < x$. $\mathbb{L}_2 = \{x \in \mathbb{Z} : x < 5 \wedge x > 2\}$

Zusammengefasst erhält man

$$\mathbb{L} = \mathbb{L}_1 \cup \mathbb{L}_2 = \{5, 6, 7\} \cup \{3, 4\} = \{3, 4, 5, 6, 7\}.$$

Man kann diese Ungleichung auch geometrisch lösen, denn gesucht sind alle ganzen Zahlen, deren Abstand zu 5 kleiner als 3, also 0, 1 oder 2, ist. Dies sind die Zahlen 5 (mit Abstand 0), 4 und 6 (mit Abstand 1) und 3 und 7 (mit Abstand 2).

BEISPIEL 2. Gesucht sind die ganzzahligen Lösungen der Ungleichung

$$4n^2 + 8n + 4 > 16.$$

Da $4n^2 + 8n + 4 = 4(n^2 + 2n + 1) = 2^2(n+1)^2$, folgt aus Satz 5

$$4n^2 + 8n + 4 > 16 \Leftrightarrow |2(n+1)| > |4| \Leftrightarrow 2|n+1| > 4 \Leftrightarrow |n+1| > 2.$$

Wegen $|n+1| = |n - (-1)|$ sind die Lösungen der Ungleichung also alle ganzen Zahlen, deren Abstand zu -1 größer als 2 ist. Die Lösungsmenge ist $\{n \in \mathbb{Z} : n < -3 \text{ oder } n > 1\}$.

BEISPIEL 3. Gesucht sind die ganzzahligen Lösungen der Ungleichung

$$(n+2)(n+4) > 0.$$

Bei dieser Ungleichung kann man die Lösungen auf zwei verschiedenen Wegen bestimmen:

- Für alle $n \in \mathbb{Z}$ gilt $n+4 > n+2$. Das Produkt $(n+2)(n+4)$ ist genau dann größer als 0, wenn beide Faktoren größer als 0 oder beide Faktoren kleiner als 0 sind. Es gilt

$$n+4 > 0 \text{ und } n+2 > 0 \Leftrightarrow n+2 > 0 \Leftrightarrow n > -2$$

und

$$n+4 < 0 \text{ und } n+2 < 0 \Leftrightarrow n+4 < 0 \Leftrightarrow n < -4.$$

Die Lösungsmenge der Ungleichung ist also

$$\{n \in \mathbb{Z} : n > -2 \text{ oder } n < -4\} = \mathbb{Z} \setminus \{-2, -3, -4\}.$$

- Es gilt $(n+2)(n+4) = n^2 + 6n + 8 = n^2 + 2 \cdot 3n + 3^2 - 1 = (n+3)^2 - 1$. Also gilt

$$(n+2)(n+4) > 0 \Leftrightarrow (n+3)^2 - 1 > 0 \Leftrightarrow (n+3)^2 > 1 \Leftrightarrow |n+3| > 1.$$

Die Lösungen der Ungleichung sind alle ganzen Zahlen, deren Abstand zu -3 größer als 1 ist, also weder Abstand 0, wie für $n = -3$, noch Abstand 1, wie für $n = -4$ und $n = -2$.

BEISPIEL 4. Gesucht sind die ganzzahligen Lösungen der Ungleichung

$$|n-3| + |n-1| < 10.$$

Um die Ungleichung zu lösen, muss man die Betragszeichen auflösen. Dazu ist eine Fallunterscheidung notwendig.

$n \geq 3$: Dann $n-3 \geq 0$ und $n-1 \geq 2 > 0$ und man kann die Betragszeichen weglassen. Die zu lösende Ungleichung ist äquivalent zu

$$n-3+n-1 < 10$$

$$2n < 14$$

$$n < 7.$$

$$\mathbb{L}_1 = \{n \in \mathbb{Z} : n \geq 3 \wedge n < 7\} = \{3, 4, 5, 6\}$$

$3 > n \geq 1$: Dann gilt $|n-3| = -(n-3) = -n+3$ und $|n-1| = n-1$. Die zu lösende Ungleichung ist äquivalent zu

$$-n+3+n-1 < 10$$

$$2 < 10.$$

Dies ist eine wahre Aussage und neben der Annahme $3 > n \geq 1$ keine weitere Bedingung.

$$\mathbb{L}_2 = \{n \in \mathbb{Z} : 1 \geq n > 3\} = \{1, 2\}$$

$1 > n$: Dann gilt $|n-3| = -(n-3) = -n+3$ und $|n-1| = -(n-1) = -n+1$. Die zu lösende Ungleichung ist äquivalent zu

$$-n+3-n+1 < 10$$

$$-2n < 6$$

$$n > -3.$$

$$\mathbb{L}_3 = \{n \in \mathbb{Z} : n < 1 \wedge n > -3\} = \{-2, -1, 0, 1\}$$

Zusammengefasst erhält man die Lösungsmenge

$$\mathbb{L} = \mathbb{L}_1 \cup \mathbb{L}_2 \cup \mathbb{L}_3 = \{6, 5, 4, 3, 2, 1, 0, -1, -2\}.$$

Auch diese Ungleichung hätte man geometrisch lösen können. Gesucht sind alle ganzen Zahlen, für die die Summe des Abstandes zu 3 und des Abstandes zu 1 kleiner als 10 ist.

3. Die rationalen Zahlen

$$\mathbb{Q} := \left\{ \frac{p}{q} : p, q \in \mathbb{Z}, q \neq 0 \right\}_{/\sim}$$

Eine rationale Zahl ist ein Quotient/Bruch ganzer Zahlen: Brüche, die durch erweitern und kürzen ineinander umwandelbar sind, repräsentieren dieselbe rationale Zahl, z.B. $\frac{1}{2} = \frac{3}{6}$. Die Äquivalenz zweier formaler Brüche kann man auch nur mit Hilfe der Rechenoperationen ganzer Zahlen beschreiben:

$$\frac{p_1}{q_1} \sim \frac{p_2}{q_2} :\Leftrightarrow p_1 q_2 = p_2 q_1$$

Die ganzen Zahlen sind eine Teilmenge der rationalen Zahlen

$$\mathbb{Z} = \left\{ \frac{p}{1} : p \in \mathbb{Z} \right\} \subset \mathbb{Q}.$$

Durch

$$\frac{p_1}{q_1} + \frac{p_2}{q_2} := \frac{p_1 q_2 + q_1 p_2}{q_1 q_2} \quad \text{und} \quad \frac{p_1}{q_1} \cdot \frac{p_2}{q_2} := \frac{p_1 p_2}{q_1 q_2}$$

werden die Rechenoperationen auf $\mathbb{Z}$ zu Rechenoperationen auf $\mathbb{Q}$ so fortgesetzt, dass die Rechengesetze (Satz 1) gelten.

Auch die Ordnungsrelation „kleiner als“ kann so zu einer strengen, totalen Ordnung auf $\mathbb{Q}$ fortgesetzt werden, dass die Ordnung mit Addition und Multiplikation verträglich ist, d.h. die Sätze 2 und 4 gelten für alle $a, b, c \in \mathbb{Q}$:

$$\frac{p_1}{q_1} < \frac{p_2}{q_2} := q_1 q_2 (p_2 q_1 - q_2 p_1) > 0$$

Setzt man

$$\left| \frac{p}{q} \right| := \frac{|p|}{|q|},$$

so ist die Betragsdefinition für rationale Zahlen mit der für ganze Zahlen verträglich und auch Satz 5 gilt für alle $a, b \in \mathbb{Q}$.

BEMERKUNG 6. Setzt man in den Definitionen der Addition, Multiplikation und Ordnungsrelation rationaler Zahlen $q_1 = q_2 = 1$, so sieht man leicht, dass man die entsprechenden Definitionen für ganze Zahlen erhält. $\square$

Für beliebige $a, b \in \mathbb{Q}$ existiert immer eine eindeutige rationale Lösung der Gleichung $a + x = b$. Für beliebige $a, b \in \mathbb{Q}$ mit $a \neq 0$ existiert immer eine eindeutige rationale Lösung der Gleichung $ax = b$. Solche Zahlenbereiche, die auch die Rechengesetze erfüllen (Satz 1), heißen **Körper**.

BEMERKUNG 7. Der Begriff der Teilbarkeit ist für rationale Zahlen nicht mehr sinnvoll. $\square$

BEMERKUNG 8. Für alle $a, b \in \mathbb{Q}$ mit $a < b$ besitzt die Menge $\{q \in \mathbb{Q} : a < q < b\}$ unendlich viele Elemente. Das bedeutet, dass sich viele Probleme für rationale Zahlen nicht mehr durch Probieren lösen lassen! $\square$

4. Die reellen Zahlen

Die reellen Zahlen sind die Vervollständigung der rationalen Zahlen, denn die Menge der rationalen Zahlen besitzt „Lücken“.

BEISPIEL 5. Betrachtet man die beiden Mengen

$$M_1 := \{q \in \mathbb{Q} : q < 0 \text{ oder } q^2 < 2\} \quad \text{und} \quad M_2 := \{q > 0 \text{ und } q^2 > 2\},$$

so gilt $\mathbb{Q} = M_1 \cup M_2$, $M_1 \cap M_2 = \emptyset$ und $m_1 < m_2$ für alle $m_1 \in M_1$ und $m_2 \in M_2$. Aber es existiert keine rationale Zahl r mit der Eigenschaft $m_1 \leq r \leq m_2$ für alle $m_1 \in M_1$ und $m_2 \in M_2$. Zwischen den beiden Mengen M_1 und M_2 , deren Vereinigung $\mathbb{Q}$ ist, befindet sich also eine Lücke. Diese wird durch eine reelle Zahl geschlossen, in diesem Fall die Zahl $\sqrt{2}$.

BEISPIEL 6. Für die Zerlegung der rationalen Zahlen $\mathbb{Q} = M_1 \cup M_2$ mit $M_1 := \{q \in \mathbb{Q} : q \leq 2\}$ und $M_2 := \{q \in \mathbb{Q} : q > 2\}$ gilt ebenfalls $M_1 \cap M_2 = \emptyset$ und $m_1 < m_2$ für alle $m_1 \in M_1$ und $m_2 \in M_2$. Hier hat jedoch die rationale Zahl 2 die Eigenschaft $m_1 \leq 2 \leq m_2$ für alle $m_1 \in M_1$ und $m_2 \in M_2$. Die Zahl 2 trennt also die beiden Mengen M_1 und M_2 .

BEMERKUNG 9. Wenn für zwei Mengen M_1 und M_2 gilt $M_1 \cap M_2 = \emptyset$, so heißen sie **disjunkt** und ihre Vereinigung $M := M_1 \cup M_2$ **disjunkte** Vereinigung. $\square$

Ein technisches Hilfsmittel zur Arbeit mit Zerlegungen von $\mathbb{Q}$ wie in den Beispielen 5 und 6 und Lücken ist der Begriff der Cauchyfolge. Eine Folge rationaler Zahlen $\{x_n : n \in \mathbb{N}\} \subset \mathbb{Q}$ heißt **Cauchyfolge**, falls für jede positive, rationale Zahl $q \in \mathbb{Q}$, $q > 0$ eine Index n_0 existiert, so dass für alle $n, m \in \mathbb{N}$ mit $n, m \geq n_0$ gilt $|x_n - x_m| < q$.

BEMERKUNG 10. Wenn eine Folge reeller Zahlen die Eigenschaften einer Cauchyfolge hat, so heißt sie konvergent und besitzt einen Grenzwert. Auch wenn alle Folgenglieder einer Cauchyfolge rationale Zahlen sind, so muss dieser Grenzwert nicht unbedingt eine rationale Zahl sein. $\square$

Eine reelle Zahl ist der Grenzwert einer Cauchyfolge. Die reelle Zahl wird durch eine Cauchyfolge repräsentiert. Es gibt verschiedene Cauchyfolgen, die den gleichen Grenzwert haben und damit dieselbe reelle Zahl darstellen.

BEISPIEL 7. Die konstante Folge $\{y_n : n \in \mathbb{N}\}$ mit $y_n = 0$ für alle n hat den Grenzwert 0. Die Folge $\{y_n : n \in \mathbb{N}\}$ mit $y_n = \frac{1}{n}$ für alle $n > 0$ ist auch eine Cauchyfolge, denn für alle $q > 0$ folgt aus $n, m \geq n_0 + 1$ mit $n_0 := \frac{2}{q}$

$$|x_n - x_m| = |x_n + (-x_m)| \leq |x_n| + |-x_m| = \frac{1}{n} + \frac{1}{m} < \frac{q}{2} + \frac{q}{2} = q.$$

Auch der Grenzwert dieser Cauchyfolge ist 0.

BEMERKUNG 11. Jede rationale Zahl $r \in \mathbb{Q}$ kann man durch eine konstante Folge $\{y_n : n \in \mathbb{N}\}$ mit $y_n = r$ für alle $n \in \mathbb{N}$ darstellen. Konstante Folgen sind Cauchyfolgen, denn $|y_n - y_m| = |r - r| = 0 < q$ für alle $q > 0$ und alle $n, m \in \mathbb{N}$. $\square$

Für beliebige reelle Zahlen x und y , die durch Cauchyfolgen repräsentiert sind, $x = \{x_n\}$ und $y = \{y_n\}$, lassen sich Addition, Subtraktion, Multiplikation und Division einfach definieren:

$$x + y := \{x_n + y_n\}, \quad x - y := \{x_n - y_n\}, \quad xy := \{x_n y_n\}, \quad \frac{x}{y} := \left\{ \frac{x_n}{y_n} \right\}$$

Dabei ist die Division nur für $y \neq 0$ möglich. Die Cauchyfolge $\{y_n\}$ muss und kann dann so gewählt werden, dass $y_n \neq 0$ für alle $n \in \mathbb{N}$.

Die Ordnungsrelation „kleiner als“ auf $\mathbb{Q}$ wird durch

$$\{x_n\} < \{y_n\} \Leftrightarrow \exists n_0 \in \mathbb{N}, r \in \mathbb{Q} : x_n < y_m \forall n, m \geq n_0$$

zu einer totalen, strengen Ordnung auf $\mathbb{R}$ fortgesetzt. Weiterhin gilt $|x| := \{|x_n|\}$.

BEMERKUNG 12. Natürlich muss man zeigen, dass auch die $\{x_n + y_n\}$, $\{x_n - y_n\}$, $\{x_n y_n\}$ und $\{x_n/y_n\}$ wieder Cauchyfolgen sind. $\square$

Die Menge der reellen Zahlen $\mathbb{R}$ ist ein **vollständiger, archimedisch geordneter Körper**.

D.h. $\mathbb{R}$ ist ein Körper (siehe Abschnitt 3), alle Cauchyfolgen reeller Zahlen besitzen reelle Grenzwerte (keine Lücken mehr) und die Rechenoperationen sind mit der Ordnungsrelation verträglich (Sätze 2 und 4). Außerdem existieren zu beliebigen $x, y \in \mathbb{R}$ mit $0 < x < y$ eine natürliche Zahl n mit $nx > y$ und eine rationale Zahl q mit $x < q < y$.

4.1. Approximation von $\sqrt{2}$. Eine positive reelle Zahl x mit der Eigenschaft $x^2 = 2$ heißt $\sqrt{2}$. Wir konstruieren mit Hilfe des Halbierungsverfahrens eine Cauchyfolge, die $\sqrt{2}$ darstellt, und erhalten damit eine beliebig gute Näherung für $\sqrt{2}$.

LEMMA 2. $\sqrt{2} \notin \mathbb{Q}$

BEWEIS. Wir nehmen an, dass $\sqrt{2}$ eine rationale Zahl ist. Dann existieren teilerfremde ganze Zahlen p, q mit der Eigenschaft $\sqrt{2} = p/q$. Es gilt

$$\sqrt{2} = \frac{p}{q} \Leftrightarrow 2 = \frac{p^2}{q^2} \Leftrightarrow 2q^2 = p^2 \Rightarrow 2|p^2 \Rightarrow 2|p,$$

da 2 eine Primzahl ist. Also gilt $p = 2m$ für ein $m \in \mathbb{N}$. Nun folgt

$$2q^2 = (2m)^2 \Leftrightarrow 2q^2 = 4m^2 \Leftrightarrow q^2 = 2m^2 \Rightarrow 2|q^2 \Rightarrow 2|q,$$

da 2 eine Primzahl ist. Dies bedeutet aber $2|p$ und $2|q$. Dies ist ein Widerspruch, da p und q keine gemeinsamen Teiler haben. Also ist die Annahme $\sqrt{2} \in \mathbb{Q}$ falsch. Wir folgern $\sqrt{2} \notin \mathbb{Q}$. $\square$

Wenn $a, b \in \mathbb{R}$ und $a, b \geq 0$, dann gilt $a^2 < b^2 \Leftrightarrow a < b$. Da $1^2 = 1 < 2 = \sqrt{2}^2$ und $2^2 = 4 > 2 = \sqrt{2}^2$, gilt $1 < \sqrt{2} < 2$.

Wir wählen $a_0 = 1$ und $b_0 = 2$ und wissen $a_0 < \sqrt{2} < b_0$. Wir definieren

$$x_n = \frac{a_n + b_n}{2}, \quad (a_{n+1}, b_{n+1}) = \begin{cases} (a_n, x_n) & , \text{ falls } x_n^2 > 2 \\ (x_n, b_n) & , \text{ falls } x_n^2 < 2 \end{cases}.$$

Aus der Definition folgt $a_n \leq a_{n+1} < \sqrt{2} < b_{n+1} \leq b_n$ und $b_n - a_n = 2^{-n}$ für alle $n \in \mathbb{N}$. Die Folgen $(a_n)_{n \in \mathbb{N}}$, $(x_n)_{n \in \mathbb{N}}$ und $(b_n)_{n \in \mathbb{N}}$ sind Cauchyfolgen, die gegen $\sqrt{2}$ konvergieren, denn für alle $m \geq n$ gilt

$$|x_m - x_n|, |a_m - a_n|, |b_n - b_m| \leq b_n - a_n = \frac{1}{2^n},$$

$a_n^2 < 2 < b_n^2$ und

$$b_n^2 - a_n^2 = (b_n - a_n)(b_n + a_n) = b_n^2 - 2 + 2 - a_n^2 \leq \frac{3}{2^n}.$$

BEMERKUNG 13. Die konstruierte Folge konvergiert sehr **langsam** gegen $\sqrt{2}$. Später lernen wir bessere Approximationsverfahren, z.B. das Newtonverfahren, kennen. $\square$

4.2. Halbierungsverfahren zur Definition k -ter Wurzeln. Für $k \in \mathbb{N}$, $k \geq 1$ und $a \in \mathbb{Q}$, $a > 0$ ist die **k -te Wurzel** aus a eine reelle Zahl x mit $x > 0$ und $x^k = a$. Sie wird mit $\sqrt[k]{a}$ bezeichnet. Mit Hilfe des Halbierungsverfahrens kann man eine Cauchyfolge erzeugen, die $\sqrt[k]{a}$ darstellt. Die Startwerte sind $x_0, x_1 \in \mathbb{Q}$ mit $x_0^k \leq a$ und $x_1^k \geq a$, z.B. $x_0 = 0$ und $x_1 = \max(1, a)$. Die anderen Folgenglieder werden rekursiv definiert:

$$(1) \quad x_{n+1} = \begin{cases} x_n - \frac{x_n^k - a}{k x_n^{k-1}} & , \text{ falls } x_n^k > a \\ x_n + \frac{x_n^k - a}{k x_n^{k-1}} & , \text{ falls } x_n^k < a \\ x_n & , \text{ falls } x_n^k = a \end{cases}$$

4.3. Polynome. Ein **Polynom** in der Variablen x vom Grad d mit reellen Koeffizienten $a_0, \dots, a_d \in \mathbb{R}$, $a_d \neq 0$, ist ein Ausdruck der Form

$$a_d x^d + a_{d-1} x^{d-1} + \dots + a_1 x + a_0.$$

Eine reelle Nullstelle eines Polynoms ist eine Zahl $x \in \mathbb{R}$ mit der Eigenschaft

$$a_d x^d + a_{d-1} x^{d-1} + \dots + a_1 x + a_0 = 0.$$

Besitzen alle Polynome vom Grad $d > 0$ mit reellen Koeffizienten eine reelle Nullstelle?

- Falls $d = 1$ so ist eine Nullstelle gerade eine Lösung der Gleichung $a_1 x + a_0 = 0$. Diese besitzt die eindeutige Lösung $x = -a_0/a_1$.
- Ein Polynom vom Grad 2 besitzt genau dann eine reelle Nullstelle, falls

$$\left(\frac{a_1}{2a_2} \right)^2 \geq \frac{a_0}{a_2}.$$

Zum Beispiel besitzt $x^2 + 1$, also $a_2 = a_0 = 1$ und $a_1 = 0$, keine reelle Nullstelle, da einerseits $0 < 1$ und $x^2 + 1 \geq 1 > 0$ für alle $x \in \mathbb{R}$.

- Mit Hilfe der Stetigkeit und des Verhaltens von Polynomen für $x \rightarrow \pm\infty$ werden wir später zeigen, dass jedes Polynom von ungeradem Grad eine reelle Nullstelle besitzt.
- Für alle $k \in \mathbb{N}$, $k > 1$, und alle $a \in \mathbb{R}$, $a > 0$, besitzt das Polynom $x^k - a$ eine Nullstelle, nämlich $\sqrt[k]{a}$.

5. Die komplexen Zahlen

Die komplexen Zahlen sind eine Körpererweiterung des Körpers $\mathbb{R}$ um die Nullstelle des Polynoms $x^2 + 1$, das keine reellen Nullstellen besitzt.

$$\mathbb{C} := \{(x, y) : x, y \in \mathbb{R}\} = \{x + iy : x, y \in \mathbb{R}\}$$

Die komplexen Zahlen bilden einen 2-dimensionalen reellen Vektorraum. Der Vektor $(1, 0)$ wird mit $1 \in \mathbb{R}$ identifiziert, der Vektor $(0, 1)$ mit einer Lösung

der Gleichung $x^2 + 1$, die wir i nennen. Die Vektoren $(1, 0)$ und $(0, 1)$ bilden eine Basis des reellen Vektorraumes $\mathbb{C}$. Mit dieser Notation ist $\mathbb{C} = \{x + iy : x, y \in \mathbb{R}\}$. Die reellen Zahlen sind eine Teilmenge der komplexen Zahlen, $\mathbb{R} = \{(x, 0) : x \in \mathbb{R}\} \subset \mathbb{C}$.

Für eine komplexe Zahl $z = x + iy$ heißt x der **Realteil** und y der **Imaginärteil** der komplexen Zahl z . Die Rechenoperationen auf $\mathbb{R}$ lassen sich in eindeutiger Weise zu Rechenoperationen auf $\mathbb{C}$ fortsetzen, wenn die Rechengesetze und $i^2 = -1$ gelten sollen. Für alle komplexen Zahlen $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$ gilt

$$\begin{aligned} z_1 + z_2 &= (x_1 + iy_1) + (x_2 + iy_2) = (x_1 + x_2) + i(y_1 + y_2) \\ z_1 - z_2 &= (x_1 + iy_1) - (x_2 + iy_2) = (x_1 - x_2) + i(y_1 - y_2) \\ z_1 z_2 &= (x_1 + iy_1)(x_2 + iy_2) = x_1 x_2 + x_1 iy_2 + iy_1 x_2 + iy_1 iy_2 \\ &= x_1 x_2 + ix_1 y_2 + iy_1 x_2 + i^2 y_1 y_2 = x_1 x_2 - y_1 y_2 + i(x_1 y_2 + y_1 x_2) \\ \frac{z_1}{z_2} &= \frac{x_1 + iy_1}{x_2 + iy_2} = \frac{(x_1 + iy_1)(x_2 - iy_2)}{(x_2 + iy_2)(x_2 - iy_2)} \\ &= \frac{x_1 x_2 + y_1 y_2 + i(-x_1 y_2 + y_1 x_2)}{x_2^2 + y_2^2} = \frac{x_1 x_2 + y_1 y_2}{x_2^2 + y_2^2} + i \frac{-x_1 y_2 + y_1 x_2}{x_2^2 + y_2^2} \end{aligned}$$

BEMERKUNG 14. Addition und Subtraktion komplexer Zahlen erfolgt komponentenweise wie Vektoraddition bzw. -subtraktion. Multiplikation und Division ist in der Vektornotation, d.h. Zerlegung in Real- und Imaginärteil, umständlich. Sie werden in Polarkoordinaten einfacher sein. $\square$

Für jede komplexe Zahl $z = x + iy$ heißen $\bar{z} := x - iy$ die zu z **konjugiert komplexe Zahl** und $|z| := \sqrt{x^2 + y^2}$ der **Betrag** der komplexen Zahl. Der Betrag misst den Abstand der komplexen Zahl zum Nullpunkt. Die konjugiert komplexe Zahl $\bar{z}$ entsteht aus z durch Spiegelung an der reellen Achse. Es gilt

$$(2) \quad z\bar{z} = (x + iy)(x - iy) = x^2 + y^2 = |z|^2$$

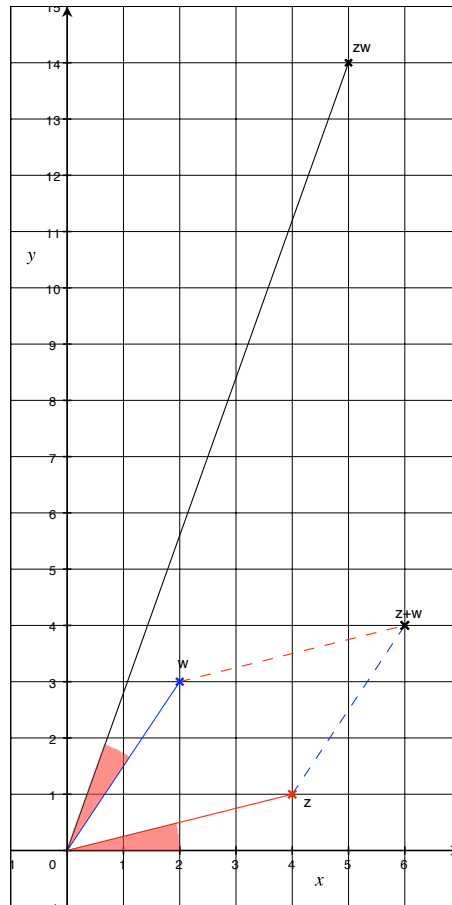
und

$$(3) \quad \frac{1}{z} = \frac{\bar{z}}{z\bar{z}} = \frac{\bar{z}}{|z|^2}.$$

Addition und Multiplikation sind mit dem Übergang zur konjugiert komplexen Zahl vertauschbar:

SATZ 6. Für alle $z_1, z_2 \in \mathbb{C}$ mit $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$ gilt

$$(4) \quad \overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2 \quad \text{und} \quad \overline{z_1 \cdot z_2} = \bar{z}_1 \cdot \bar{z}_2.$$

ABBILDUNG 1. Graphische Darstellung von z , w , $z + w$, zw

BEWEIS.

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

□

BEISPIEL 8. Für $z = 4 + i$ und $w = 2 + 3i$ gilt $z + w = 4 + i + 2 + 3i = 6 + 4i$ und $zw = (4 + i)(2 + 3i) = 4 \cdot 2 - 1 \cdot 3 + i(1 \cdot 2 + 4 \cdot 3) = 5 + 14i$. Abbildung 1 zeigt die graphische Darstellung der komplexen Zahlen z , w , $z + w$ und zw als Punkte bzw. Vektoren in der Ebene.

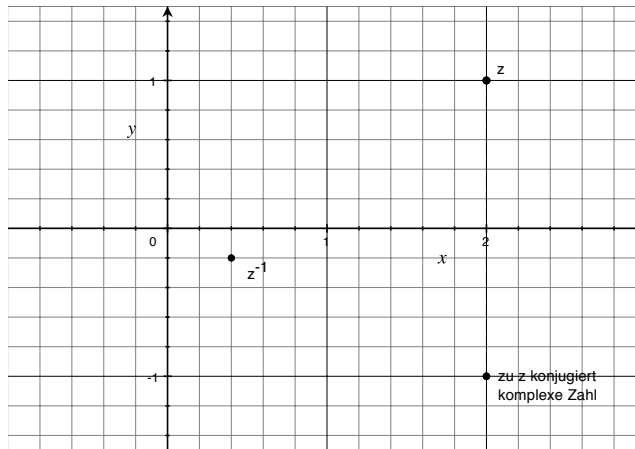


ABBILDUNG 2. Graphische Darstellung von z , $\bar{z}$ und $\frac{1}{z}$

BEISPIEL 9. Für $z = 2 + i$ gilt $|z| = \sqrt{2^2 + 1^2} = \sqrt{5}$, $\bar{z} = 2 - i$ und

$$\frac{1}{z} = \frac{2 - i}{\sqrt{5}^2} = \frac{2}{5} - i \frac{1}{5}.$$

Abbildung 2 zeigt die komplexen Zahlen z , $\bar{z}$ und $1/z$ in der komplexen Zahlenebene.

BEISPIEL 10. Gesucht sind alle $z \in \mathbb{C}$ mit $z^2 - 2z + 5 = 0$.

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

$$z^2 - 2z + 1 + 4 = 0$$

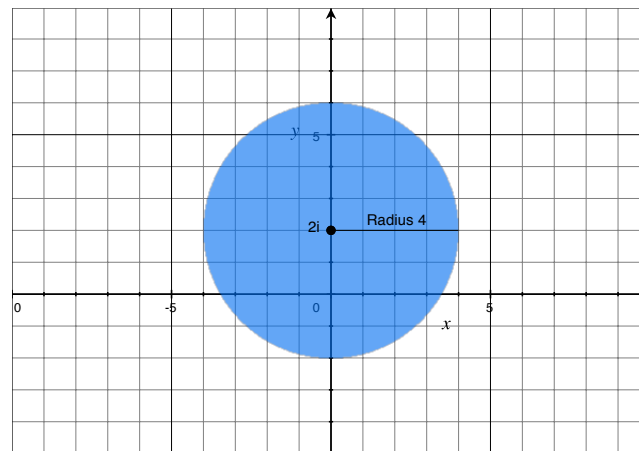
$$(z - 1)^2 = -4 = -1 \cdot 2^2$$

$$z - 1 = \pm i \cdot 2$$

$$z = 1 + \pm 2i$$

Es gibt also genau zwei Lösungen der Gleichung, $z = 1 + 2i$ und $z = 1 - 2i$.

BEISPIEL 11. Die Lösungen der Ungleichung $|z - 2i| < 4$ sind alle komplexen Zahlen, deren Abstand zu $2i$ kleiner als 4 ist, also alle Punkte (x, y) die innerhalb eines Kreises mit Radius 4 um $(0, 2)$ liegen (siehe Abbildung 3).

ABBILDUNG 3. Lösungen der Ungleichung $|z - 2i| < 4$

6. Symbole

Symbol	Bezeichnung oder Sprechweise	Bemerkung
$\{ \}$	Mengenklammern	Begrenzung der Aufzählung der Elemente einer Menge
$\in$	ist Element der Menge	
$M \times N$	Produkt der Mengen M und N $M \times N = \{(m, n) : m \in M, n \in N\}$	
(m, n)	Paar	Elemente der Menge $M \times N$
$:$	mit der Eigenschaft	
$\rightarrow$	Abbildungspfeil (für Mengen)	
$\mapsto$	Abbildungspfeil (für Elemente)	
$\exists$	es existiert ein Element	
$\forall$	für alle, für beliebige	
$:\Leftrightarrow$	genau dann, wenn	zur Definition von Begriffen
$\Leftrightarrow$	genau dann, wenn	in Aussagen
$\Rightarrow$	daraus folgt	
$\exists!$	es existiert genau ein Element	
$:=$	ist gleich, ist definiert durch	zur Definition von Objekten
$\cup$	Vereinigung von Mengen	
$\cap$	Durchschnitt von Mengen	
$\emptyset$	leere Menge	
$\subset$	Teilmenge	
$\max$	Maximum	z.B. $\max(2, 3, -4) = 3$
$\min$	Minimum	z.B. $\min(2, 3, -4) = -4$
$\sum$	Summenzeichen	z.B. $\sum_{k=2}^4 2k = 4 + 6 + 8$

KAPITEL II

Formeln, Gleichungen und Ungleichungen

1. Summenformel der geometrischen Reihe

SATZ 7. Für alle $n \in \mathbb{N}$ und alle $q \in \mathbb{C}$ gilt

$$(1 - q)(1 + q + q^2 + \dots + q^n) = 1 - q^{n+1}.$$

BEWEIS. Man kann diese Aussage direkt oder mit vollständiger Induktion beweisen.

- Direkter Beweis:

$$\begin{aligned}(1 - q)(1 + q + q^2 + \dots + q^n) &= (1 - q) \sum_{k=0}^n q^k = \sum_{k=0}^n q^k - q \sum_{k=0}^n q^k \\ &= \sum_{k=0}^n q^k - \sum_{k=0}^n q^{k+1} = \sum_{k=0}^n q^k - \sum_{l=1}^{n+1} q^l \\ &= q^0 + \sum_{k=1}^n q^k - \sum_{l=1}^n q^l - q^{n+1} = 1 - q^{n+1}\end{aligned}$$

- Beweis mit vollständiger Induktion: Für $n = 0$ ist die Behauptung wahr, denn $(1 - q)(q^0) = (1 - q) \cdot 1 = 1 - q = 1 - q^1$. Nun gilt

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

□

FOLGERUNG 1. Für alle $n \in \mathbb{N}$ und alle $q \in \mathbb{C}$ mit $q \neq 1$ gilt

$$(5) \quad \sum_{k=0}^n q^k = \frac{1 - q^{n+1}}{1 - q}.$$

FOLGERUNG 2. Für alle $n, m \in \mathbb{N}$ mit $m \leq n$ und alle $q \in \mathbb{C}$ mit $q \neq 1$ gilt

$$\sum_{k=m}^n q^k = \frac{q^m - q^{n+1}}{1 - q}.$$

BEWEIS.

$$\sum_{k=m}^n q^k = \sum_{k=0}^n q^k - \sum_{k=0}^{m-1} q^k = \frac{1-q^{n+1}}{1-q} - \frac{1-q^m}{1-q} = \frac{q^m - q^{n+1}}{1-q}.$$

□

BEISPIEL 12. Für $q = 1/2$ erhält man für alle $m \leq n$

$$\sum_{k=m}^n \frac{1}{2^k} = \frac{\frac{1}{2^m} - \frac{1}{2^{n+1}}}{1 - \frac{1}{2}} = 2 \left(\frac{1}{2^m} - \frac{1}{2^{n+1}} \right) < \frac{1}{2^{m-1}}.$$

2. Gaußsche Summenformel

SATZ 8. Für alle $n \in \mathbb{N}$ gilt

$$(6) \quad 1 + 2 + 3 + \dots + n = \frac{1}{2}n(n+1).$$

BEWEIS. Auch diese Aussage kann man direkt oder mit vollständiger Induktion beweisen.

- Beweis mit vollständiger Induktion: Für $n = 0$ ist die Behauptung wahr, denn $0 = \frac{1}{2} \cdot 0 \cdot (0+1)$. Für $n = 1$ ist die Behauptung auch wahr, denn $1 = \frac{1}{2} \cdot 1 \cdot (1+1)$. Nun gilt

$$\begin{aligned} \sum_{k=1}^{n+1} k &= \sum_{k=1}^n k + (n+1) = \frac{1}{2}n(n+1) + (n+1) = (n+1) \left(\frac{n}{2} + 1 \right) \\ &= \frac{1}{2}(n+1)(n+2) = \frac{1}{2}(n+1)((n+1)+1) \end{aligned}$$

- Direkter Beweis:

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

□

3. Dreiecksungleichung

SATZ 9 (Dreiecksungleichung für reelle Zahlen). Für alle $x, y \in \mathbb{R}$ gilt

$$(7) \quad |x+y| \leq |x| + |y|.$$

BEWEIS. Wenn $x+y \geq 0$, dann $|x+y| = x+y \leq |x|+|y|$, da $x \leq |x|$ und $y \leq |y|$. Wenn $x+y < 0$, dann $|x+y| = -(x+y) = -x+(-y) \leq |x|+|y|$, da $-x \leq |x|$ und $-y \leq |y|$. □

BEMERKUNG 15. In der Dreiecksungleichung für reelle Zahlen gilt $|x + y| = |x| + |y|$ genau dann, wenn x und y das gleiche Vorzeichen haben. $\square$

SATZ 10 (Dreiecksungleichung für komplexe Zahlen). *Für alle $z_1, z_2 \in \mathbb{C}$ gilt*

$$(8) \quad |z_1 + z_2| \leq |z_1| + |z_2|.$$

BEWEIS. Es seien $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$ mit $x_1, x_2, y_1, y_2 \in \mathbb{R}$. Aus der Schwarzschen Ungleichung (Satz 13) folgt mit $a_1 = x_1, a_2 = y_1, b_1 = x_2$ und $b_2 = y_2$

$$\begin{aligned} (x_1x_2 + y_1y_2)^2 &\leq (x_1^2 + y_1^2)(x_2^2 + y_2^2) \\ 2(x_1x_2 + y_1y_2) &\leq 2\sqrt{x_1^2 + y_1^2}\sqrt{x_2^2 + y_2^2} \\ x_1^2 + y_1^2 + 2(x_1x_2 + y_1y_2) + x_2^2 + y_2^2 &\leq x_1^2 + y_1^2 + 2\sqrt{x_1^2 + y_1^2}\sqrt{x_2^2 + y_2^2} + x_2^2 + y_2^2 \\ (x_1 + x_2)^2 + (y_1 + y_2)^2 &\leq \left(\sqrt{x_1^2 + y_1^2} + \sqrt{x_2^2 + y_2^2}\right)^2 \\ |z_1 + z_2|^2 &\leq (|z_1| + |z_2|)^2 \\ |z_1 + z_2| &\leq |z_1| + |z_2|. \end{aligned}$$

$\square$

BEMERKUNG 16. In der Dreiecksungleichung für komplexe Zahlen gilt $|z_1 + z_2| = |z_1| + |z_2|$ genau dann, wenn die Vektoren z_1 und z_2 in die gleiche Richtung zeigen, also $z_1 = \lambda z_2$ oder $z_2 = \lambda z_1$ für ein $\lambda \in \mathbb{R}$ mit $\lambda \geq 0$. $\square$

FOLGERUNG 3 (Verallgemeinerte Dreiecksungleichung). *Für alle $z_1, \dots, z_n \in \mathbb{C}$ gilt*

$$|z_1 + \dots + z_n| \leq |z_1| + \dots + |z_n|.$$

4. Arithmetisches und geometrisches Mittel

SATZ 11 (Ungleichung zwischen arithmetischem und geometrischem Mittel). *Für alle $a, b \in \mathbb{R}$ mit $a, b \geq 0$ gilt*

$$(9) \quad \frac{a+b}{2} \geq \sqrt{ab}.$$

BEWEIS. Für alle $a, b \in \mathbb{R}$ gilt

$$\begin{aligned} (a-b)^2 &\geq 0 \\ a^2 - 2ab + b^2 &\geq 0 \\ a^2 + 2ab + b^2 &\geq 4ab \\ (a+b)^2 &\geq 4ab. \end{aligned}$$

Aus $a, b \geq 0$ folgt $4ab \geq 0$ und $a+b \geq 0$, also $a+b = |a+b| \geq \sqrt{4ab} = 2\sqrt{ab}$. $\square$

Für zwei positive, reelle Zahlen a, b heißt $\frac{1}{2}(a+b)$ ihr **arithmetisches Mittel** oder Mittelwert und $\sqrt{ab}$ ihr **geometrisches Mittel**. Das geometrische Mittel ist die Seitenlänge des Quadrates, das den gleichen Flächeninhalt hat wie das Rechteck mit den Seitenlängen a und b .

Aus dem Beweis des Satzes 11 folgt, dass in der Ungleichung $\frac{1}{2}(a+b) \geq \sqrt{ab}$ genau dann das Gleichheitszeichen gilt, wenn $a = b$.

SATZ 12 (Allgemeine Ungleichung zwischen arithmetischem und geometrischem Mittel). *Für alle $a_1, \dots, a_n \in \mathbb{R}$ mit $a_k \geq 0$ für alle $k = 1, \dots, n$ gilt*

$$(10) \quad \frac{1}{n} \sum_{k=1}^n a_k \geq \sqrt[n]{\prod_{k=1}^n a_k}.$$

BEWEIS. Siehe Beispiel 130. □

5. Schwarzsche Ungleichung

SATZ 13 (Schwarzsche Ungleichung). *Für alle $a_k, b_k \in \mathbb{R}$ mit $k = 1, \dots, n$ gilt*

$$(11) \quad \left(\sum_{k=1}^n a_k b_k \right)^2 \leq \left(\sum_{k=1}^n a_k^2 \right) \left(\sum_{k=1}^n b_k^2 \right).$$

BEWEIS. Für alle $\mu, \lambda \in \mathbb{R}$ gilt

$$\begin{aligned} \sum_{k=1}^n (\lambda a_k - \mu b_k)^2 &\geq 0 \\ \sum_{k=1}^n (\lambda^2 a_k^2 - 2\lambda\mu a_k b_k + \mu^2 b_k^2) &\geq 0 \\ \lambda^2 \sum_{k=1}^n a_k^2 - 2\lambda\mu \sum_{k=1}^n (a_k b_k) + \mu^2 \sum_{k=1}^n b_k^2 &\geq 0 \\ \lambda^2 \sum_{k=1}^n a_k^2 + \mu^2 \sum_{k=1}^n b_k^2 &\geq 2\lambda\mu \sum_{k=1}^n (a_k b_k). \end{aligned}$$

Nun setzt man

$$\lambda := \sqrt{\sum_{k=1}^n b_k^2} \quad \text{und} \quad \mu := \begin{cases} \sqrt{\sum_{k=1}^n a_k^2} & , \text{ falls } \sum_{k=1}^n (a_k b_k) \geq 0 \\ -\sqrt{\sum_{k=1}^n a_k^2} & , \text{ falls } \sum_{k=1}^n (a_k b_k) < 0 \end{cases}.$$

Dann gilt

$$\begin{aligned} \left(\sum_{k=1}^n b_k^2\right) \left(\sum_{k=1}^n a_k^2\right) + \left(\sum_{k=1}^n a_k^2\right) \left(\sum_{k=1}^n b_k^2\right) &\geq 2 \sqrt{\sum_{k=1}^n b_k^2} \sqrt{\sum_{k=1}^n a_k^2} \left| \sum_{k=1}^n (a_k b_k) \right| \\ \sqrt{\sum_{k=1}^n b_k^2} \sqrt{\sum_{k=1}^n a_k^2} &\geq \left| \sum_{k=1}^n (a_k b_k) \right| \geq 0 \\ \left(\sum_{k=1}^n b_k^2\right) \left(\sum_{k=1}^n a_k^2\right) &\geq \left(\sum_{k=1}^n (a_k b_k)\right)^2, \end{aligned}$$

falls $\sum_{k=1}^n a_k^2 \neq 0$ und $\sum_{k=1}^n b_k^2 \neq 0$.

Es gilt $\sum_{k=1}^n a_k^2 = 0$ genau dann, wenn $a_k = 0$ für alle k , und $\sum_{k=1}^n b_k^2 = 0$ genau dann, wenn $b_k = 0$ für alle k . In beiden Fällen lautet die zu beweisende Aussage $0 \leq 0$ und ist wahr. $\square$

BEMERKUNG 17. Mit Hilfe der Begriffe Skalarprodukt und Betrag von Vektoren läßt sich die Schwarzsche Ungleichung knapp als $\langle a, b \rangle^2 \leq \|a\|^2 \|b\|^2$ schreiben. Sie vergleicht also das Skalarprodukt zweier Vektoren mit dem Produkt ihrer Längen. (siehe Abschnitt 3). $\square$

6. Der binomische Satz

Multipliziert man das Produkt $(a + b)^n$ vollständig aus, so erhält man Summanden der Form $c_k a^k b^{n-k}$ für $k = 0, \dots, n$, wobei jeder Koeffizient c_k gerade die Anzahl der Möglichkeiten ist, aus den n Faktoren die k Faktoren auszuwählen, die einen Faktor a zum Produkt $a^k b^{n-k}$ liefern. Zum Beispiel

$$\begin{aligned} (a + b)^2 &= a^2 + 2ab + b^2 \quad (c_0 = c_2 = 1, c_1 = 2) \\ (a + b)^3 &= a^3 + 3a^2b + 3ab^2 + b^3 \quad (c_0 = c_3 = 1, c_1 = c_2 = 3). \end{aligned}$$

Wir wollen eine allgemeine Formel für die Koeffizienten c_k herleiten. Dazu müssen wir zwei Grundprobleme der Kombinatorik lösen.

6.1. Permutationen und Kombinationen. Die Anzahl der verschiedenen Anordnungen von k verschiedenen Elementen wird mit $P(k)$ bezeichnet und Anzahl der **Permutationen** von k Elementen genannt.

$$P(k) = |\{\Phi : \Phi : M \rightarrow M \text{ bijektiv}\}|, \quad |M| = k$$

BEMERKUNG 18. Hat man k paarweise verschiedene Buchstaben zur Verfügung, so ist $P(k)$ gerade die Anzahl der verschiedenen Wörter, die nur diese Buchstaben und jeden dieser Buchstaben genau einmal enthalten. $\square$

Das Produkt der ersten k von 0 verschiedenen natürlichen Zahlen bezeichnet man mit $k!$ und nennt es k **Fakultät**:

$$k! := 1 \cdot 2 \cdot \dots \cdot (k-1) \cdot k = \prod_{m=1}^k m, \quad \forall k \in \mathbb{N}, (0! := 1).$$

SATZ 14. Für alle $k \in \mathbb{N}$ gilt $P(k) = k!$.

BEWEIS. Vollständige Induktion: Für $k = 1$ gilt $P(k) = 1$, denn aus einem Buchstaben kann man nur ein Wort bilden, das diesen Buchstaben genau einmal enthält. Außerdem gilt auch $1! = 1$. Hat man $k + 1$ Buchstaben und will ein Wort aus diesen bilden, so hat man $k + 1$ Möglichkeiten für den ersten Buchstaben des Wortes. Für die restlichen k Plätze im Wort stehen dann noch die k anderen Buchstaben zur Verfügung, also $P(k + 1) = (k + 1)P(k) = (k + 1)k! = (k + 1)!$. $\square$

BEISPIEL 13. Aus den drei Buchstaben I, S und T, also $k = 3$ lassen sich die sechs Wörter ITS, IST, SIT, STI, TIS und TSI bilden und auch $P(3) = 1 \cdot 2 \cdot 3 = 6$.

Die Anzahl k -elementigen Teilmengen einer n -elementigen Menge heißt **Kombinationen** von k Elementen aus einer Gesamtheit von n Elementen und wird mit $C(k, n)$ bezeichnet.

BEMERKUNG 19. Hat man n Personen zur Verfügung, so ist $C(k, n)$ gerade die Anzahl der verschiedenen Mannschaften zu k Spielern, die man aus diesen n Personen bilden kann. $\square$

SATZ 15. Für alle $n, k \in \mathbb{N}$ mit $n \geq k$ gilt

$$(12) \quad C(k, n) = \frac{n(n-1) \dots (n-k+1)}{k!}.$$

BEWEIS. Wählt man die Spieler nacheinander aus, so hat man für die Wahl des m -ten Spielers noch $n - m + 1$ Möglichkeit, insgesamt also $n(n-1) \dots (n-k+1)$ Möglichkeiten. Da es aber auf der Reihenfolge der Auswahl nicht ankommt, muss diese Zahl durch $P(k)$ geteilt werden. $\square$

6.2. Der Binomialkoeffizient.

BEMERKUNG 20. Die Formel für $C(k, n)$ in Satz 15 ist sogar für $k > n$ sinnvoll. Falls $k > n$, so stehen nicht genug Personen zur Bildung einer Mannschaft zur Verfügung. Es gibt also keine Möglichkeit, eine reguläre Mannschaft zu bilden. In der Tat taucht im Zähler der Formel für $C(k, n)$ ein Faktor Null auf, wenn $k > n$.

Für $n \geq k$ gilt

$$\begin{aligned} \frac{n(n-1) \dots (n-k+1)}{k!} &= \frac{\prod_{m=n-k+1}^n m}{k!} = \frac{(\prod_{m=n-k+1}^n m)(\prod_{m=1}^{n-k} m)}{k!(n-k)!} \\ &= \frac{n!}{k!(n-k)!}. \end{aligned}$$

Die Größe $C(k, n)$ wird auch **Binomialkoeffizient** genannt und mit $\binom{n}{k}$ bezeichnet:

$$\binom{n}{k} := \frac{n(n-1) \dots (n-k+1)}{k!} = \begin{cases} \frac{n!}{k!(n-k)!} & , \text{ falls } n \geq k \\ 0 & , \text{ falls } n < k \end{cases} \quad \forall n, k \in \mathbb{N}$$

SATZ 16 (Binomischer Satz). Für alle $a, b \in \mathbb{R}$ und $n \in \mathbb{N}$, $n > 0$, gilt

$$(13) \quad (a+b)^n = \sum_{k=0}^n \binom{n}{k} a^k b^{n-k}.$$

BEWEIS.

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

Man könnte diesen Satz auch mit vollständiger Induktion beweisen. $\square$

BEMERKUNG 21. Der Begriff des Binomialkoeffizienten kann sinnvoll für $n \in \mathbb{R}$ verallgemeinert werden, so dass eine entsprechende binomische Formel auch für reelle Exponenten gilt. $\square$

BEISPIEL 14. Es gilt $\binom{3}{1} = \frac{3!}{1!} = 3$.

LEMMA 3. Für alle $n \in \mathbb{N}$ gilt

$$(14) \quad \binom{n}{0} = 1 = \binom{n}{n} \text{ und } \binom{n}{1} = n = \binom{n}{n-1} = n, \text{ falls } n > 0.$$

BEWEIS.

$$\binom{n}{0} = \frac{n!}{0!(n-0)!} = 1, \quad \binom{n}{n} = \frac{n!}{n!(n-n)!} = 1$$

und falls $n > 0$

$$\binom{n}{1} = \frac{n!}{1!(n-1)!} = n, \quad \binom{n}{n-1} = \frac{n!}{(n-1)!(n-(n-1))!} = n$$

$\square$

LEMMA 4. Für alle $n, k \in \mathbb{N}$ mit $n \geq k$ gilt

$$(15) \quad \binom{n}{k} = \binom{n}{n-k}.$$

BEWEIS.

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

$\square$

BEMERKUNG 22. Lemma 4 kann man auch anschaulich beweisen, denn die Anzahl der Möglichkeiten aus n Personen k für eine Mannschaft auszuwählen ist gleich der Anzahl der Möglichkeiten aus n Personen $n-k$ auszuwählen, die nicht mitspielen.

6.3. Rekursionsformel.

SATZ 17 (Rekursionsformel für Binomialkoeffizienten). *Für alle $n, k \in \mathbb{N}$ mit $k > 0$ gilt*

$$(16) \quad \binom{n+1}{k} = \binom{n}{k} + \binom{n}{k-1}.$$

BEWEIS. Falls $k > n+1$, so sind alle auftretenden Terme Null. Falls $k = n+1$, so gilt $\binom{n+1}{k} = \binom{n}{k-1} = 1$ und $\binom{n}{k} = 0$. Falls $k \leq n$, so gilt

$$\begin{aligned} \binom{n}{k} + \binom{n}{k-1} &= \frac{n!}{k!(n-k)!} + \frac{n!}{(k-1)!(n-(k-1))!} \\ &= \frac{n!}{k!(n-k)!} + \frac{n!}{(k-1)!(n-k+1)!} \\ &= \frac{n!}{k!(n-k+1)!}(n-k+1+k) = \frac{(n+1)!}{k!(n+1-k)!} = \binom{n+1}{k}. \end{aligned}$$

□

Mit Hilfe der Rekursionsformel wird das Pascalsche Dreieck aufgebaut:

$$\begin{array}{lcl} n=0: & \binom{0}{0} & 1 \\ n=1: & \binom{1}{0} + \binom{1}{1} & 1 + 1 \\ n=2: & \binom{2}{0} + \binom{2}{1} + \binom{2}{2} & 1 + 2 + 1 \\ n=3: & \binom{3}{0} + \binom{3}{1} + \binom{3}{2} + \binom{3}{3} & 1 + 3 + 3 + 1 \\ n=4: & \binom{4}{0} + \binom{4}{1} + \binom{4}{2} + \binom{4}{3} + \binom{4}{4} & 1 + 4 + 6 + 4 + 1 \\ n=5: & \binom{5}{0} + \binom{5}{1} + \binom{5}{2} + \binom{5}{3} + \binom{5}{4} + \binom{5}{5} & 1 + 5 + 10 + 10 + 5 + 1 \\ & \vdots & \end{array}$$

BEMERKUNG 23. Um für ein festes $n \in \mathbb{N}$ alle Binomialkoeffizienten $\binom{n}{k}$ zu berechnen benötigt man

- mit der Formel $\binom{n}{k} = \frac{n(n-1)\dots(n-k+1)}{k!}$ für jedes $k \leq n$ genau $2(k-1) + 1$ Multiplikationen, also insgesamt ($k=0, n$ werden weggelassen)

$$\sum_{k=1}^{n-1} (2k+1) = 2\left(\sum_{k=1}^{n-1} k\right) + (n-1) = (n-1)(n+1)$$

Multiplikationen.

- mit der Rekursionsformel

$$\sum_{k=1}^n (k-1) = \sum_{k=1}^{n-1} k = \frac{1}{2}(n-1)n$$

Additionen.

Da eine Multiplikation auch aufwendiger ist als eine Addition, ist die Berechnung der Binomialkoeffizienten mit Hilfe der Rekursionsformel effektiver.

□

6.4. Summenformeln für Binomialkoeffizienten. Setzt man in die binomische Formel spezielle Werte für a und b ein, so erhält man Summenformeln für Binomialkoeffizienten.

FOLGERUNG 4. Für alle $n \in \mathbb{N}$ gilt

$$(17) \quad \sum_{k=0}^n \binom{n}{k} = 2^n.$$

BEWEIS. Für $a = b = 1$ folgt aus der binomischen Formel (Satz 16)

$$2^n = (1 + 1)^n = \sum_{k=0}^n \binom{n}{k} 1^k 1^{n-k} = \sum_{k=0}^n \binom{n}{k}.$$

□

FOLGERUNG 5. Für alle $n \in \mathbb{N}$, $n > 0$, gilt

$$(18) \quad \sum_{k=0}^{n/2} \binom{n}{2k} = \sum_{k=0}^{(n-1)/2} \binom{n}{2k+1} = 2^{n-1}.$$

BEWEIS. Für $a = 1$, $b = -1$ und $n \in \mathbb{N}$ mit $n > 0$ folgt aus der binomischen Formel (Satz 16)

$$\begin{aligned} (1 - 1)^n &= \sum_{k=0}^n \binom{n}{k} 1^k (-1)^{n-k} \\ 0^n &= \sum_{2|k} \binom{n}{k} + \sum_{2 \nmid k} \binom{n}{k} (-1) \\ 0 &= \sum_{l=0}^{n/2} \binom{n}{2l} - \sum_{l=0}^{(n-1)/2} \binom{n}{2l+1} \end{aligned}$$

und

$$\sum_{k=0}^{n/2} \binom{n}{2k} = \frac{1}{2} \sum_{k=0}^n \binom{n}{k} = \frac{1}{2} 2^n = 2^{n-1}.$$

□

7. Bernoullische Ungleichung

SATZ 18 (Bernoullische Ungleichung). Für alle $n \in \mathbb{N}$ und $x \in \mathbb{R}$ mit $x > -1$ gilt

$$(19) \quad (1 + x)^n \geq 1 + nx.$$

BEWEIS. Wir beweisen diese Aussage mit vollständiger Induktion.

$n = 0$: $(1 + x)^n = (1 + x)^0 = 1$ und $1 + nx = 1 + 0 \cdot x = 1$ also $(1 + x)^n \geq 1 + nx$.

$n = 1$: $(1 + x)^n = (1 + x)^1 = 1 + x$ und $1 + nx = 1 + 1 \cdot x = 1 + x$ also $(1 + x)^n \geq 1 + nx$.

$$(1+x)^{n+1} = (1+x)^n(1+x) \geq (1+nx)(1+x) = 1+x+nx+nx^2 \geq 1+x(n+1),$$

da $1+x > 0$ und $nx^2 > 0$ für alle x, n .

9

Wir beobachten n^k , a^n und $n!$ für festes $k \in \mathbb{N}$, $k \neq 0$, $a \in \mathbb{R}$, $a > 1$ und für beliebige $n \in \mathbb{N}$.

Eine Folge $\{x_n\}_{n \in \mathbb{N}}$ heißt **monoton wachsend**, falls $x_{n+1} \geq x_n$ für alle $n \in \mathbb{N}$. Sie heißt **streng** monoton wachsend, falls $x_{n+1} > x_n$ für alle $n \in \mathbb{N}$. Eine Folge $\{x_n\}_{n \in \mathbb{N}}$ heißt **monoton fallend**, falls $x_{n+1} \leq x_n$ für alle $n \in \mathbb{N}$. Sie heißt **streng** monoton fallend, falls $x_{n+1} < x_n$ für alle $n \in \mathbb{N}$.

BEWEIS. Aus $n+1 > n$ folgt $(n+1)^k > n^k$ für alle $n \in \mathbb{N}$. Wegen $a > 1$ gilt auch $a^{n+1} = a^n a > a^n$ für alle $n \in \mathbb{N}$. Es gilt $0! = 1! = 1$. Für $n > 0$ gilt $(n+1)! = (n+1)n! > n!$. $\square$

Eine Folge $\{x_n\}_{n \in \mathbb{N}}$ heißt **nach oben beschränkt**, falls ein $S \in \mathbb{R}$ mit $x_n \leq S$ für alle $n \in \mathbb{N}$ existiert. Die Zahl S ist eine **obere Schranke** der Folge. Eine Folge $\{x_n\}_{n \in \mathbb{N}}$ heißt **nach unten beschränkt**, falls ein $S \in \mathbb{R}$ mit $x_n \geq S$ für alle $n \in \mathbb{N}$ existiert. Die Zahl S ist eine **untere Schranke** der Folge.

BEWEIS. Wir zeigen, dass die Ungleichungen $n^k > S$, $a^n > S$ und $n! > S$ für beliebige $S \in \mathbb{R}$ Lösungen besitzen.

Es gilt genau dann $n^k > S$, wenn $S < 0$ oder $n > \sqrt[k]{S}$. Es gibt natürliche Zahlen, die größer als $\sqrt[k]{S}$ sind (Archimedisches Axiom).

Da $a > 1$ ist, gilt $a = 1 + x$ für ein $x > 0$. Aus der Bernoullischen Ungleichung (Satz 18) folgt $a^n = (1+x)^n \geq 1+nx$. Es gilt $1+nx > S$ genau dann, wenn $n > (S-1)/x$. Für alle $n > (S-1)/x$ folgt $a^n \geq 1+nx > S$.

Es gilt $0! = 1! = 1$. Für $n \geq 2$ gilt $n! \geq 2^{n-1}$. Da die Folge $\{2^n\}$ keine obere Schranke besitzt, ist auch die Folge $\{n!\}$ nicht nach oben beschränkt.

□

BEISPIEL 15. Für welche $n \in \mathbb{N}$ gilt $2^n < n^3$? Wir testen, ob diese Ungleichung für kleine n gilt und formulieren und beweisen dann eine allgemeine Aussage.

[illegible]

Mit vollständiger Induktion folgt $2^n > n^3$ für alle $n \geq 10$, denn für $n \geq 10$ gilt $2^{n+1} = 2 \cdot 2^n > 2n^3$ (Induktionsvoraussetzung) und, wegen $n/3 > 3$,

$$(n+1)^3 = n^3 + 3n^2 + 3n + 1 < n^3 + \frac{n}{3}n^2 + \frac{n^2}{3}n + \frac{n^3}{3}1 = 2n^3.$$

Die Ungleichung $2^n > n^3$ gilt also für $n = 0, 1$ und für alle $n \geq 10$.

SATZ 19. Für alle $a \in \mathbb{R}$ mit $a > 1$ existiert ein $n_0 \in \mathbb{N}$, so dass $a^n > n$ $\forall n \geq n_0$.

BEWEIS. Da $a > 1$ ist, gilt $a = 1 + x$ für ein $x > 0$. Für $n \geq 2$ folgt aus der binomischen Formel (Satz 16)

$$a^n = (1+x)^n \geq 1 + nx + \frac{1}{2}n(n-1)x^2 > nx + \frac{1}{2}n(n-1)x^2.$$

Für $n > 0$ gilt

$$\begin{aligned} nx + \frac{1}{2}n(n-1)x^2 > n &\Leftrightarrow x + \frac{1}{2}(n-1)x^2 > 1 \\ &\Leftrightarrow 2x + (n-1)x^2 > 2 \Leftrightarrow n > \frac{2-2x+x^2}{x^2}. \end{aligned}$$

Nun sei n_0 die kleinste natürliche Zahl, die größer als $(2-2x+x^2)/x^2$ ist. Dann gilt $a^n > n$ für alle $n \geq n_0$. $\square$

FOLGERUNG 6. Für alle $a \in \mathbb{R}$ mit $a > 1$ und alle $k \in \mathbb{N}$ existiert ein $n_0 \in \mathbb{N}$, so dass $a^n > n^k$ $\forall n \geq n_0$.

BEWEIS. $a^n > n^k$ genau dann, wenn $\sqrt[k]{a^n} > n$. Wenn $a > 1$, dann $\sqrt[k]{a} > 1$. $\square$

FOLGERUNG 7. Für alle $a \in \mathbb{R}$ mit $a > 1$, alle $k \in \mathbb{N}$ und alle $c \in \mathbb{R}$ existiert ein $n_0 \in \mathbb{N}$, so dass $a^n > cn^k$ $\forall n \geq n_0$.

BEWEIS. Für alle $n \geq c$ gilt $cn^k < n^{k+1}$. Folgerung 6 liefert ein $n_1 \in \mathbb{N}$ mit der Eigenschaft $a^n > n^{k+1}$ für alle $n \geq n_1$. Nun gilt $a^n > n^{k+1} \geq cn^k$ für alle $n \geq \max(c, n_1)$, also $n_0 = \max(c, n_1)$. $\square$

Das folgende Lemma bedeutet, dass das qualitative Wachstumsverhalten von Polynomen durch den führenden Term dominiert wird.

LEMMA 7. Für alle $d \in \mathbb{N}$, $d > 0$, alle $c_k \in \mathbb{R}$ mit $c_d \neq 0$ und alle $\varepsilon > 0$ existiert ein $n_0 \in \mathbb{N}$, so dass

$$(20) \quad (c_d - \varepsilon)n^d < \sum_{k=0}^d c_k n^k < (c_d + \varepsilon)n^d \quad \forall n \geq n_0.$$

BEWEIS. Es gilt

$$\sum_{k=0}^d c_k n^k = n^d \left(c_d + \frac{c_{d-1}}{n} + \dots + \frac{c_1}{n^{d-1}} + \frac{c_0}{n^d} \right).$$

Da $\lim_{n \rightarrow \infty} n^{-k} = 0$ für alle $k \in \mathbb{N}$, $k > 0$, existiert ein Index n_1 , zum Beispiel $n_0 := d \max(|c_k| : k = 0, \dots, d-1)$, mit

$$-\varepsilon < \frac{c_{d-1}}{n} + \dots + \frac{c_1}{n^{d-1}} = \sum_{k=0}^{d-1} \frac{c_k}{n^{d-k}} < \varepsilon \quad \forall n \geq n_0.$$

Nun folgt

$$n^d(c_d - \varepsilon) < n^d \left(c_d + \sum_{k=0}^{d-1} \frac{c_k}{n^{d-k}} \right) < n^d(c_d + \varepsilon) \quad \forall n \geq n_0.$$

□

FOLGERUNG 8. Für jedes Polynom $P(n)$ mit reellen Koeffizienten und alle $a \in \mathbb{R}$ mit $a > 1$ existiert ein $n_0 \in \mathbb{N}$, so dass $a^n > P(n) \quad \forall n \geq n_0$.

BEWEIS. Wenn P ein Polynom vom Grad k ist, dann existieren nach Lemma 7 ein Index $n_1 \in \mathbb{N}$ und eine Konstante $c \in \mathbb{R}$ mit $P(n) < cn^k$ für alle $n \geq n_1$. Nach Folgerung 7 existiert ein Index $n_2 \in \mathbb{N}$ mit $a^n > cn^k \quad \forall n \geq n_0$. Für alle $n \geq \max(n_1, n_2)$ gilt nun $a^n > cn^k > P(n)$, also $n_0 = \max(n_1, n_2)$. □

BEMERKUNG 24. Beachtet man, dass $a^n = e^{n \ln a}$, so besagt Folgerung 8, dass die Exponentialfunktion stärker/schneller wächst als jedes Polynom. Deshalb bevorzugt man Lösungsalgorithmen, deren Laufzeit polynomial und nicht exponentiell von der Länge der Eingangsdaten abhängen. Es ist jedoch möglich, dass dieser „Laufzeitvorteil“ erst für sehr große Eingangsdaten auftritt. So gibt es zum Beispiel für das Problem der linearen, diskreten Optimierung zwei Lösungsalgorithmen, den Simplexalgorithmus und den Ellipsoidalgorithmus. Der Simplexalgorithmus besitzt (im ungünstigsten Fall) exponentielle Laufzeit. Er ist aber relativ leicht zu programmieren. Der Ellipsoidalgorithmus besitzt polynomiale Laufzeit, allerdings mit großen Koeffizienten. Die Programmierung des Ellipsoidalgorithmus ist relativ aufwendig. Obwohl der Ellipsoidalgorithmus qualitativ besser ist als der Simplexalgorithmus, wird in der Praxis oft der Simplexalgorithmus bevorzugt, weil er einfacher umzusetzen ist, für große (aber nicht zu große) Datenmengen schneller ist als der Ellipsoidalgorithmus und der ungünstigste Fall, der zur exponentiellen Laufzeit führt, nur selten eintritt. □

BEISPIEL 16. Für welche $n \in \mathbb{N}$ gilt $n! > 3^n$? Wir testen, ob diese Ungleichung für kleine n gilt und formulieren und beweisen dann eine allgemeine Aussage.

n	0	1	2	3	4	5	6	7	8
$n!$	1	1	2	6	24	120	720	> 5000	> 40000
3^n	1	3	9	27	81	243	729	< 2400	< 7500
$n! ? 3^n$	=	<	<	<	<	<	<	>	>

Mit vollständiger Induktion folgt $n! > 3^n$ für alle $n \geq 8$, denn für $n \geq 8$ gilt $(n+1)! = (n+1)n! > 3 \cdot 3^n = 3^{n+1}$. Die Ungleichung $n! > 3^n$ gilt also für alle $n \geq 7$.

SATZ 20. Für alle $a \in \mathbb{R}$ existiert ein $n_0 \in \mathbb{N}$, so dass $n! > a^n$ für alle $n \geq n_0$.

BEWEIS. Wir brauchen die Voraussetzung $a > 1$ nicht, weil $a^n < 1$ für alle a mit $|a| < 1$ und $a^n < |a|^n$ für alle $a \in \mathbb{R}$ gilt.

Sei $n_1 \in \mathbb{N}$ mit $n_1 > 2a$. Für alle $k \geq n_1$ gilt $\frac{a}{k} < \frac{1}{2}$ und damit

$$\frac{a^n}{n!} = \prod_{k=1}^n \frac{a}{k} = \left(\prod_{k=1}^{n_1-1} \frac{a}{k} \right) \left(\prod_{k=n_1}^n \frac{a}{k} \right) < \left(\prod_{k=1}^{n_1-1} \frac{a}{k} \right) \frac{1}{2^{n-n_1+1}}.$$

Hat man a und k_1 gewählt, so ist der Faktor $S := \prod_{k=1}^{n_1-1} \frac{a}{k}$ konstant. Wir wissen schon, dass die Folge 2^n unbeschränkt wächst. Also existiert ein Index $n_0 \geq n_1$, so dass $2^{n-n_1+1} > S$ für alle $n \geq n_0$. Für alle $n \geq n_0$ gilt nun $\frac{a^n}{n!} < 1$, also $a^n < n!$. $\square$

KAPITEL III

Primzahlen

Eine natürliche Zahl $p > 1$ heißt **Primzahl**, falls sie nur durch 1 und sich selbst teilbar ist. Wenn p eine Primzahl ist, so folgt aus $a \in \mathbb{N}$ und $a|p$, dass $a = 1$ oder $a = p$ ist.

SATZ 21 (Existenz der Primzahlzerlegung). *Jede natürliche Zahl $n > 1$ lässt sich als Produkt von Primzahlen schreiben oder ist selbst eine Primzahl.*

BEWEIS. Diese Aussage beweist man mit vollständiger Induktion. Die Zahl 2 ist eine Primzahl. Für $n > 2$ gilt entweder n ist eine Primzahl, oder n ist keine Primzahl. Im ersten Fall ist nichts weiter zu zeigen. Im zweiten Fall existieren $a, b \in \mathbb{N}$ mit $1 < a, b < n$ und $n = ab$. Da sich nach Induktionsvoraussetzung a und b selbst Primzahlen oder Produkte von Primzahlen sind, ist auch $n = ab$ ein Produkt von Primzahlen. $\square$

Für jede natürliche Zahl $n > 1$ existieren also Primzahlen $p_1 < p_2 < \dots < p_m$ und $k_j \in \mathbb{N}$, $k_j > 1$, so dass

$$(21) \quad n = p_1^{k_1} p_2^{k_2} \cdots p_m^{k_m} = \prod_{j=1}^m p_j^{k_j}.$$

Die Zerlegung in ein Produkt nennt man **Primzahlzerlegung**.

LEMMA 8. *Wenn p eine Primzahl ist und $p|ab$, dann gilt $p|a$ oder $p|b$.*

BEWEIS. Wir nehmen an, dass die Behauptung für eine Primzahl p nicht für alle a, b gilt. Sei p die kleinste Primzahl, für die Behauptung nicht gilt. Dann gibt es $a \leq b \in \mathbb{N}$ mit $p|ab$ und $p \nmid a$ und $p \nmid b$. Wir wählen a und b so, dass das Produkt minimal wird.

Wenn $p|ab$, dann $p|(a - mp)(b - np)$ für alle $m, n \in \mathbb{N}$, und wenn $p \nmid a$ und $p \nmid b$, dann $p \nmid (a - mp)$ und $p \nmid (b - np)$ für alle $m, n \in \mathbb{N}$. Also können wir $0 \leq a, b < p$ annehmen.

Wenn $p|ab$, dann $ab = cp$ für ein $c \in \mathbb{N}$. Wegen $0 \leq a, b < p$, folgt $1 < c < a$, da p eine Primzahl ist. Nun ist aber auch c eine Primzahl oder Produkt von Primzahlen. D.h. es gibt eine Primzahl q mit $q|ab$. Aus der Minimalität von p folgt $q|a$ oder $q|b$. Es gibt also Zahlen a' und b' mit $p|a'b'$, $p \nmid a'$, $p \nmid b'$ und $a'b' = cp/q < ab$. Dies ist ein Widerspruch zur Minimalität von ab . $\square$

SATZ 22 (Eindeutigkeit der Primzahlzerlegung). *Die Primzahlzerlegung ist eindeutig.*

BEWEIS. Wir müssen zeigen, dass aus

$$\prod_{j=1}^m p_j^{k_j} = \prod_{l=1}^r q_l^{l_j}$$

für Primzahlen $p_1 < \dots < p_m$, $q_1 < \dots < q_r$ und $k_j, l_j > 0$ folgt, dass $m = r$, $p_j = q_j$ und $k_j = l_j$ für alle j . Dies kann bequem mit vollständiger Induktion gezeigt werden. Da es nur eine Primzahl ≤ 2 gibt, folgt $2 = 2^1$, also $m = 1$, $p_1 = 2$ und $k_1 = 1$. Wenn

$$n = \prod_{j=1}^m p_j^{k_j} = \prod_{l=1}^r q_l^{l_j},$$

so folgt $p_1 | n$ und damit $p_1 = q_i$ für einen Index i und

$$m = p_1^{k_1-1} \prod_{j=2}^m p_j^{k_j} = q_i^{k_i-1} \prod_{l \neq i} q_l^{l_j} < n.$$

Aus der Induktionsvoraussetzung folgt $i = 1$, $p_j = q_j$ für alle j , $k_1 - 1 = l_1 - 1$ und $k_j = l_j$ für alle $j > 1$. $\square$

SATZ 23. *Es gibt unendlich viele Primzahlen.*

BEWEIS. Wir nehmen an, dass es nur endlich viele Primzahlen gibt. Diese sind $p_1, \dots, p_N$. Nun sei $n := \prod_{k=1}^N p_k + 1$. Dann gilt $n \not\equiv 0 \pmod{p_k}$ für alle k . Also ist n eine Primzahl. Andererseits gilt $n > p_k$ für alle k . Dies ist ein Widerspruch. $\square$

FOLGERUNG 9. *Eine natürliche Zahl n mit der Primzahlzerlegung $n = \prod_{j=1}^m p_j^{k_j}$ hat*

$$(22) \quad \prod_{j=1}^m (k_j + 1)$$

Teiler.

BEWEIS. Wenn $m | n$, dann $l_j \leq k_j$ für alle j in der Primzahlzerlegung $m = \prod_{j=1}^m p_j^{l_j}$, also $l_j \in \{0, 1, 2, \dots, k_j\}$ für alle j . $\square$

KAPITEL IV

Zahldarstellungen bezüglich verschiedener Basen

Für $b \in \mathbb{N}$, $b \neq 0, 1$, und Symbole $a_k \in \{0, 1, \dots, b-1\}$ heißt $(a_m a_{m-1} \dots a_1 a_0)_b$ die Darstellung der Zahl $\sum_{k=0}^m a_k b^k$ zur **Basis b**:

$$(23) \quad (a_m a_{m-1} \dots a_1 a_0)_b = \sum_{k=0}^m a_k b^k$$

1. Darstellung zur Basis b für natürliche Zahlen

SATZ 24. *Für alle $b \in \mathbb{N}$ mit $b \neq 0, 1$ besitzt jede natürliche Zahl eine eindeutige Darstellung zur Basis b .*

BEWEIS. Wir zeigen zuerst die Eindeutigkeit und dann die Existenz einer Darstellung zur Basis b .

Wenn $n = (a_m a_{m-1} \dots a_1 a_0)_b = (a'_m a'_{m-1} \dots a'_1 a'_0)_b$, dann sei d der größte Index mit $a_d \neq a'_d$. Nun gilt

$$\begin{aligned} (a_m a_{m-1} \dots a_1 a_0)_b &= (a'_m a'_{m-1} \dots a'_1 a'_0)_b \\ \sum_{k=0}^m a_k b^k &= \sum_{k=0}^m a'_k b^k \\ \sum_{k=0}^d a_k b^k + \sum_{k=d+1}^m a_k b^k &= \sum_{k=0}^d a'_k b^k + \sum_{k=d+1}^m a'_k b^k \\ \sum_{k=0}^d a_k b^k &= \sum_{k=0}^d a'_k b^k \\ (a_d - a'_d) b^d &= \sum_{k=0}^{d-1} (a'_k - a_k) b^k. \end{aligned}$$

Da $a_k, a'_k \in \{0, 1, \dots, b-1\}$ für alle k , folgt $|a'_k - a_k| \leq b-1$ für alle k . Nun folgt

$$\begin{aligned} \left| \sum_{k=0}^{d-1} (a'_k - a_k) b^k \right| &\leq \sum_{k=0}^{d-1} |a'_k - a_k| b^k \leq \sum_{k=0}^{d-1} (b-1) b^k \\ &= (b-1) \sum_{k=0}^{d-1} b^k = (b-1) \frac{1-b^d}{1-b} = b^d - 1 < b^d. \end{aligned}$$

Daraus folgt $a_d = a'_d$.

Wählt man $a_k \in \{0, 1, \dots, b-1\}$ für $k = 0, \dots, m$ beliebig, so kann man 2^{m+1} verschiedene natürliche Zahlen $\leq 2^{m+1} - 1$ darstellen. Es gibt aber genau $2^{m+1} - 1$ natürliche Zahlen $\leq 2^{m+1} - 1$. Darum besitzt jede natürliche Zahl eine Darstellung zur Basis b . $\square$

Für natürliche Zahlen erhält man die Darstellung zur Basis b mit Hilfe des Euklidischen Algorithmus, der auf der Division mit Rest beruht. Für $n \in \mathbb{N}$ und $b \neq 0$ existieren eindeutige natürliche Zahlen m, r mit $n = mb + r$ und $r < b$. Für $n \in \mathbb{N}$ erhält man die Darstellung zur Basis b , $(a_m a_{m-1} \dots a_1 a_0)_b$, rekursiv. Man setzt $n_0 = n$ und definiert n_k und a_k durch die Gleichung

$$(24) \quad n_k = n_{k+1}b + a_k \quad \text{mit } n_{k+1}, a_k \in \mathbb{N} \text{ und } 0 \leq a_k < b$$

bis $n_{m+1} = 0$ ist.

BEISPIEL 17. Wir suchen die Binärdarstellung, also $b = 2$, der Dezimalzahl 327.

$$\begin{aligned} n_0 &= 327 = 163 \cdot 2 + 1, & a_0 &= 1 \\ n_1 &= 163 = 81 \cdot 2 + 1, & a_1 &= 1 \\ n_2 &= 81 = 40 \cdot 2 + 1, & a_2 &= 1 \\ n_3 &= 40 = 20 \cdot 2 + 0, & a_3 &= 0 \\ n_4 &= 20 = 10 \cdot 2 + 0, & a_4 &= 0 \\ n_5 &= 10 = 5 \cdot 2 + 0, & a_5 &= 0 \\ n_6 &= 5 = 2 \cdot 2 + 1, & a_6 &= 1 \\ n_7 &= 2 = 1 \cdot 2 + 0, & a_7 &= 0 \\ n_8 &= 1 = 0 \cdot 2 + 1, & a_8 &= 1 \end{aligned}$$

Also gilt $(327)_{10} = (101000111)_2$.

2. Darstellung zur Basis b für reelle Zahlen

Für $x \in \mathbb{R}$ mit $x \geq 0$ gilt $x = [x] + x_0$ mit $0 \leq x_0 < 1$ und $[x] \in \mathbb{N}$. Dabei ist $[x]$ die größte natürliche Zahl $\leq x$. $(0, a_{-1}a_{-2}\dots)_b$ ist eine Darstellung der reellen Zahl x_0 mit $0 \leq x_0 < 1$ zu Basis b , falls $a_{-k} \in \{0, 1, \dots, b-1\}$ für alle k und $\sum_{k=1}^{\infty} a_{-k}b^{-k} = x_0$.

Die a_{-k} können wieder rekursiv berechnet werden durch

$$(25) \quad bx_k = a_{k-1} + x_{k-1} \text{ mit } a_{k-1} \in \{0, 1, \dots, b-1\} \text{ und } 0 \leq x_{k-1} < 1.$$

BEISPIEL 18. Für $x = (0, 25)_{10} = \frac{1}{4} = \frac{1}{2^2}$ muss $x = (0, 01)_2$ gelten. Auch der Algorithmus liefert $2x_0 = 0,5$, also $x_{-1} = 0,5$ und $a_{-1} = 0$, und $2x_{-1} = 1$, also $a_{-2} = 1$ und $x_{-2} = 0$.

BEISPIEL 19. Für $x = \frac{1}{3}$ bricht der Algorithmus zur Umwandlung in eine Binärzahl nie ab. Er liefert aber eine periodische Binärzahlentwicklung.

$$\begin{aligned} 2 \cdot \frac{1}{3} &= \frac{2}{3}, & a_{-1} &= 0, & x_{-1} &= \frac{2}{3} \\ 2 \cdot \frac{2}{3} &= \frac{4}{3} = 1 + \frac{1}{3}, & a_{-2} &= 1, & x_{-2} &= \frac{1}{3} \\ 2 \cdot \frac{1}{3} &= \frac{2}{3}, & a_{-3} &= 0, & x_{-3} &= \frac{2}{3} \\ 2 \cdot \frac{2}{3} &= \frac{4}{3} = 1 + \frac{1}{3}, & a_{-4} &= 1, & x_{-4} &= \frac{1}{3}. \end{aligned}$$

Daraus folgt $a_{-2k} = 1$ und $a_{-2k-1} = 0$ für alle $k \in \mathbb{N}$, d.h. $(1/3)_{10} = (0, \overline{01})_2$.

3. Schriftliches Rechnen im Binärsystem

Schriftliches Rechnen zur Basis b funktioniert (mit dem richtigen Übertrag) wie im Dezimalsystem.

BEISPIEL 20. Wir berechnen $(13)_{10} + (7)_{10} = (20)_{10}$ und $(13)_{10}(3)_{10} = (39)_{10}$ im Binärsystem. Es gilt $(13)_{10} = (1101)_2$, $(7)_{10} = (111)_2$, $(1101)_2 + (111)_2 = (10100)_2$, $(3)_{10} = (11)_2$ und $(1101)_2(11)_2 = (100111)_2$, denn

$$\begin{array}{r} \\ + \\ \hline \text{Übertrag} \\ \\ \hline 10100 \end{array} \quad \text{und} \quad \begin{array}{r} 1101 \cdot 11 \\ \\ \\ \hline \\ + 1101 \\ \hline 100111 \end{array}.$$

Die Subtraktion von Binärzahlen kann auf die Addition zurückgeführt werden. Es seien $a = (a_m \dots a_1 a_0)_2$ und $b = (b_m \dots b_1 b_0)_2$ zwei $m+1$ -stellige Binärzahlen. Wir definieren $\bar{b} := (\bar{b}_m \dots \bar{b}_1 \bar{b}_0)_2$ mit $\bar{b}_k = 1 - b_k$ für alle k . Das Komplement $\bar{b}$ entsteht also aus b durch Vertauschung der Nullen und Einsen. Nun gilt

$$\begin{aligned} b + \bar{b} &= \sum_{k=0}^m b_k 2^k + \sum_{k=0}^m (1 - b_k) 2^k = \sum_{k=0}^m 2^k = \frac{1 - 2^{m+1}}{1 - 2} = 2^{m+1} - 1 \\ -b &= \bar{b} + 1 - 2^{m+1} \\ a - b &= a + \bar{b} + 1 - 2^{m+1}. \end{aligned}$$

Da a und b nur $m+1$ Stellen haben, bedeutet die Subtraktion von 2^{m+1} nur die Streichung der führenden 1.

BEISPIEL 21. Wir berechnen $(13)_{10} - (7)_{10} = (6)_{10}$. Es gilt $a = (1101)_2$ und $b = (0111)_2$, also $m = 3$. Nun folgt

$$\begin{aligned} \bar{b} &= (1000)_2 \\ a + \bar{b} &= (10101)_2 \\ a + \bar{b} + 1 &= (10110)_2 \\ a + \bar{b} + 1 - 2^4 &= (0110)_2 = (6)_{10}. \end{aligned}$$

KAPITEL V

Funktionen – Grundbegriffe

Eine (reellwertige) Funktion ist eine Abbildung $f : \mathbb{R} \supset D \rightarrow \mathbb{R}$ die jedem Element $x \in D$ eine reelle Zahl $f(x)$ zuordnet. Die Teilmenge D heißt **Definitionsbereich** der Funktion f .

BEMERKUNG 25. Formal ist eine Abbildung $f : X \rightarrow Y$ von einer Menge X in eine Menge Y definiert durch eine Teilmenge $\Gamma_f \subset X \times Y$ mit der Eigenschaft

$$\forall x \in X \exists! y \in Y \text{ mit } (x, y) \in \Gamma_f.$$

Dann $f(x) := y$, wobei $(x, y) \in \Gamma_f$. $\square$

BEISPIEL 22. Für die kinetische Energie E gilt $E = \frac{1}{2}mv^2$. Dabei sind m die Masse und v die Geschwindigkeit des Körpers. Nimmt man an, dass die Masse konstant ist, $m = m_0 \in \mathbb{R}$, so kann man die kinetische Energie als Funktion der Geschwindigkeit auffassen: $E(v) = \frac{1}{2}m_0v^2$, $E : \mathbb{R} \rightarrow \mathbb{R}$. Die Geschwindigkeit v wird als unabhängige Größe, die konstante Masse m_0 als Parameter und E als abhängige Größe angesehen.

BEISPIEL 23. Das Ohmsche Gesetz besagt $RI = U$.

Nimmt man an, dass der Widerstand konstant ist, $R = R_0 \in \mathbb{R}$, so kann man die Stromstärke als Funktion der Spannung auffassen: $I(U) = \frac{1}{R_0}U$, $I : \mathbb{R} \rightarrow \mathbb{R}$.

Nimmt man an, dass die Spannung konstant ist, $U = U_0 \in \mathbb{R}$, so kann man die Stromstärke als Funktion des Widerstandes auffassen: $I(R) = U_0 \frac{1}{R}$, $I : \{x > 0\} \rightarrow \mathbb{R}$. Hier darf die unabhängige Größe R nur von Werten annehmen, die größer als 0 sind.

Eine Funktion, deren Definitionsbereich $\mathbb{N}$ ist, heißt **Folge**. Man bezeichnet die Funktionswerte als Folgenglieder $x_n := f(n)$ und schreibt $\{x_n\}_{n \in \mathbb{N}}$ oder $(x_n)_{n \in \mathbb{N}}$ statt $f : \mathbb{N} \rightarrow \mathbb{R}$, $f(n) = x_n$.

Wir haben in Abschnitt 8 die Folgen $\{n^k\}$, $\{a^n\}$ und $\{n!\}$ untersucht.

BEMERKUNG 26. Die Folge $f(n) = n^k$ kann sinnvoll zu einer Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ fortgesetzt werden, nämlich $f(x) = x^k$ für $x \in \mathbb{R}$. Auch wenn die ursprüngliche Größe, die hinter der unabhängigen Variablen n steckt, nur natürliche Werte annehmen kann, so ist es doch oft sinnvoll, den Definitionsbereich der Funktion zu vergrößern, um Werkzeuge aus der Differentialrechnung, z.B. zur Extremwertbestimmung, anwenden zu können. Wir werden später a^x und $x!$ für $x \in \mathbb{R}$ definieren. $\square$

Für $a, b \in \mathbb{R}$ mit $a < b$ heißt $(a, b) := \{x \in \mathbb{R} : a < x < b\}$ **offenes Intervall** und $[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\}$ **abgeschlossenes Intervall**. Diese Bezeichnungen kann man auch mischen, $[a, b) = \{x \in \mathbb{R} : a \leq x < b\}$, oder für $a = -\infty$ bzw. $b = \infty$ verwenden, z.B. $(0, \infty) = \{x \in \mathbb{R} : 0 < x\}$. Eine Teilmenge $D \subset \mathbb{R}$ heißt **offen**, wenn sie Vereinigung beliebig vieler offener Intervalle ist, z.B. $\emptyset$, $(0, 1) \cup (1, 2)$ oder $\mathbb{R}$.

1. Graphische Darstellung einfacher Funktionen

Der Graph Γ_f einer Funktion $f : D \rightarrow \mathbb{R}$ ist die Punktmenge

$$\Gamma_f := \{(x, f(x)) : x \in D\} \subset \mathbb{R}^2.$$

Will man den Graph einer Funktion skizzieren, so muss man sich für einen charakteristischen, aussagekräftigen Ausschnitt der Ebene $\mathbb{R}^2$ entscheiden. Für einfache Funktionen, wie in den Beispielen dieses Abschnitts, reicht es oft aus

- spezielle Funktionswerte auszurechnen und Nullstellen zu bestimmen,
- die Steigung (siehe Abschnitt IX) in wichtigen Punkten auszurechnen
- und die berechneten Funktionswerte in den Stetigkeitsbereichen (siehe Abschnitt VI) so zu verbinden, dass der Graph die richtige Steigung hat,

um eine grobe Skizze des Funktionsgraphen anzufertigen.

1.1. Konstante Funktionen. Für $c \in \mathbb{R}$ bezeichnet $f \equiv c$ die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = c$ für alle $x \in \mathbb{R}$.

1.2. Lineare Funktionen. Falls $f(x) = ax + b$ für $a, b \in \mathbb{R}$, so heißt $f : \mathbb{R} \rightarrow \mathbb{R}$ lineare Funktion. Es gilt $f(0) = b$. Der Anstieg der Funktion f ist konstant a .

Der Graph Γ_f der Funktion f ist also die Gerade durch die Punkte $(0, b)$ und $(1, a + b)$.

Falls $a \neq 0$, so ist der Graph Γ_f die Gerade durch die Punkte $(0, b)$ und $(-b/a, 0)$.

Abbildung 1 zeigt den Graph der Funktion $f(x) = \frac{1}{2}x + b$ für $b = 1$, $b = 0$ und $b = -1$.

Wenn $a = 0$, so ist $f \equiv b$ eine konstante Funktion. Wenn $a \neq 0$, so nimmt f jede reelle Zahl genau einmal als Funktionswert an.

1.3. Die Funktionen $f(x) = ax^2$. Für alle $a \neq 0$ hat die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = ax^2$, genau eine Nullstelle, $x = 0$. Außerdem gilt $f(\pm 1) = a$. Falls $a > 0$, so gilt $f(x) \geq 0$ für alle $x \in \mathbb{R}$. Falls $a < 0$, so gilt $f(x) \leq 0$ für alle $x \in \mathbb{R}$. Abbildung 2 zeigt den Graph der Funktion $f(x) = ax^2$ für $a = 1$, $a = 1/2$ und $a = -1$.

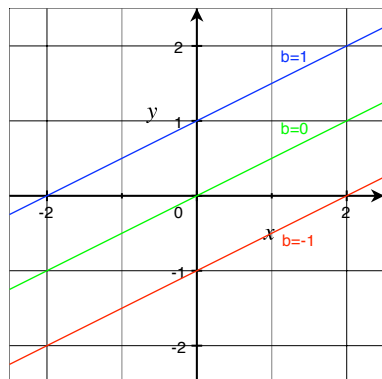


ABBILDUNG 1. $\frac{1}{2}x + b$
für $b = -1, 0, 1$

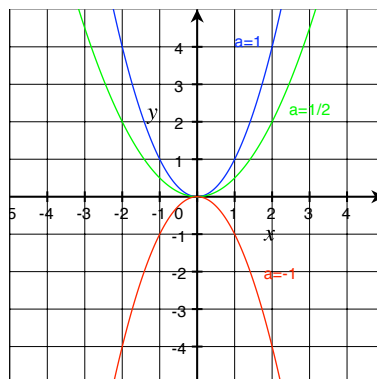


ABBILDUNG 2. ax^2
für $a = 1, \frac{1}{2}, -1$

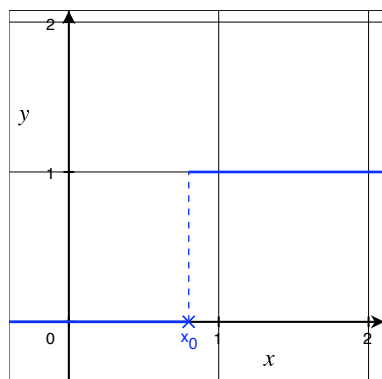


ABBILDUNG 3. Sprungfunktion

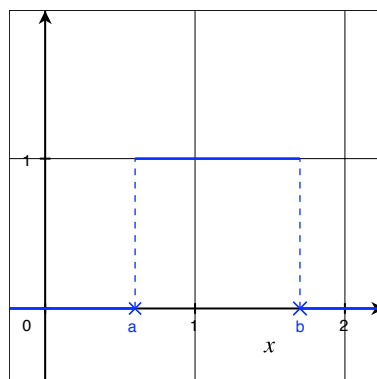


ABBILDUNG 4. Rechteckfunktion

1.4. Die Funktion $f(x) = 1/x$. Wir betrachten $f(x) = 1/x$ für $x \neq 0$, also auf dem Definitionsbereich $\mathbb{R} \setminus \{0\}$.

Es gilt $f(1) = 1$, $f(2) = 1/2$, $f(4) = 1/4$, $f(1/2) = 2$, $f(1/4) = 4$, $\lim_{n \rightarrow \infty} f(n) = \lim_{n \rightarrow \infty} 1/n = 0$ und $\lim_{n \rightarrow \infty} f(1/n) = \lim_{n \rightarrow \infty} n = \infty$.

1.5. Die Sprungfunktion. Für jedes $x_0 \in \mathbb{R}$ hat $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} 1 & , \text{ falls } x \geq x_0 \\ 0 & , \text{ falls } x < x_0 \end{cases}$$

eine Sprungstelle in $x = x_0$. Abbildung 3 zeigt den Graph einer Sprungfunktion mit Sprungstelle x_0 .

2. Verknüpfung von Funktionen

Es seien $f, g : D \rightarrow \mathbb{R}$ zwei Funktionen. Dann kann man durch Anwendung der **Rechenoperationen** auf die Funktionswerte $f(x)$ und $g(x)$ neue

Funktionen bilden, nämlich die Summe $f + g$, das Produkt fg und, falls $g(x) \neq 0$ für alle $x \in D$, auch den Quotienten f/g :

$$(26) \quad (f + g)(x) := f(x) + g(x), \quad (fg)(x) := f(x)g(x), \quad \left(\frac{f}{g}\right)(x) := \frac{f(x)}{g(x)}$$

BEISPIEL 24. Betrachtet man Stromstärke $I(t)$ und Spannung $U(t)$ als zeitabhängige Funktionen, so ist die Leistung $P(t)$ als zeitabhängige Funktion das Produkt der Funktionen U und I , also $P = UI$ und $P(t) = U(t)I(t)$.

BEMERKUNG 27. Der Wert der Funktionen $f + g$, fg und f/g an einer Stelle x_0 hängt nur von den Funktionswerten $f(x_0)$ und $g(x_0)$ ab. $\square$

BEMERKUNG 28. Durch Addition und Multiplikation kann man aus den Funktionen $g \equiv c$ für beliebiges $c \in \mathbb{R}$ und $f(x) = x$ alle Polynome konstruieren. Die Funktionen, die man aus der Identität $f(x) = x$ und den konstanten Funktionen $g \equiv c$ durch Addition, Multiplikation und Division bilden kann, heißen **rationale** Funktionen. $\square$

Sind $f : D_2 \rightarrow \mathbb{R}$ und $g : D_1 \rightarrow \mathbb{R}$ zwei Funktionen mit $g(x) \in D_2$ für alle $x \in D_1$, so heißt die Funktion

$$(27) \quad f \circ g : D_1 \rightarrow \mathbb{R}, \quad (f \circ g)(x) := f(g(x)),$$

die **Verkettung** bzw. Hintereinanderausführung der Funktionen f und g .

BEISPIEL 25. Für $g : \mathbb{R} \rightarrow \mathbb{R}$ mit $g(x) = x + 2$ und $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ gilt $(f \circ g)(x) = f(g(x)) = f(x + 2) = (x + 2)^2 = x^2 + 4x + 4$.

BEISPIEL 26. Für jedes $a \in \mathbb{R}$ ist die Verkettung der Funktionen $g : \mathbb{R} \rightarrow \mathbb{R}$, $g(x) = x - a$, mit der Sprungfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$,

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

die Funktion

$$(f \circ g)(x) = f(g(x)) = f(x - a) = \begin{cases} 1 & , \text{ falls } x \geq a \\ 0 & , \text{ falls } x < a \end{cases}.$$

Die Verkettung $f(x - a)$ verschiebt die Sprungstelle $x = 0$ von f nach $x = a$.

Für alle $a, b \in \mathbb{R}$ mit $a < b$ erhält man die Rechteckfunktion als

$$f(x - a) - f(x - b) = \begin{cases} 0 & , \text{ falls } x < a \\ 1 & , \text{ falls } a \leq x < b \\ 0 & , \text{ falls } b \leq x \end{cases}.$$

Abbildung 4 zeigt den Graph einer Rechteckfunktion mit Sprungstellen a und b .

3. Die Exponentialfunktion

Die Exponentialfunktion $\exp : \mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto e^x$, kann auf mehrere Arten definiert werden:

- als Reihe

$$e^x := \sum_{n=0}^{\infty} \frac{1}{n!} x^n$$

- als Grenzwert einer Folge

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

- als Lösung der Differentialgleichung $y' = y$ mit $y(0) = 1$

Jeder der drei Zugänge hat seine Vor- und Nachteile. In jedem Fall benötigt man mathematische Sätze, um die Existenz des durch eine Reihe, Folge bzw. Differentialgleichung definierten Objektes „Exponentialfunktion“ zu zeigen und seine Eigenschaften, z.B. die Potenzgesetze $e^{a+b} = e^a e^b$, zu ermitteln.

BEMERKUNG 29. Wir definieren die Exponentialfunktion hier als Lösung einer Differentialgleichung mit Anfangsbedingung, weil dies der klassischen Definition entspricht und Differentialgleichungen, neben den Gleichungen wie in den Beispielen 22 und 23, ein Hilfsmittel sind, um Zusammenhänge zwischen physikalischen Größen zu beschreiben. $\square$

BEMERKUNG 30. Die Differentialgleichung $y' = y$ ist eine der einfachsten Differentialgleichungen der Form $y' = f(y)$ und ein Spezialfall der Differentialgleichungen $y' = ky$ mit dem Parameter k . Die Differentialgleichung $y' = ky$ beschreibt für $k > 0$ Wachstumsprozesse und für $k < 0$ Zerfallsprozesse. $\square$

Die Exponentialfunktion ist die Lösung der Differentialgleichung $y' = y$ mit der Anfangsbedingung $y(0) = 1$.

Wir benutzen die Bezeichnung $e^x := y(x)$. Es gilt also $e^0 = 1$ und $(e^x)' = e^x$.

SATZ 25 (Funktionalgleichung). Für alle $a, b \in \mathbb{R}$ gilt $e^a > 0$ und

$$(28) \quad e^{a+b} = e^a e^b.$$

BEWEIS. Wenn $e^a = 0$ für ein $a \in \mathbb{R}$, so hätte das Anfangswertproblem $y' = y$ mit $y(a) = 0$ zwei Lösungen, nämlich e^x und $y \equiv 0$. Da jede Anfangswertaufgabe zur Differentialgleichung $y' = y$ eindeutig lösbar ist, $e^0 = 1 > 0$ und die Exponentialfunktion als differenzierbare Funktion stetig ist, folgt $e^x > 0$ für alle $x \in \mathbb{R}$.

Wir betrachten die Hilfsfunktion $h(x) = \frac{e^{x+a}}{e^x}$. Diese ist wohldefiniert, da $e^x > 0$ für alle $x \in \mathbb{R}$. Da die Exponentialfunktion eine Lösung der Differentialgleichung $y' = y$ ist, folgt

$$h'(x) = \frac{(e^{x+a})' e^x - e^{x+a} (e^x)'}{e^x e^x} = \frac{e^{x+a} e^x - e^{x+a} e^x}{e^x e^x} = 0.$$

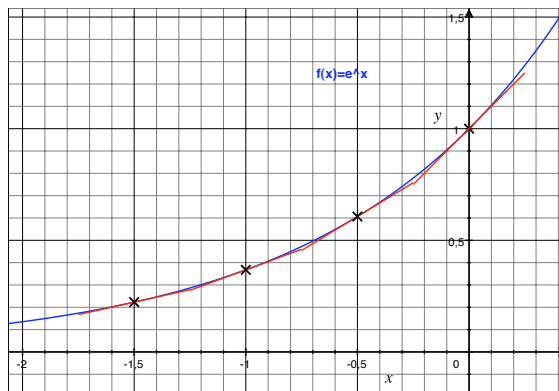


ABBILDUNG 5. Die Exponentialfunktion

Aus $h' \equiv 0$ folgt $h \equiv c$ für eine Konstante $c \in \mathbb{R}$. Wegen $h(0) = \frac{e^a}{e^0} = e^a$ und $e^0 = 1$, folgt $c = e^a$, also $e^{x+a} = e^x e^a$ für alle $a, x \in \mathbb{R}$. $\square$

Da die Exponentialfunktion die Differentialgleichung $y' = y$ erfüllt, gilt $(e^x)' = e^x$, $(e^x)'' = (e^x)' = e^x$ usw. In jedem Punkt (x, e^x) ist der Anstieg der Tangente an den Graph der Exponentialfunktion e^x . Insbesondere hat die Tangente im Punkt $(0, 1)$ an den Graph der Exponentialfunktion den Anstieg 1. Da $(e^x)'' > 0$ für alle $x \in \mathbb{R}$, liegt der Graph der Exponentialfunktion immer über den Tangenten.

BEMERKUNG 31. Für den Beweis der Funktionalgleichung braucht man den Existenz- und Eindeigkeitsatz für die Lösung von Anfangswertproblemen, den Begriff der Ableitung, die Ableitungsregel für Quotienten, den Satz „Wenn $h' \equiv 0$, dann ist h konstant“, den Begriff der Stetigkeit und den Zwischenwertsatz. Um den Graph der Exponentialfunktion zu skizzieren brauchten wir die geometrische Bedeutung der ersten und der zweiten Ableitung. $\square$

Abbildung 5 zeigt den Graph der Exponentialfunktion (blau), spezielle Funktionswerte (schwarz) e^x und Tangenten an den Graph in diesen Punkten (rot) für $x = 0$, $x = -1/2$, $x = -1$ und $x = -3/2$.

4. Die Winkelfunktionen

Die Differentialgleichung $y'' + y = 0$ ist der Spezialfall $k = 1$ der Schwingungsgleichung $y'' + ky = 0$ für $k \in \mathbb{R}$, $k > 0$, die die Bewegung eines Federpendels oder auch den zeitlichen Verlauf der Stromstärke in einem elektrischen Schwingkreis beschreibt.

Wir definieren die Funktionen $\sin$ und $\cos$ als Lösungen dieser Differentialgleichung mit speziellen Anfangsbedingungen.

Die **Sinusfunktion** ist die Lösung der Differentialgleichung

$$y'' + y = 0 \text{ mit } y(0) = 0 \text{ und } y'(0) = 1.$$

Es gilt also $(\sin x)'' + \sin x = 0$, $\sin 0 = 0$ und $\sin' 0 = 1$.

Die **Kosinusfunktion** ist die Lösung der Differentialgleichung

$$y'' + y = 0 \text{ mit } y(0) = 1 \text{ und } y'(0) = 0.$$

Es gilt also $(\cos x)'' + \cos x = 0$, $\cos 0 = 1$ und $\cos' 0 = 0$.

BEMERKUNG 32. Jede Lösung der Differentialgleichung $y'' + y = 0$ ist von der Form $a \sin x + b \cos x$ für $a, b \in \mathbb{R}$. $\square$

4.1. Ableitungen der Winkelfunktionen.

SATZ 26. Für alle $x \in \mathbb{R}$ gilt

$$(29) \quad (\sin x)' = \cos x \text{ und } (\cos x)' = -\sin x.$$

BEWEIS. Für die Hilfsfunktion $h(x) := (\sin x)'$ gilt $h'(x) = (\sin x)'' = -\sin x$ und $h''(x) = -(\sin x)' = -h(x)$. Die Funktion h ist also auch eine Lösung der Differentialgleichung $y'' + y = 0$. Es gilt $h(0) = \sin' 0 = 1$ und $h'(0) = -\sin 0 = 0$. Die Funktion h erfüllt also die gleichen Anfangsbedingungen wie $\cos$. Daraus folgt $h(x) = \cos x$. Weiterhin folgt aus $\cos x = (\sin x)'$ sofort $(\cos x)' = (\sin x)'' = -\sin x$. $\square$

4.2. Kreisgleichung.

SATZ 27. Für alle $x \in \mathbb{R}$ gilt

$$(30) \quad \sin^2 x + \cos^2 x = 1.$$

BEWEIS. Für die Hilfsfunktion $h(x) := \sin^2 x + \cos^2 x$ gilt

$$h'(x) = 2 \sin x (\sin x)' + 2 \cos x (\cos x)' = 2 \sin x \cos x - 2 \cos x \sin x = 0.$$

Daraus folgt, dass h eine konstante Funktion ist. Da $h(0) = (\sin 0)^2 + (\cos 0)^2 = 1$, gilt $h \equiv 1$. $\square$

4.3. Beschränktheit.

FOLGERUNG 10. Für alle $x \in \mathbb{R}$ gilt $|\sin x| \leq 1$ und $\cos x \leq 1$.

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt (auf D) **beschränkt**, falls ein $S \in \mathbb{R}$ mit $|f(x)| \leq S$ für alle $x \in D$ existiert.

Die Funktionen $\sin$ und $\cos$ sind auf $\mathbb{R}$ beschränkt.

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt (auf D) **nach oben beschränkt**, falls ein $S \in \mathbb{R}$ mit $f(x) \leq S$ für alle $x \in D$ existiert.

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt (auf D) **nach unten beschränkt**, falls ein $S \in \mathbb{R}$ mit $f(x) \geq S$ für alle $x \in D$ existiert.

Beschränkte Funktionen sind sowohl nach unten als auch nach oben beschränkt.

BEISPIEL 27. Die Funktion $f(x) = 3x + 1$ ist auf $\mathbb{R}$ nicht beschränkt. Aber sie ist z.B. auf dem Intervall $[-2, 1]$ beschränkt, denn aus $-2 \leq x \leq 1$ folgt $-5 \leq 3x + 1 \leq 4$, also $|f(x)| \leq 5$ für alle $x \in [-2, 1]$.

BEISPIEL 28. Die Funktion $f(x) = 1/x$ ist auf der Menge $\{x > 0\}$ nach unten beschränkt, denn $1/x > 0$ für alle $x > 0$. Sie ist auf der Menge $\{x > 0\}$ nicht nach oben beschränkt, denn $\lim_{x \rightarrow 0} 1/x = \infty$. Jedoch ist die Funktion $f(x) = 1/x$ auf dem Intervall $[1, \infty)$ beschränkt, denn $1/x \leq 1$ für alle $x \geq 1$.

4.4. Symmetrie.

SATZ 28. Für alle $x \in \mathbb{R}$ gilt

$$(31) \quad \sin(-x) = -\sin x \text{ und } \cos(-x) = \cos x.$$

BEWEIS. Für die Hilfsfunktion $h(x) = \sin(-x)$ gilt $h'(x) = -\cos(-x)$ und $h''(x) = -\sin(-x) = -h(x)$. Die Funktion h erfüllt die Differentialgleichung $y'' + y = 0$. Da $h(0) = \sin(-0) = 0$ und $h'(0) = -\cos 0 = -1$ folgt $h(x) = -\sin x$, da auch $-\sin x$ die Differentialgleichung und diese Anfangsbedingungen erfüllt.

Leitet man die Gleichung $\sin(-x) = -\sin x$ ab, so erhält man $-\cos(-x) = -\cos x$. $\square$

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **gerade**, falls $f(-x) = f(x)$ für alle $x \in \mathbb{R}$ gilt.

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **ungerade**, falls $f(-x) = -f(x)$ für alle $x \in \mathbb{R}$ gilt.

BEISPIEL 29. Die Funktionen $\cos x$, $f(x) = x^2$ und jede konstante Funktion $g \equiv c$ sind gerade, denn $\cos(-x) = \cos x$ (siehe 28), $f(-x) = (-x)^2 = (-1)^2 x^2 = x^2 = f(x)$ und $g(-x) = c = g(x)$.

BEISPIEL 30. Die Funktionen $\sin x$ und $f(x) = x$ sind ungerade, denn $\sin(-x) = -\sin x$ (siehe Satz 28) und $f(-x) = -x = -f(x)$.

Der Graph einer geraden Funktion ist symmetrisch bezüglich Spiegelung an der y -Achse. Der Graph einer ungeraden Funktion ist symmetrisch bezüglich Drehung um 180° um den Koordinatenursprung. Der Graph einer ungeraden Funktion bleibt unverändert, wenn man ihn erst an der x -Achse und dann das Bild an der y -Achse spiegelt.

4.5. Additionstheoreme.

SATZ 29. Für alle $x, y \in \mathbb{R}$ gilt

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

$$\cos(x + y) = \cos(x) \cos(y) - \sin(x) \sin(y).$$

BEWEIS. Für jedes $a \in \mathbb{R}$ besitzt die Hilfsfunktion

$$h(x) := \sin(x + a) - \sin x \cos a - \sin a \cos x$$

folgende Eigenschaften: $h(0) = 0$,

$$h'(x) = \cos(x + a) - \cos x \cos a + \sin a \sin x,$$

$$h'(0) = \cos a - \cos 0 \cos a + \sin a \sin 0 = 0,$$

$$h''(x) = -\sin(x + a) + \sin x \cos a + \sin a \cos x = -h(x).$$

Die Funktion h erfüllt die Differentialgleichung $y'' + y = 0$ und die Anfangsbedingungen $y(0) = y'(0) = 0$. Daraus folgt $h \equiv 0$ und $h' \equiv 0$. $\square$

FOLGERUNG 11. Für alle $x, y \in \mathbb{R}$ gilt

$$(32) \quad \sin(2x) = 2 \sin(x) \cos(x), \quad \cos(2x) = \cos^2 x - \sin^2 x = 2 \cos^2 x - 1.$$

BEWEIS. Die Behauptung folgt aus Satz 29 mit $x = y$ und der Gleichung $\sin^2 x = 1 - \cos^2 x$. $\square$

LEMMA 9. Die Funktion $\cos x$ besitzt eine Nullstelle.

BEWEIS. Wir nehmen $\cos x \neq 0$ für alle $x \in \mathbb{R}$ an. Dann gilt $\cos x > 0$ für alle $x \in \mathbb{R}$, weil $\cos 0 = 1 > 0$ und $\cos x$ stetig ist. Daraus folgt, dass $\sin x$ monoton wachsend ist, weil $(\sin x)' = \cos x > 0$. Daraus folgt $\sin x > 0$ für alle $x > 0$, weil $\sin 0 = 0$ und $\sin' 0 > 0$. Also ist $\cos x$ monoton fallend für $x > 0$, da $(\cos x)' = -\sin x < 0$.

Da $\cos$ eine nichtkonstante Funktion ist und $\cos x \leq 1$ für alle $x \in \mathbb{R}$ gilt, existieren ein $x_0 \in \mathbb{R}$ und $q \in \mathbb{R}$ mit $0 < q < 1$ und $\cos x_0 = q$. Da $\cos(x + x_0) = \cos x \cos x_0 - \sin x \sin x_0 \leq q \cos x$ für alle $x > 0$, folgt mit vollständiger Induktion $\cos(nx_0) \leq q^n$. Daraus folgt $\lim_{x \rightarrow \infty} \cos x = 0$ und damit $\lim_{x \rightarrow \infty} \sin x = 1$. Es existiert also ein $x_1 \in \mathbb{R}$ mit $\sin x_1 = \cos x_1$. Nun gilt $\cos(2x_1) = \cos x_1 \cos x_1 - \sin x_1 \sin x_1 = 0$. $\square$

$$(33) \quad \frac{\pi}{2} := \text{kleinste positive Nullstelle der Funktion } \cos x$$

SATZ 30 (Spezielle Funktionswerte). Es gilt

$$(34) \quad \sin\left(\frac{\pi}{2}\right) = 1, \sin \pi = 0, \cos \pi = -1, \sin(2\pi) = 0 \text{ und } \cos(2\pi) = 1.$$

BEWEIS. Aus $\cos^2 x + \sin^2 x = 1$ folgt $\sin(\pi/2) = \pm 1$. Da $\sin x$ für $x \in [0, \pi/2]$ monoton wachsend ist, gilt $\sin \frac{\pi}{2} = 1$. Aus den Additionstheoremen für den doppelten Winkel folgt

$$\begin{aligned} \sin \pi &= \sin\left(2\frac{\pi}{2}\right) = 2 \sin\left(\frac{\pi}{2}\right) \cos\left(\frac{\pi}{2}\right) = 0 \\ \cos \pi &= \cos\left(2\frac{\pi}{2}\right) = \cos^2\left(\frac{\pi}{2}\right) - \sin^2\left(\frac{\pi}{2}\right) = -1 \\ \sin(2\pi) &= 2 \sin(\pi) \cos(\pi) = 0 \\ \cos(2\pi) &= \cos^2(\pi) - \sin^2(\pi) = 1. \end{aligned}$$

$\square$

Aus dem Additionstheorem für den doppelten Winkel folgt auch, dass π die kleinste positive Nullstelle der Funktion $\sin x$ ist.

Abbildung 6 zeigt die Graphen der Winkelfunktionen auf dem Intervall $[0, \pi/2]$. Die speziellen Funktionswerte $\sin 0$, $\sin(\pi/2)$, $\cos 0$ und $\cos(\pi/2)$ und die Tangenten (schwarz) an den Graph der Funktionen in den Punkten $x = 0$ und $x = \pi/2$ sind eingezeichnet. Sie liefern einen Anhaltspunkt für das Aussehen der Graphen.

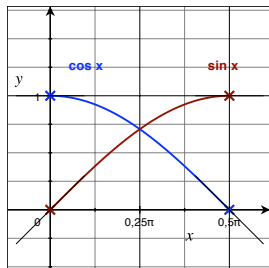


ABBILDUNG
6. $\sin$ und $\cos$
auf dem Intervall
 $[0, \pi/2]$

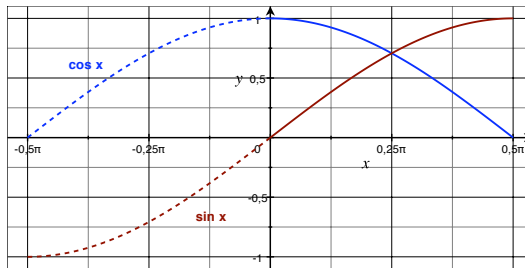


ABBILDUNG 7. Gerade
bzw. ungerade Fortsetzung
auf das Intervall
 $[-\pi/2, \pi/2]$

Abbildung 7 zeigt die gerade Fortsetzung der Funktion $\cos x$ und die ungerade Fortsetzung der Funktion $\sin x$ vom Intervall $[0, \pi/2]$ auf das Intervall $[-\pi/2, \pi/2]$.

SATZ 31 (Verschiebung und Periodizität). Für alle $x \in \mathbb{R}$ gilt

$$(35) \quad \sin\left(x + \frac{\pi}{2}\right) = \cos x, \quad \sin(x + \pi) = -\sin x,$$

$$(36) \quad \cos\left(x + \frac{\pi}{2}\right) = -\sin x, \quad \cos(x + \pi) = -\cos x,$$

$$(37) \quad \sin(x + 2\pi) = \sin x, \quad \cos(x + 2\pi) = \cos x.$$

BEWEIS. Aus den Additionstheoremen folgt mit $y = \pi/2$

$$\begin{aligned} \sin\left(x + \frac{\pi}{2}\right) &= \sin(x) \cos\left(\frac{\pi}{2}\right) + \sin\left(\frac{\pi}{2}\right) \cos(x) = \cos x \\ \cos\left(x + \frac{\pi}{2}\right) &= \cos(x) \cos\left(\frac{\pi}{2}\right) - \sin\left(\frac{\pi}{2}\right) \sin(x) = -\sin x, \end{aligned}$$

mit $y = \pi$

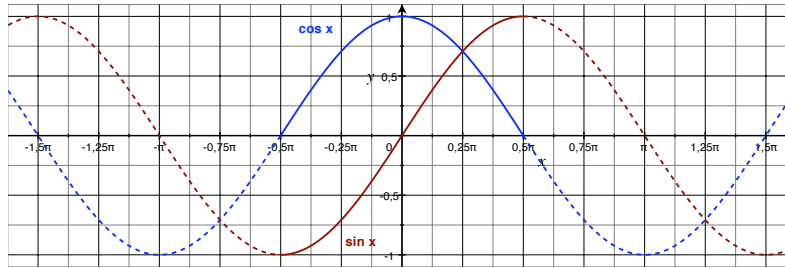
$$\begin{aligned} \sin(x + \pi) &= \sin(x) \cos(\pi) + \sin(\pi) \cos(x) = -\sin x \\ \cos(x + \pi) &= \cos(x) \cos(\pi) - \sin(\pi) \sin(x) = -\cos x, \end{aligned}$$

und mit $y = 2\pi$

$$\begin{aligned} \sin(x + 2\pi) &= \sin(x) \cos(2\pi) + \sin(2\pi) \cos(x) = \sin x \\ \cos(x + 2\pi) &= \cos(x) \cos(2\pi) - \sin(2\pi) \sin(x) = \cos x. \end{aligned}$$

□

Abbildung 8 zeigt die Fortsetzung der Funktionen $\sin x$ und $\cos x$ vom Intervall $[-\pi/2, \pi/2]$ auf $\mathbb{R}$ mit Hilfe der Beziehungen $\sin(x + \pi) = -\sin x$ und $\cos(x + \pi) = -\cos x$.

ABBILDUNG 8. Die 2π -periodischen Funktionen $\sin x$ und $\cos x$

5. Injektivität, Monotonie und Umkehrfunktion

Für eine Funktion $f : D \rightarrow \mathbb{R}$ heißt die Menge aller Funktionswerte $f(D) := \{f(x) : x \in D\}$ das **Bild** von f bzw. der Menge D unter f .

BEISPIEL 31. Das Bild der Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist das Intervall $[0, \infty)$, denn $x^2 \geq 0$ für alle $x \in \mathbb{R}$ und zu jedem $y \geq 0$ existiert $\sqrt{y}$ und $f(\sqrt{y}) = y$.

BEISPIEL 32. Das Bild der Funktion $f : \{x > 0\} \rightarrow \mathbb{R}$ mit $f(x) = 1/x$ ist das offene Intervall $(0, \infty)$, denn $1/x > 0$ für alle $x > 0$ und zu jedem $y > 0$ existiert $1/y$ und $f(1/y) = y$.

BEISPIEL 33. Das Bild der Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \cos$ ist das abgeschlossene Intervall $[-1, 1]$, denn $|\cos x| \leq 1$ für alle $x \in \mathbb{R}$, $\cos 0 = 1$, $\cos \pi = -1$ und $\cos x$ ist als differenzierbare Funktion stetig, nimmt also alle Werte im Intervall $[-1, 1]$ als Funktionswerte an (siehe Abschnitt 3).

BEISPIEL 34. Das Bild der Sprungfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

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

besteht nur aus zwei Elementen $f(\mathbb{R}) = \{0, 1\}$, denn diese Sprungfunktion nimmt nur die Funktionswerte 0 und 1 an.

BEISPIEL 35. Das Bild der Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} x & , \text{ falls } x \leq 0 \\ x + 2 & , \text{ falls } x > 0 \end{cases}$$

ist die Menge $(-\infty, 0] \cup (2, \infty)$.

Für eine Funktion $f : D \rightarrow \mathbb{R}$ heißt eine Funktion $g : f(D) \rightarrow \mathbb{R}$ **Umkehrfunktion** der Funktion f auf D , falls $f(x) = y \Leftrightarrow g(y) = x$ für alle $x \in D$. Die Umkehrfunktion wird auch mit f^{-1} bezeichnet.

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt **injektiv** auf D , falls für alle $x_1, x_2 \in D$ gilt: $f(x_1) = f(x_2) \Rightarrow x_1 = x_2$.

SATZ 32. *Eine Funktion $f : D \rightarrow \mathbb{R}$ besitzt genau dann eine Umkehrfunktion auf D , wenn f auf D injektiv ist.*

BEWEIS. Durch $g(f(x)) := x$ ist eine Funktion $g : f(D) \rightarrow \mathbb{R}$ wohldefiniert, da es zu jedem $y \in f(D)$ genau ein $x \in D$ mit $f(x) = y$ gibt. $\square$

BEMERKUNG 33. Man erhält den Graph der Umkehrfunktion f^{-1} durch Spiegelung des Graphen von f an der Geraden $y = x$, falls beide Koordinatenachsen linear und gleich skaliert sind. $\square$

BEISPIEL 36. Die lineare Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = ax + b$ mit $a, b \in \mathbb{R}$ und $a \neq 0$ ist injektiv auf $\mathbb{R}$, denn die Gleichung $y = ax + b$ lässt sich nach x umformen $x = (y - b)/a$, da $a \neq 0$. Die Umkehrfunktion ist $g(y) = (y - b)/a$.

BEISPIEL 37. Die quadratische Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ ist nicht injektiv auf $\mathbb{R}$, denn $f(-x) = (-x)^2 = x^2 = f(x)$ für alle $x \in \mathbb{R}$ und $-x \neq x$ für alle $x \neq 0$. Jedoch ist $f(x) = x^2$ injektiv auf $D = [0, \infty)$. Es gilt $f(D) = [0, \infty)$. Die Umkehrfunktion ist $g : [0, \infty) \rightarrow \mathbb{R}$ mit $g(y) = \sqrt{y}$. Die Abbildung 10 zeigt den Graph der Funktion $f(x) = x^2$ für $x \geq 0$ und den Graph der Umkehrfunktion $g(x) = \sqrt{x}$ für $x \geq 0$ und die Spiegelgerade $y = x$.

BEISPIEL 38. Die Funktion $f : (0, \infty) \rightarrow \mathbb{R}$ mit $f(x) = 1/x$ ist injektiv auf $D = (0, \infty)$, denn die Gleichung $y = 1/x$ lässt sich nach x umformen $x = 1/y$, da $y \in f(D) = (0, \infty)$. Die Umkehrfunktion ist $g(y) = 1/y$.

BEISPIEL 39 (Definition $\arccos$). $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \cos x$ ist, wie jede periodische Funktion, auf $\mathbb{R}$ nicht injektiv, weil $\cos(x + 2\pi) = \cos x$ für alle $x \in \mathbb{R}$ und $x \neq x + 2\pi$.

Aber die Funktion $\cos x$ ist auf dem Intervall $[0, \pi]$ injektiv, denn wenn $\cos x_1 = \cos x_2$ für $0 \leq x_1 \leq x_2 \leq \pi$, dann gilt $\sin x_1 = \sin x_2$, da $\sin x \geq 0$ für alle $x \in [0, \pi]$, und $\sin(x_2 - x_1) = \sin x_1 \cos x_2 - \sin x_2 \cos x_1 = 0$, also $x_2 = x_1$, weil die $x_2 = \pi$ und $x_1 = 0$ wegen $\cos 0 \neq \cos \pi$ ausscheidet.

Die Umkehrfunktion zu $\cos : [0, \pi] \rightarrow \mathbb{R}$ wird $\arccos$ genannt. Ihr Definitionsbereich ist das abgeschlossene Intervall $[-1, 1]$. Da $\arccos$ als Umkehrfunktion definiert ist, gilt $\arccos y \in [0, \pi]$ für alle $y \in [-1, 1]$. Die Abbildung 9 zeigt den Graph der Funktion $f(x) = \cos x$ für $x \in [0, \pi]$ und den Graph der Umkehrfunktion $g(x) = \arccos x$ für $x \in [-1, 1]$ und die Spiegelgerade $y = x$.

6. Polarkoordinaten der komplexen Zahlen

Für $\phi \in \mathbb{R}$ definieren wir

$$(38) \quad e^{i\phi} := \cos \phi + i \sin \phi.$$

Jeder reellen Zahl ϕ wird also eine komplexe Zahl zugeordnet, nämlich $\cos \phi + i \sin \phi$.

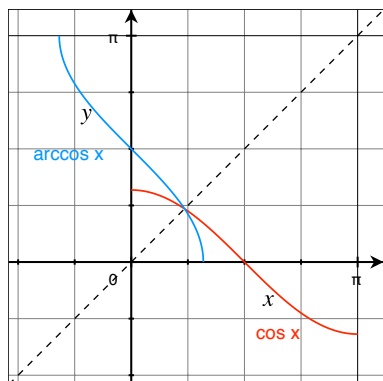


ABBILDUNG 9. $\cos$ und Umkehrfunktion $\arccos$

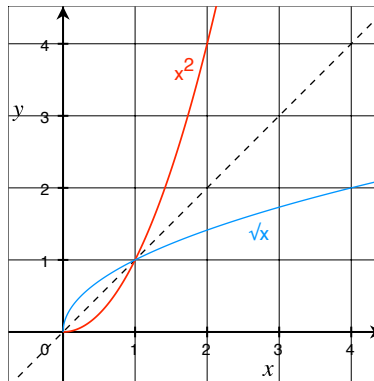


ABBILDUNG 10. $\sqrt{x}$ als Umkehrfunktion von x^2

BEMERKUNG 34. Die Zuordnung $\phi \mapsto \cos \phi + i \sin \phi$ ist ein Beispiel für eine komplexwertige Funktion einer reellen Variablen. $\square$

SATZ 33. Für alle $\phi \in \mathbb{R}$ gilt $|e^{i\phi}| = 1$.

BEWEIS. Für eine komplexe Zahl $z = x + iy$ mit $x, y \in \mathbb{R}$ gilt $|z|^2 = x^2 + y^2$. Nun folgt aus Satz 27

$$|e^{i\phi}|^2 = |\cos \phi + i \sin \phi|^2 = \cos^2 \phi + \sin^2 \phi = 1.$$

$\square$

SATZ 34. Für alle $k \in \mathbb{Z}$ und alle $\phi \in \mathbb{R}$ gilt

$$(39) \quad e^{i\phi} = e^{i(\phi+2k\pi)}$$

BEWEIS. Da $e^{i(\phi+2k\pi)} = \cos(\phi+2k\pi) + i \sin(\phi+2k\pi)$ folgt die Behauptung, weil $\sin$ und $\cos$ beide 2π -periodische Funktionen sind. $\square$

SATZ 35. Für alle $r, \phi \in \mathbb{R}$ gilt $\overline{re^{i\phi}} = re^{-i\phi}$.

BEWEIS. Es gilt

$$\overline{re^{i\phi}} = \overline{r \cos \phi + ir \sin \phi} = r \cos \phi - ir \sin \phi = r \cos(-\phi) + ir \sin(-\phi) = re^{-i\phi},$$

da $\cos$ eine gerade und $\sin$ eine ungerade Funktion ist. $\square$

6.1. Die komplexe Zahl $e^{i\phi}$ für spezielle Winkel. Aus den speziellen Funktionswerten $\sin(k\pi) = 0$, $\sin(\pi/2 + 2k\pi) = 1$, $\sin(3\pi/2 + 2k\pi) = -1$ für $k \in \mathbb{Z}$ und $\cos(x) = \sin(x + \pi/2)$ für alle $x \in \mathbb{R}$ folgt

$$e^{i \cdot 0} = \cos 0 + i \sin 0 = 1$$

$$e^{i\pi/2} = \cos(\pi/2) + i \sin(\pi/2) = 0 + i \cdot 1 = i$$

$$e^{i\pi} = \cos(\pi) + i \sin(\pi) = -1 + i \cdot 0 = -1$$

$$e^{i3\pi/2} = \cos(3\pi/2) + i \sin(3\pi/2) = 0 + i \cdot (-1) = -i.$$

Aus dem Additionstheorem für den doppelten Winkel $\cos(2x) = 2\cos^2(x) - 1$ folgt mit $x = \pi/4$

$$0 = \cos(\pi/2) = 2\cos^2(\pi/4) - 1, \text{ also } \cos(\pi/4) = \frac{1}{\sqrt{2}} = \frac{1}{2}\sqrt{2},$$

weil $\cos x > 0$ für alle $x \in [0, \pi/2)$. Nun ergibt sich aus dem Additionstheorem $\sin(2x) = 2\sin x \cos x$ für $x = \pi/4$

$$1 = \sin(\pi/2) = 2\sin(\pi/4)\cos(\pi/4), \text{ also } \sin(\pi/4) = \frac{1}{2\cos(\pi/4)} = \frac{1}{2}\sqrt{2}.$$

Zusammen erhalten wir mit Hilfe der Beziehungen $\sin(\pi - x) = \sin x$, $\sin(\pi + x) = -\sin x$, $\sin(2\pi - x) = \sin(\pi + (\pi - x)) = -\sin(\pi - x) = -\sin x$

$$\begin{aligned} e^{i\pi/4} &= \frac{1}{2}\sqrt{2}(1+i) \\ e^{i3\pi/4} &= \frac{1}{2}\sqrt{2}(-1+i) \\ e^{i5\pi/4} &= \frac{1}{2}\sqrt{2}(-1-i) = e^{-i3\pi/4} \\ e^{i7\pi/4} &= \frac{1}{2}\sqrt{2}(1-i) = e^{-i\pi/4}. \end{aligned}$$

LEMMA 10 (Additionstheorem für den dreifachen Winkel). *Für alle $x \in \mathbb{R}$ gilt*

$$(40) \quad \sin(3x) = 3\sin(x) - 4\sin^3(x).$$

BEWEIS. Aus den Additionstheoremen und $\sin^2 x + \cos^2 x = 1$ folgt

$$\begin{aligned} \sin(3x) &= \sin(x+2x) = \sin(x)\cos(2x) + \sin(2x)\cos(x) \\ &= \sin(x)(\cos^2(x) - \sin^2(x)) + 2\sin(x)\cos(x)\cos(x) \\ &= -\sin^3(x) + 3\sin(x)\cos^2(x) \\ &= -\sin^3(x) + 3\sin(x)(1 - \sin^2(x)) = 3\sin(x) - 4\sin^3(x). \end{aligned}$$

□

Mit $x = \pi/3$ erhalten wir aus dem Additionstheorem für den dreifachen Winkel

$$0 = \sin \pi = \sin(\pi/3)(3 - 4\sin^2(\pi/3)), \text{ also } \sin(\pi/3) = \frac{\sqrt{3}}{2},$$

weil $\sin x > 0$ für alle $x \in (0, \pi)$. Es folgt $\cos^2(\pi/3) = 1 - \sin^2(\pi/3) = 1/4$, also $\cos(\pi/3) = 1/2$, da $\cos x > 0$ für alle $x \in [0, \pi/2)$. Aus $\sin(\pi/2 - x) = \cos x$ folgt wegen $\pi/2 = \pi/3 + \pi/6$ sofort $\sin(\pi/6) = 1/2$ und $\cos(\pi/6) = \sqrt{3}/2$. Durch Ausnutzung der Symmetrien der Winkelfunktionen erhalten

wir

$$\begin{aligned} e^{i\pi/6} &= \frac{\sqrt{3}}{2} + i\frac{1}{2} & e^{i\pi/3} &= \frac{1}{2} + i\frac{\sqrt{3}}{2} \\ e^{i5\pi/6} &= -\frac{\sqrt{3}}{2} + i\frac{1}{2} & e^{i2\pi/3} &= -\frac{1}{2} + i\frac{\sqrt{3}}{2} \\ e^{i7\pi/6} &= -\frac{\sqrt{3}}{2} - i\frac{1}{2} & e^{i4\pi/3} &= -\frac{1}{2} - i\frac{\sqrt{3}}{2} \\ e^{i11\pi/6} &= \frac{\sqrt{3}}{2} - i\frac{1}{2} & e^{i5\pi/3} &= \frac{1}{2} - i\frac{\sqrt{3}}{2}. \end{aligned}$$

BEMERKUNG 35. Man kann die Sinus- und Kosinuswerte für spezielle Winkel auch durch elementargeometrische Überlegungen an gleichseitigen bzw. gleichschenkligen Dreiecken herleiten. $\square$

BEMERKUNG 36. Möchte man die Werte $\sin x$ und $\cos x$ für beliebige $x \in \mathbb{R}$ (näherungsweise) bestimmen, so kann man die Taylorreihen (siehe Abschnitt 2)

$$\sin x = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}, \quad \cos x = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n}$$

benutzen oder zwischen tabellierten Werten interpolieren (siehe Abschnitt 4). $\square$

6.2. Funktionalgleichung.

SATZ 36. Für alle $\phi_1, \phi_2 \in \mathbb{R}$ gilt

$$(41) \quad e^{i(\phi_1+\phi_2)} = e^{i\phi_1} e^{i\phi_2}.$$

BEWEIS. Aus der Definition von $e^{i\phi}$, der Multiplikation komplexer Zahlen und den Additionstheoremen für die Winkelfunktionen folgt

$$\begin{aligned} e^{i\phi_1} e^{i\phi_2} &= (\cos \phi_1 + i \sin \phi_1)(\cos \phi_2 + i \sin \phi_2) \\ &= \cos \phi_1 \cos \phi_2 + i \sin \phi_2 \cos \phi_1 + i \sin \phi_1 \cos \phi_2 - \sin \phi_1 \sin \phi_2 \\ &= (\cos \phi_1 \cos \phi_2 - \sin \phi_1 \sin \phi_2) + i(\sin \phi_2 \cos \phi_1 + \sin \phi_1 \cos \phi_2) \\ &= \cos(\phi_1 + \phi_2) + i \sin(\phi_1 + \phi_2) \\ &= e^{i(\phi_1+\phi_2)}. \end{aligned}$$

$\square$

Für $z \in \mathbb{C}$ mit $z = x + iy$ mit $x, y \in \mathbb{R}$ definieren wir

$$(42) \quad e^z = e^{x+iy} := e^x e^{iy}.$$

BEMERKUNG 37. Diese Definition der Exponentialfunktion für komplexe Argumente ist eine natürliche Fortsetzung der Exponentialfunktion für reelle Argumente. Mit dieser Definition gilt die Funktionalgleichung auch

für komplexe Zahlen $z_1, z_2 \in \mathbb{C}$, denn falls $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$, so gilt

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

6.3. Existenz und Eindeutigkeit der Polarkoordinaten. Es sei $z \in \mathbb{C}$. Wenn $z = re^{i\phi}$ für ein $r \in \mathbb{R}$ mit $r \geq 0$ und ein $\phi \in \mathbb{R}$ gilt, so heißt $z = re^{i\phi}$ eine **Polarkoordinatendarstellung** der komplexen Zahl z . Die Zerlegung in Real- und Imaginärteil ist dann $z = re^{i\phi} = r \cos \phi + ir \sin \phi$. Da $e^{i\phi} \neq 0$ für alle $\phi \in \mathbb{R}$, gilt

$$(43) \quad re^{i\phi} = 0 \Leftrightarrow r = 0 \quad \forall r, \phi \in \mathbb{R}.$$

LEMMA 11. Für alle $r_1, r_2, \phi_1, \phi_2 \in \mathbb{R}$ mit $r_1, r_2 > 0$ gilt

$$(44) \quad r_1 e^{i\phi_1} = r_2 e^{i\phi_2} \Leftrightarrow r_1 = r_2 \text{ und } \phi_2 = \phi_1 + 2k\pi \text{ für ein } k \in \mathbb{Z}.$$

BEWEIS. Es reicht $\Rightarrow$ zu zeigen. Da $|re^{i\phi}| = |r| = r$, folgt aus $r_1 e^{i\phi_1} = r_2 e^{i\phi_2}$ sofort $r_1 = r_2$. Wenn nun $e^{i\phi_1} = e^{i\phi_2}$, so gilt $\cos \phi_1 = \cos \phi_2$ und $\sin \phi_1 = \sin \phi_2$. Daraus folgt $\sin(\phi_1 - \phi_2) = \sin \phi_1 \cos \phi_2 - \sin \phi_2 \cos \phi_1 = 0$, also $\phi_1 - \phi_2 = k\pi$ für ein $k \in \mathbb{Z}$. Nun folgt aus dem Additionstheorem mit $\sin(k\pi) = 0$ und $\cos(k\pi) = (-1)^k$ für alle $k \in \mathbb{Z}$ sofort $\cos \phi_1 = \cos(\phi_2 + k\pi) = \cos \phi_2 \cos(k\pi) = (-1)^k \cos \phi_2$. Da die Bedingung aber $\cos \phi_1 = \cos \phi_2$ lautet, muss k gerade sein. $\square$

SATZ 37. Für jede komplexe Zahl $z \in \mathbb{C}$ existieren $r, \phi \in \mathbb{R}$ mit $r \geq 0$ und $z = re^{i\phi}$. Falls $z \neq 0$, so gilt $z = re^{i\phi}$ für

$$(45) \quad r := |z| \text{ und } \phi := \begin{cases} \arccos \frac{x}{r} & (y \geq 0) \\ -\arccos \frac{x}{r} & (y < 0) \end{cases} \text{ bzw. } \begin{cases} \arctan \frac{y}{x} & (x > 0) \\ \pi + \arctan \frac{y}{x} & (x < 0) \\ \frac{1}{2}\pi & (x = 0, y > 0) \\ \frac{3}{2}\pi & (x = 0, y < 0) \end{cases}.$$

Die Polarkoordinatendarstellung einer komplexen Zahl $z \neq 0$ ist eindeutig, wenn man $\phi \in [0, 2\pi)$ fordert.

BEWEIS. Wenn $z = 0$, so gilt $z = re^{i\phi}$ für $r = 0$ und beliebiges $\phi \in \mathbb{R}$. Wenn $z = x + iy \neq 0$, so gilt $r = |z| > 0$,

$$\frac{z}{r} = \frac{x}{r} + i\frac{y}{r}, \quad \left| \frac{z}{r} \right| = \frac{|z|}{r} = 1, \text{ also } \frac{x}{r}, \frac{y}{r} \in [-1, 1] \text{ und } \left(\frac{x}{r} \right)^2 + \left(\frac{y}{r} \right)^2 = 1.$$

Die Funktion $\arccos$ ist als Umkehrfunktion (siehe Abschnitt 5) des $\cos$ auf dem Intervall $[-1, 1]$ definiert und liefert Werte im Intervall $[0, \pi]$. Wenn $\phi = \pm \arccos(x/r)$, dann gilt

$$\sin \phi = \pm \sqrt{1 - \cos^2 \phi} = \pm \sqrt{1 - (x/r)^2} = \pm \sqrt{(y/r)^2} = \pm \frac{y}{r}.$$

Wenn $\phi = \arccos(x/r)$, dann $\sin \phi \geq 0$, da $\phi \in [0, \pi]$. Falls $y \geq 0$, so gilt also $\sin(\arccos(x/r)) = y/r$.

Wenn $\phi = -\arccos(x/r)$, dann $\sin \phi \leq 0$, da $\phi \in [-\pi, 0]$. Falls $y < 0$, so gilt also $\sin(-\arccos(x/r)) = y/r$. $\square$

6.4. Multiplikation und Potenzen in Polarkoordinaten.

SATZ 38. Für alle $z_1, z_2, z \in \mathbb{C}$ mit $z_1 = r_1 e^{i\phi_1}$, $z_2 = r_2 e^{i\phi_2}$, $z = r e^{i\phi}$ und $n \in \mathbb{Z}$ gilt

$$(46) \quad z_1 z_2 = r_1 r_2 e^{i(\phi_1 + \phi_2)} \text{ und } z^n = r^n e^{in\phi}.$$

BEWEIS. In Polarkoordinaten folgt aus der Funktionalgleichung

$$z_1 z_2 = r_1 e^{i\phi_1} r_2 e^{i\phi_2} = r_1 r_2 e^{i\phi_1} e^{i\phi_2} = r_1 r_2 e^{i(\phi_1 + \phi_2)}.$$

Mit vollständiger Induktion folgt für $z = z_1 = z_2$ die Behauptung $z^n = r^n e^{in\phi}$ für alle $n \in \mathbb{N}$. Da $r^{-1} e^{-i\phi} r e^{i\phi} = 1$ für alle $r > 0$ gilt, folgt $z^{-1} = r^{-1} e^{-i\phi}$ für alle $z \neq 0$. Wiederum folgt mit vollständiger Induktion die Behauptung $z^{-n} = r^{-n} e^{-in\phi}$ für alle $n \in \mathbb{N}$. $\square$

6.5. n -te Einheitswurzeln. Für $n \in \mathbb{N}$, $n > 0$ wird eine komplexe Zahl $z \in \mathbb{C}$ mit der Eigenschaft $z^n = 1$ eine n -te **Einheitswurzel** genannt. Der Fundamentalsatz der Algebra besagt, dass das Polynom $z^n - 1$ genau n Nullstellen besitzt. Eine n -te Einheitswurzel ist $z_0 = 1$, denn $z_0^n = 1$ für alle $n \in \mathbb{N}$. Wir bestimmen alle n -ten Einheitswurzeln.

SATZ 39. Die n -ten Einheitswurzeln sind

$$(47) \quad z_k = e^{ik \frac{2\pi}{n}} \text{ mit } k = 0, 1, \dots, n-1.$$

BEWEIS. Wir betrachten z in Polarkoordinaten, also $z = r e^{i\phi}$ mit $r, \phi \in \mathbb{R}$ und $\phi \in [0, 2\pi)$. Da $0^n \neq 1$, gilt $r > 0$. Nun folgt

$$\begin{aligned} z^n = 1 &\Leftrightarrow r^n e^{in\phi} = 1 \cdot e^{i \cdot 0} \\ &\Leftrightarrow r^n = 1 \text{ und } n\phi = 2k\pi \text{ für ein } k \in \mathbb{Z} \\ &\Leftrightarrow r = 1 \text{ und } \phi = k \frac{2\pi}{n} \text{ für ein } k \in \mathbb{Z}, \end{aligned}$$

da $r \in \mathbb{R}$ und $r > 0$. Aus der Bedingung $\phi \in [0, 2\pi)$ folgt $k \in \{0, 1, \dots, n-1\}$. $\square$

BEISPIEL 40. Die komplexen Lösungen der Gleichung $z^3 = 1$ sind

$$z_0 = e^{i \cdot 0 \frac{2\pi}{3}} = 1, \quad z_1 = e^{i \frac{2\pi}{3}} = -\frac{1}{2} + i \frac{\sqrt{3}}{2}, \quad z_2 = e^{i 2 \frac{2\pi}{3}} = e^{i \frac{4\pi}{3}} = -\frac{1}{2} - i \frac{\sqrt{3}}{2}.$$

Abbildung 11 zeigt die dritten Einheitswurzeln.

6.6. n -te Wurzeln. Für jedes $a \in \mathbb{C}$ mit $a \neq 0$ und alle $n \in \mathbb{N}$ mit $n > 0$ besitzt die Gleichung $z^n = a$ genau n verschiedene komplexe Lösungen. Diese lassen sich einfach mit Hilfe einer Lösung der Gleichung $z^n = a$ und den n -ten Einheitswurzeln konstruieren.

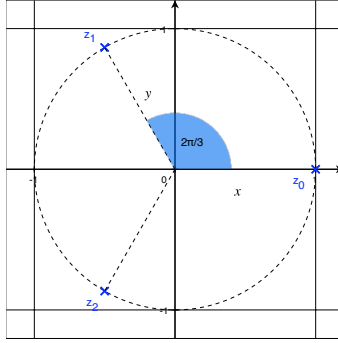


ABBILDUNG 11. Die dritten Einheitswurzeln

SATZ 40. Wenn $a = r_0 e^{i\phi_0} \in \mathbb{C}$ mit $r_0 > 0$ und $n \in \mathbb{N}$ mit $n > 0$, so hat die Gleichung $z^n = a$ genau die n Lösungen

$$(48) \quad z_k = \sqrt[n]{r_0} e^{i\left(\frac{\phi_0}{n} + k \frac{2\pi}{n}\right)} \text{ mit } k = 0, 1, \dots, n-1.$$

BEWEIS. Es gilt

$$\left(\sqrt[n]{r_0} e^{i\frac{\phi_0}{n}} \right)^n = r_0 e^{i\phi_0} = a.$$

Wenn z eine Lösung der Gleichung $z^n = a$ ist und $b := \sqrt[n]{r_0} e^{i\frac{\phi_0}{n}}$, dann ist z/b eine n -te Einheitswurzel, denn

$$\left(\frac{z}{b} \right)^n = \frac{z^n}{b^n} = \frac{a}{a} = 1.$$

Aus Satz 39 folgt

$$z = b e^{ik \frac{2\pi}{n}} = \sqrt[n]{r_0} e^{i\frac{\phi_0}{n}} e^{ik \frac{2\pi}{n}} = \sqrt[n]{r_0} e^{i\left(\frac{\phi_0}{n} + k \frac{2\pi}{n}\right)}$$

für $k = 0, 1, \dots, n-1$. □

BEISPIEL 41. Da $-8 = 8e^{i\pi}$, ergeben sich die komplexen Lösungen der Gleichung $z^3 = -8$ mit $n = 3$, $r_0 = 8$ und $\phi_0 = \pi$:

$$z_0 = \sqrt[3]{8} e^{i\frac{\pi}{3}} = 2 \left(\frac{1}{2} + i \frac{\sqrt{3}}{2} \right) = 1 + i\sqrt{3}$$

$$z_1 = \sqrt[3]{8} e^{i\left(\frac{\pi}{3} + \frac{2\pi}{3}\right)} = 2e^{i\pi} = -2$$

$$z_2 = \sqrt[3]{8} e^{i\left(\frac{\pi}{3} + \frac{4\pi}{3}\right)} = 2e^{i\frac{5\pi}{3}} = 2 \left(\frac{1}{2} - i \frac{\sqrt{3}}{2} \right) = 1 - i\sqrt{3}.$$

Abbildung 12 zeigt die Lösungen der Gleichung $z^3 = -8$.

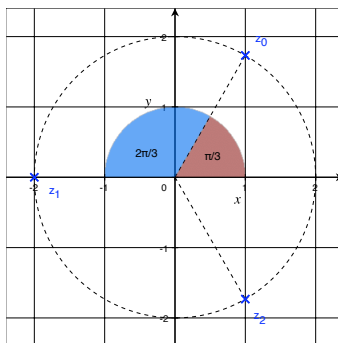


ABBILDUNG 12. Lösungen der Gleichung $z^3 = -8$

KAPITEL VI

Stetigkeit

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt **stetig** im Punkt $x_0 \in D$, falls für jedes $\varepsilon > 0$ ein $\delta > 0$ existiert, so dass $|f(x) - f(x_0)| < \varepsilon$ für alle $x \in D$ mit $|x - x_0| < \delta$ gilt. Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt **stetig** auf D , falls f in allen Punkten $x \in D$ stetig ist.

BEMERKUNG 38. Anschaulich ist eine Funktion stetig, wenn man ihren Graph zeichnen kann, ohne den Stift abzusetzen. Aber es gibt auch andere Unstetigkeitsphänomene (siehe Beispiel 51).

BEISPIEL 42 (Stetigkeit der konstanten Funktion). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f \equiv c$ ist auf $\mathbb{R}$ stetig, denn für jeden Punkt x_0 und alle $\varepsilon > 0$ gilt $|f(x) - f(x_0)| = |c - c| = 0 < \varepsilon$ für alle $x \in \mathbb{R}$, also $\delta = \infty$.

BEISPIEL 43 (Stetigkeit der Funktion $f(x) = x$). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x$ ist auf $\mathbb{R}$ stetig, denn für jeden Punkt x_0 und $\varepsilon > 0$ gilt $|f(x) - f(x_0)| = |x - x_0| < \varepsilon$ für alle $x \in \mathbb{R}$ mit $|x - x_0| < \varepsilon$, also $\delta = \varepsilon$.

BEISPIEL 44 (Unstetigkeit der Sprungfunktion). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

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

ist stetig in allen Punkten $x_0 \neq 0$, denn $|f(x) - f(x_0)| = 0$ für alle x mit $|x - x_0| < |x_0| = \delta$. Die Funktion ist in $x_0 = 0$ nicht stetig, denn für $\varepsilon = 1/2$ findet man kein $\delta > 0$ mit $|f(x) - f(0)| < 1/2$ für alle $|x| < \delta$, weil für alle $x < 0$ gilt $|f(x) - 1| = 1$.

1. Folgen

Eine Folge $\{x_n\}$ reeller Zahlen ist eine Funktion $\mathbb{N} \rightarrow \mathbb{R}$, $n \mapsto x_n$, zum Beispiel $\{n^2\}$ oder $\{1/n\}$ oder $\{3^n\}$.

Eine Folge $\{x_n\}$ heißt **konvergent**, falls ein $c \in \mathbb{R}$ existiert, so dass für alle $\varepsilon > 0$ ein Index $n_0 \in \mathbb{N}$ mit $|x_n - c| < \varepsilon$ für alle $n \geq n_0$ existiert. Die Zahl c heißt **Grenzwert** der Folge. Man schreibt $\lim_{n \rightarrow \infty} x_n = c$. Eine konvergente Folge mit $\lim_{n \rightarrow \infty} x_n = 0$ heißt **Nullfolge**.

Eine Folge ist genau dann konvergent, wenn zu jedem $\varepsilon > 0$ ein $n_0 \in \mathbb{N}$ mit $|x_n - x_m| < \varepsilon$ für alle $n, m \geq n_0$ existiert. Eine Folge konvergiert genau dann gegen den Grenzwert c , wenn die Folge $\{x_n - c\}$ eine Nullfolge ist.

BEISPIEL 45. Die Folge $\{x_n\}$ mit $x_n = 1/n$ für $n > 0$ ist eine Nullfolge, denn für jedes $\varepsilon > 0$ gilt $|x_n - 0| = |x_n| = 1/n < \varepsilon$ genau dann, wenn $n > 1/\varepsilon$. Also kann man n_0 als kleinste natürliche Zahl, die größer als $1/\varepsilon$ ist, wählen.

BEISPIEL 46. Die konstante Folge $\{x_n\}$ mit $x_n = c$ für alle $n \in \mathbb{N}$ konvergiert gegen c , denn $|x_n - c| = |c - c| = 0 < \varepsilon$ für alle $\varepsilon > 0$ und alle $n \in \mathbb{N}$.

SATZ 41. *Jede konvergente Folge ist beschränkt.*

BEWEIS. Wenn $\lim_{n \rightarrow \infty} x_n = c$, dann existiert ein $n_0 \in \mathbb{N}$ mit $|x_n - c| < 1$ für alle $n \geq n_0$. Aus der Dreiecksungleichung folgt $|x_n| = |x_n - c + c| \leq |x_n - c| + |c| < c + 1$ für alle $n \geq n_0$. Die endlich vielen Folgenglieder $x_0, \dots, x_{n_0-1}$ sind beschränkt, also $|x_n| < S$ für alle $n < n_0$. Daraus folgt $|x_n| < \max(S, |c| + 1)$ für alle $n \in \mathbb{N}$. $\square$

BEISPIEL 47. Für alle $k \in \mathbb{N}$ mit $k > 0$ und alle $a \in \mathbb{R}$ mit $a > 1$ konvergieren die Folgen $\{n^k\}$ und $\{a^n\}$ nicht, denn sie sind nicht beschränkt (siehe Abschnitt 8).

SATZ 42 (Rechenregeln für konvergente Folgen). *Wenn $\{x_n\}$ und $\{y_n\}$ konvergente Folgen sind, so sind auch die Folgen $\{x_n + y_n\}$ und $\{x_n y_n\}$ konvergent und für die Grenzwerte gilt*

$$(49) \quad \lim_{n \rightarrow \infty} (x_n + y_n) = \left(\lim_{n \rightarrow \infty} x_n \right) + \left(\lim_{n \rightarrow \infty} y_n \right)$$

und

$$(50) \quad \lim_{n \rightarrow \infty} (x_n y_n) = \left(\lim_{n \rightarrow \infty} x_n \right) \left(\lim_{n \rightarrow \infty} y_n \right).$$

Falls $\lim_{n \rightarrow \infty} y_n \neq 0$, so konvergiert auch die Folge $\{x_n/y_n\}$ und es gilt

$$(51) \quad \lim_{n \rightarrow \infty} \left(\frac{x_n}{y_n} \right) = \frac{\lim_{n \rightarrow \infty} x_n}{\lim_{n \rightarrow \infty} y_n}.$$

BEWEIS. Mit den Bezeichnungen $\lim_{n \rightarrow \infty} x_n = a$ und $\lim_{n \rightarrow \infty} y_n = b$ gilt

$$|(x_n + y_n) - (a + b)| = |x_n - a + y_n - b| \leq |x_n - a| + |y_n - b|.$$

Zu gegebenem $\varepsilon > 0$ existieren $n_1, n_2 \in \mathbb{N}$ mit $|x_n - a| < \varepsilon/2$ für alle $n \geq n_1$ und $|y_n - b| < \varepsilon/2$ für alle $n \geq n_2$. Daraus folgt $|(x_n + y_n) - (a + b)| < \varepsilon$ für alle $n \geq n_0 := \max(n_1, n_2)$.

Für die Produktfolge $\{x_n y_n\}$ betrachten wir

$$|x_n y_n - ab| = |x_n y_n - a y_n + a y_n - ab| \leq |y_n| |x_n - a| + |a| |y_n - b|.$$

Da $\{y_n\}$ konvergiert, ist $\{y_n\}$ beschränkt, $|y_n| \leq S$ für alle $n \in \mathbb{N}$. Zu gegebenem $\varepsilon > 0$ existieren $n_1, n_2 \in \mathbb{N}$ mit $|x_n - a| < \varepsilon/(2S)$ für alle $n \geq n_1$ und $|y_n - b| < \varepsilon/(2|a|)$ für alle $n \geq n_2$. Daraus folgt $|x_n y_n - ab| < \varepsilon$ für alle $n \geq n_0 := \max(n_1, n_2)$.

Für die Quotientenfolge $\{x_n y_n\}$ betrachten wir

$$\left| \frac{x_n}{y_n} - \frac{a}{b} \right| = \left| \frac{x_n b - y_n a}{y_n b} \right| = \left| \frac{x_n b - ab - y_n a + ab}{y_n b} \right| \leq \frac{|x_n - a|}{|b|} + \frac{|y_n - b| |a|}{|b|^2},$$

da $|y_n| \geq S$ für alle $n \geq n_1$. Zu gegebenem $\varepsilon > 0$ existieren $n_2, n_3 \in \mathbb{N}$ mit $|x_n - a| < S\varepsilon/2$ für alle $n \geq n_1$ und $|y_n - b| < S|b|\varepsilon/(2|a|)$ für alle $n \geq n_2$. Daraus folgt $|x_n/y_n - a/b| < \varepsilon$ für alle $n \geq n_0 := \max(n_1, n_2, n_3)$. $\square$

BEISPIEL 48. Für alle $k \in \mathbb{N}$ mit $k > 0$ ist die Folge $x_n = 1/n^k = (\frac{1}{n})^k$ eine Nullfolge.

LEMMA 12 (Sandwich-Lemma). Wenn für die Folgen $\{a_n\}$, $\{x_n\}$ und $\{A_n\}$ ein $n_0 \in \mathbb{N}$ existiert, so dass $a_n \leq x_n \leq A_n$ für alle $n \geq n_0$ und $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} A_n = c$, dann ist c auch Grenzwert der Folge $\{x_n\}$.

BEWEIS. Für jedes $\varepsilon > 0$ existiert ein $n_1 \in \mathbb{N}$ mit $|a_n - c| < \varepsilon$ für alle $n \geq n_1$ und ein $n_2 \in \mathbb{N}$ mit $|A_n - c| < \varepsilon$ für alle $n \geq n_2$. Daraus folgt

$$|x_n - c| \leq \max(|A_n - c|, |a_n - c|) < \varepsilon$$

für alle $n \geq \max(n_1, n_2)$. $\square$

BEISPIEL 49. Die Folge $\{x_n\}$ mit $x_n = 1/(n^2 + 3n)$ ist eine Nullfolge, denn

$$\frac{1}{n^2 + 3n} = \frac{1}{n^2} \cdot \frac{1}{1 + \frac{3}{n}}, \quad \lim_{n \rightarrow \infty} \frac{1}{n^2} = 0, \quad \lim_{n \rightarrow \infty} \frac{3}{n} = 0, \quad \lim_{n \rightarrow \infty} \frac{1}{1 + \frac{3}{n}} = \frac{1}{1 + 0} = 1.$$

Alternativ kann man mit Lemma 12 argumentieren, weil $3n^2 > n^2 + 3n > n^2$ für $n \geq 3$. Damit

$$\frac{1}{3n^2} < x_n < \frac{1}{n^2},$$

wobei $1/(3n^2)$ und $1/n^2$ Nullfolgen sind.

LEMMA 13. Wenn $\{x_n\}$ eine Nullfolge und $\{y_n\}$ beschränkt ist, dann ist auch $\{x_n y_n\}$ eine Nullfolge.

BEWEIS. Wenn $|y_n| \leq S$ für alle $n \in \mathbb{N}$, dann $|x_n y_n| \leq |x_n| S$. Zu gegebenem $\varepsilon > 0$ existiert $n_0 \in \mathbb{N}$ mit $|x_n| < \varepsilon/S$ für alle $n \geq n_0$. Daraus folgt $|x_n y_n| < \varepsilon$ für alle $n \geq n_0$. $\square$

2. Folgenstetigkeit

SATZ 43. Eine Funktion $f : D \rightarrow \mathbb{R}$ ist genau dann in $x_0 \in D$ stetig, wenn

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

für alle Folgen $\{x_n\} \subset D$ mit $\lim_{n \rightarrow \infty} x_n = x_0$ gilt.

BEWEIS. Wenn f in x_0 stetig ist und die Folge $\{x_n\} \subset D$ gegen x_0 konvergiert, dann existiert zu jedem $\varepsilon > 0$ ein $\delta > 0$, so dass $|f(x) - f(x_0)| < \varepsilon$ für alle x mit $|x - x_0| < \delta$. Da $\lim_{n \rightarrow \infty} x_n = x_0$, so existiert ein Index $n_0 \in \mathbb{N}$ mit $|x_n - x_0| < \delta$ für alle $n \geq n_0$. Also gilt für alle $n \geq n_0$ auch $|f(x_n) - f(x_0)| < \varepsilon$. D.h. $\lim_{n \rightarrow \infty} f(x_n) = f(x_0)$.

Wenn für jede Folge mit $\lim_{n \rightarrow \infty} x_n = x_0$ auch $\lim_{n \rightarrow \infty} f(x_n) = f(x_0)$ gilt und $\varepsilon > 0$, dann muss es ein $\delta > 0$ mit $|f(x) - f(x_0)| < \varepsilon$ für alle $|x - x_0| < \delta$ geben. Anderenfalls befände sich für jedes $\delta = 1/n$ im Intervall $(x_0 - 1/n, x_0 + 1/n)$ einen Punkt x_n mit $|f(x_n) - f(x_0)| \geq \varepsilon$. Diese Punktfolge konvergiert gegen x_0 , aber $|f(x_n) - f(x_0)| \geq \varepsilon$ für alle $n \in \mathbb{N}$. $\square$

BEMERKUNG 39. Die Beschreibung der Stetigkeit mit Hilfe von Folgen ist besonders hilfreich, wenn man zeigen möchte, dass eine Funktion f in einem Punkt x_0 nicht stetig ist. Es genügt dann nämlich, eine gegen x_0 konvergente Folge $\{x_n\}$ mit $\lim_{n \rightarrow \infty} f(x_n) \neq f(x_0)$ anzugeben (siehe Beispiele 50 und 51). $\square$

BEISPIEL 50. Die Sprungfunktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = 1$, falls $x \geq 0$, und $f(x) = 0$, falls $x < 0$, ist in $x_0 = 0$ nicht stetig (siehe Beispiel 44). Für die Folge $\{x_n\}$ mit $x_n = -1/n$ gilt $\lim_{n \rightarrow \infty} x_n = x_0 = 0$ aber $\lim_{n \rightarrow \infty} f(x_n) = \lim_{n \rightarrow \infty} 0 = 0 \neq 1 = f(x_0)$, da $x_n = -1/n < 0$ für alle $n \in \mathbb{N}$.

BEISPIEL 51. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

ist in $x_0 = 0$ nicht stetig, denn für die Folge $\{x_n\}$ mit

$$x_n = \frac{1}{2n\pi + \pi/2} \text{ gilt } \lim_{n \rightarrow \infty} x_n = x_0 = 0$$

und

$$\lim_{n \rightarrow \infty} f(x_n) = \lim_{n \rightarrow \infty} f(\sin(2n\pi + \pi/2)) = 1 \neq 0 = f(x_0).$$

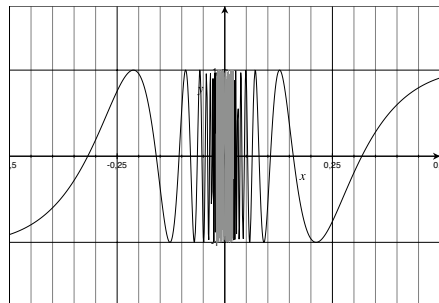
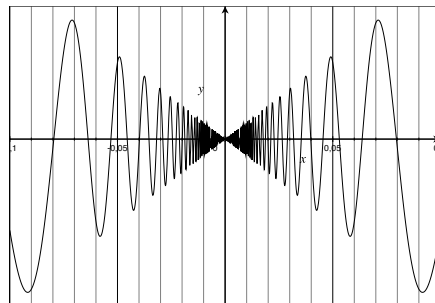
BEMERKUNG 40. Die Beispiele 50 und 51 sind typisch für Unstetigkeitsstellen.

In Beispiel 50 liegt in $x_0 = 0$ eine Sprungstelle vor, d.h. für alle Nullfolgen $\{x_n\}$ mit $x_n > 0$ für alle $n \in \mathbb{N}$ existieren die Grenzwerte $\lim_{n \rightarrow \infty} f(x_0 + x_n)$ und $\lim_{n \rightarrow \infty} f(x_0 - x_n)$. Sie sind jedoch verschieden voneinander oder vom Funktionswert $f(x_0)$.

Die Ursache für die Unstetigkeit im Beispiel 51 ist eine in der Nähe der Unstetigkeitsstelle immer schneller oszillierende Funktion (mit konstanter Amplitude). $\square$

BEISPIEL 52. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} x \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

ABBILDUNG 1. $\sin(1/x)$ ABBILDUNG 2. $x \sin(1/x)$

ist in $x_0 = 0$ stetig, denn für alle $x \neq 0$ gilt

$$|f(x) - f(0)| = |x \sin(1/x) - 0| = |x| |\sin(1/x)| \leq |x|.$$

Zu gegebenem $\varepsilon > 0$ gilt $|f(x) - f(0)| < \varepsilon$ für alle $x \in \mathbb{R}$ mit $|x| < \delta := \varepsilon$.

3. Eigenschaften stetiger Funktionen

Für eine Funktion $f : D \rightarrow \mathbb{R}$ und eine Menge $U \subset \mathbb{R}$ heißt die Menge $f^{-1}(U) := \{x \in D : f(x) \in U\}$ das **Urbild** von U unter f .

LEMMA 14. *Eine Funktion $f : D \rightarrow \mathbb{R}$ ist genau dann stetig auf D , wenn für jede offene Menge $U \subset \mathbb{R}$ das Urbild $f^{-1}(U)$ eine offene Menge in D ist.*

BEWEIS. Es seien $f : D \rightarrow \mathbb{R}$ eine stetige Funktion und $U \subset \mathbb{R}$ eine offene Menge. Für jedes $y_0 \in U$ existiert ein $\varepsilon > 0$ mit $\{y : |y - y_0| < \varepsilon\} \subset U$. Für jedes $x_0 \in D$ mit $f(x_0) = y_0$ existiert ein $\delta_{y_0, x_0} > 0$ mit $|f(x) - f(x_0)| < \varepsilon$ für alle $|x - x_0| < \delta_{y_0, x_0}$. Alle Mengen $V_{y_0, x_0} := \{x : |x - x_0| < \delta_{y_0, x_0}\}$ sind offene Intervalle. Es gilt $f^{-1}(U) := \cup_{y_0 \in U} V_{y_0, x_0}$. Darum ist $f^{-1}(U)$ offen.

Es sei $f : D \rightarrow \mathbb{R}$ eine Funktion mit der Eigenschaft $f^{-1}(U)$ ist offen für alle offenen Mengen $U \subset \mathbb{R}$. Wenn $y_0 \in f(D)$, dann ist $\{y : |y - y_0| < \varepsilon\} =: U_0$ eine offene Menge für jedes $\varepsilon > 0$. Daraus folgt, dass $f^{-1}(U_0)$ eine offene Menge ist. Da $x_0 \in f^{-1}(U_0)$ für alle $x_0 \in D$ mit $f(x_0) = y_0$, existiert zu jedem x_0 ein δ_{x_0} mit $f(x) \in U_0$ für alle x mit $|x - x_0| < \delta_{x_0}$. $\square$

SATZ 44 (Zwischenwertsatz). *Wenn eine Funktion f auf dem Intervall (a, b) stetig ist und $f(x_1) < f(x_2)$ für $x_1, x_2 \in (a, b)$ gilt, dann enthält das Bild von f das Intervall $[f(x_1), f(x_2)]$, also $[f(x_1), f(x_2)] \subset f([x_1, x_2]) \subset f((a, b))$.*

BEWEIS. Wenn $y_0 \in [f(x_1), f(x_2)]$ und $y_0 \notin f((a, b))$, dann sind die Intervalle $U_- = (-\infty, y_0)$ und $U_+ = (y_0, \infty)$ offene Mengen. Nun sind $f^{-1}(U_-)$ und $f^{-1}(U_+)$ offene Mengen mit $(a, b) = f^{-1}(U_-) \cup f^{-1}(U_+)$ und $f^{-1}(U_-) \cap f^{-1}(U_+) = \emptyset$. Dies bedeutet, dass (a, b) nicht zusammenhängend ist. Aber jedes Intervall ist zusammenhängend. $\square$

BEISPIEL 53. Für die Funktion $f : \{x \neq 0\} \rightarrow \mathbb{R}$ mit $f(x) = 1/x$ gilt $f(-1) = -1$ und $f(1) = 1$. Jedoch ist das Intervall $[-1, 1]$ nicht vollständig im Bild enthalten, denn $f(x) \neq 0$ für alle $x \neq 0$. Das Definitionsgebiet $\{x \neq 0\}$ ist nicht zusammenhängend, es besteht aus den zwei Komponenten $\{x < 0\}$ und $\{x > 0\}$.

SATZ 45 (Existenz von Extremwerten). *Wenn $f : [a, b] \rightarrow \mathbb{R}$ stetig, dann existieren $x_0, x_1 \in [a, b]$ mit $f(x_0) \leq f(x) \leq f(x_1)$ für alle $x \in [a, b]$.*

BEWEIS. Entscheidende Voraussetzung für den Beweis ist die Kompaktheit des abgeschlossenen Intervalls $[a, b]$. $\square$

BEISPIEL 54. Die Funktion $f : \{x > 0\} \rightarrow \mathbb{R}$, $f(x) = 1/x$, ist auf $\{x > 0\}$ stetig. Sie besitzt jedoch kein Maximum, denn $\lim_{x \rightarrow 0} f(x) = \infty$. Das Definitionsintervall $(0, \infty)$ ist nicht abgeschlossen.

Wenn $f(x_0) \leq f(x)$ für alle $x \in [a, b]$ für ein $x_0 \in [a, b]$, so heißt der Funktionswert $f(x_0)$ **Minimum** der Funktion f auf dem Intervall $[a, b]$. Wenn $f(x_1) \geq f(x)$ für alle $x \in [a, b]$ für ein $x_1 \in [a, b]$, so heißt der Funktionswert $f(x_1)$ **Maximum** der Funktion f auf dem Intervall $[a, b]$. Jedes $x \in [a, b]$ mit $f(x) = f(x_0)$ oder $f(x) = f(x_1)$ heißt **Extremstelle** der Funktion f auf dem Intervall $[a, b]$.

BEISPIEL 55. Die Funktion $f(x) = x^2$ besitzt auf dem Intervall $[-1, 2]$ das Maximum $f(2) = 4$ und das Minimum $f(0) = 0$. Die Extremstelle $x = 2$ liegt am Rand des Intervalls, die Extremstelle $x = 0$ im Inneren des Intervalls. Die Funktion $f(x) = x^2$ besitzt auf dem Intervall $[1, 2]$ das Maximum $f(2) = 4$ und das Minimum $f(1) = 1$. Beide Extremstellen, $x = 1$ und $x = 2$ liegen am Rand des Intervalls.

FOLGERUNG 12. *Wenn $f : [a, b] \rightarrow \mathbb{R}$ stetig, dann ist f auf $[a, b]$ beschränkt, d.h. es existiert ein $S \in \mathbb{R}$ mit $|f(x)| \leq S$ für alle $x \in [a, b]$.*

4. Verknüpfung stetiger Funktionen

Aus den Sätzen 43 und 42 folgt sofort:

SATZ 46. *Wenn $f, g : D \rightarrow \mathbb{R}$ auf D stetige Funktionen sind, dann sind auch $f + g$ und fg auf D stetige Funktionen. Wenn zusätzlich $g(x) \neq 0$ für alle $x \in D$, dann ist auch f/g auf D stetig.*

FOLGERUNG 13. *Alle Polynome sind stetig auf $\mathbb{R}$.*

BEWEIS. Jedes Polynom entsteht durch die Verknüpfung Addition und Multiplikation aus den stetigen Funktionen $g \equiv c$ mit $c \in \mathbb{R}$ und $f(x) = x$. $\square$

FOLGERUNG 14. *Wenn p und q zwei Polynome sind, dann ist der Quotient p/q auf $D := \{x \in \mathbb{R} : q(x) \neq 0\}$ stetig.*

SATZ 47 (Stetigkeit verketteter Funktionen). *Wenn für zwei Funktionen f, g gilt, dass g stetig in x_0 und f stetig in $g(x_0)$ ist, dann ist $f \circ g$ stetig in x_0 .*

BEWEIS. Für jedes $\varepsilon > 0$ existiert ein $\delta_1 > 0$ mit $|f(g(x)) - f(g(x_0))| < \varepsilon$ für alle x mit $|g(x) - g(x_0)| < \delta_1$, da f in $g(x_0)$ stetig ist. Da g in x_0 stetig ist, existiert zu diesem δ_1 ein $\delta > 0$ mit $|g(x) - g(x_0)| < \delta_1$ für alle x mit $|x - x_0| < \delta$. Zusammen gilt $|f(g(x)) - f(g(x_0))| < \varepsilon$ für alle x mit $|x - x_0| < \delta$. $\square$

BEISPIEL 56. Die Funktionen $\sin(3x+5)$ und e^{x^2} sind als Verkettung der Funktionen $f(x) = \sin x$ und $g(x) = 3x+5$ bzw. $f(x) = e^x$ und $g(x) = x^2$ stetige Funktionen.

Möchte man eine Funktion $f : D \rightarrow \mathbb{R}$ mit Definitionsbereich D nur auf der Teilmenge $M \subset D$ betrachten, so schränkt man die Funktion auf diese Teilmenge M ein. Man schreibt $f|_M$ um anzudeuten, dass man nur Argumente aus M betrachtet.

LEMMA 15. *Es seien f und g zwei Funktionen, deren Definitionsbereich das offene Intervall (a, b) enthält.*

Wenn $f|_{(a,b)} = g|_{(a,b)}$ und g auf (a, b) stetig ist, dann ist auch f auf (a, b) stetig.

BEWEIS. Für alle $\delta > 0$ und alle $x_0 \in (a, b)$ ist $\{x : |x - x_0| < \delta\} \cap (a, b)$ wieder eine offene Menge, die x_0 enthält. Darum existiert ein $\delta \geq \delta' > 0$ mit

$$(\{x : |x - x_0| < \delta\} \cap (a, b)) \supset \{x : |x - x_0| < \delta'\}.$$

$\square$

BEISPIEL 57. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

ist in allen Punkten $x_0 \neq 0$ stetig, denn auf den offenen Intervallen $(-\infty, 0)$ und $(0, \infty)$ stimmt sie mit der Funktion $\sin(1/x)$ überein, die für $x \neq 0$ als Verkettung der stetigen Funktionen $\sin x$ und $1/x$ stetig ist.

LEMMA 16. *Wenn die Funktion $f : (a, b) \rightarrow \mathbb{R}$ auf (a, b) stetig ist und eine Umkehrfunktion besitzt, dann gilt entweder $f(x_1) < f(x_2)$ für alle $a \leq x_1 < x_2 \leq b$ oder $f(x_1) > f(x_2)$ für alle $a \leq x_1 < x_2 \leq b$.*

BEWEIS. Es genügt zu zeigen, dass für $a \leq x_1 < x < x_2 \leq b$ immer $f(x_1) < f(x) < f(x_2)$ oder $f(x_1) > f(x) > f(x_2)$ gilt.

Wenn $f(x_1) > f(x) < f(x_2)$, so folgt aus Satz 44

$$[f(x), \min(f(x_1), f(x_2))] \subset f((x_1, x)) \text{ und } [f(x), \min(f(x_1), f(x_2))] \subset f((x, x_2)).$$

Wenn $f(x_1) < f(x) > f(x_2)$, so folgt aus Satz 44

$$[\max(f(x_1), f(x_2)), f(x)] \subset f((x_1, x)) \text{ und } [\max(f(x_1), f(x_2)), f(x)] \subset f((x, x_2)).$$

Dies widerspricht der Injektivität der Funktion f . $\square$

SATZ 48 (Stetigkeit der Umkehrfunktion). *Wenn die Funktion $f : (a, b) \rightarrow \mathbb{R}$ auf (a, b) stetig ist und eine Umkehrfunktion besitzt, so ist die Umkehrfunktion stetig.*

BEWEIS. Wenn $y_0 \in f((a, b))$, dann gilt $y_0 = f(x_0)$ für $x_0 := f^{-1}(y_0)$. Für $\varepsilon > 0$ existieren $x_1, x_2 \in [a, b]$ mit $x_0 - \varepsilon < x_1 < x_0 < x_2 < x_0 + \varepsilon$. Aus dem Zwischenwertsatz folgt

$$y_0 \in [f(x_1), f(x_2)] \subset f([x_1, x_2]), \text{ falls } f(x_1) < f(x_2),$$

oder

$$y_0 \in [f(x_2), f(x_1)] \subset f([x_1, x_2]), \text{ falls } f(x_2) < f(x_1).$$

Aus der Injektivität der Funktion f folgt $f^{-1}(y) \in (x_1, x_2)$ für alle y mit $|y - y_0| < \delta < \min(|y_0 - f(x_1)|, |y_0 - f(x_2)|)$. $\square$

FOLGERUNG 15. *Die Funktion $g : [0, \infty) \rightarrow \mathbb{R}$, $g(x) = \sqrt[k]{x}$, ist für alle $k \in \mathbb{N}$ mit $k > 0$ stetig.*

BEWEIS. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^k$ ist stetig auf $\mathbb{R}$ und injektiv auf $[0, \infty)$. Ihre Umkehrfunktion ist g . $\square$

5. Nullstellen von Polynomen

Eine Funktion $p : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$p(x) = a_d x^d + a_{d-1} x^{d-1} + \dots + a_1 x + a_0 = \sum_{n=0}^d a_n x^n \text{ mit } a_n \in \mathbb{R}, a_d \neq 0$$

heißt **reelles Polynom** bzw. Polynom mit reellen Koeffizienten vom Grad d . Jedes reelle Polynom ist stetig auf $\mathbb{R}$. Ein Polynom vom Grad 0 ist eine konstante Funktion.

5.1. Lineare Funktionen ($d = 1$). Jedes Polynom vom Grad 1 ist injektiv auf $\mathbb{R}$. Die Umkehrfunktion zu $p(x) = ax + b$ mit $a \neq 0$ ist die lineare Funktion $p^{-1}(y) = (y - b)/a$. Das lineare Polynom p hat genau eine Nullstelle, $p^{-1}(0) = -b/a$.

5.2. Quadratische Polynome ($d = 2$). Es genügt Polynome vom Grad 2 mit $a_2 = 1$ gut zu verstehen, denn

$$a_2 x^2 + a_1 x + a_0 = a_2 \left(x^2 + \frac{a_1}{a_2} x + \frac{a_0}{a_2} \right) = a_2 (x^2 + px + q) \text{ mit } p = \frac{a_1}{a_2}, q = \frac{a_0}{a_2}.$$

Für das quadratische Polynom $f(x) = x^2 + px + q$ mit $p, q \in \mathbb{R}$ gilt

$$f(x) = x^2 + 2\frac{p}{2}x + q = x^2 + 2\frac{p}{2}x + \frac{p^2}{4} - \frac{p^2}{4} + q = \left(x + \frac{p}{2}\right)^2 + q - \frac{p^2}{4}.$$

Daraus folgt:

- $f(-p/2) = q - p^2/4$ ist das Minimum der Funktion f .
- $f(-p/2 + x) = f(-p/2 - x)$ für alle $x \in \mathbb{R}$, d.h. f ist symmetrisch bezüglich Spiegelung an der Geraden $x = -p/2$.

- f besitzt genau dann eine Nullstelle, wenn $D := p^2 - 4q \geq 0$. Diese sind $x_{1,2} = (-p \pm \sqrt{D})/2$.

5.3. Polynome höheren Grades ($d > 2$). Wenn $a_d = 1$, dann gilt

$$p(x) = x^d + a_{d-1}x^{d-1} + \dots + a_1x + a_0 = x^d \left(1 + \frac{a_{d-1}}{x} + \dots + \frac{a_1}{x^{d-1}} + \frac{a_0}{x^d} \right) = x^d h(x).$$

Für $|x| > 2d \max(|a_{d-1}|, \dots, |a_0|)$ gilt $\frac{1}{2} < h(x) < \frac{3}{2}$ und damit

$$\frac{x^d}{2} < f(x) < \frac{3}{2}x^d, \text{ falls } x > 0 \text{ oder } 2|d$$

und

$$\frac{x^d}{2} > f(x) > \frac{3}{2}x^d, \text{ falls } x < 0 \text{ und } 2 \nmid d,$$

Daraus folgt

$$\lim_{x \rightarrow \pm\infty} f(x) = \infty, \text{ falls } d \text{ gerade ist,}$$

und

$$\lim_{x \rightarrow \infty} f(x) = \infty, \lim_{x \rightarrow -\infty} f(x) = -\infty, \text{ falls } d \text{ ungerade ist.}$$

Aus dem Zwischenwertsatz folgt

SATZ 49. *Jedes reelle Polynom ungeraden Grades besitzt mindestens eine reelle Nullstelle.*

BEISPIEL 58. Das reelle Polynom $x^4 + 1$ hat Grad 4 und besitzt keine reelle Nullstelle, denn $x^4 + 1 \geq 1 > 0$ für alle $x \in \mathbb{R}$.

5.4. Zerlegung in lineare und quadratische Faktoren.

BEISPIEL 59. Wir betrachten das reelle Polynom $f(x) = x^3 - 8x + 3$. Wir suchen Nullstellen und Intervalle, die Nullstellen der Funktion f enthalten. Es gilt $f(0) = 3 > 0$. Um schnell, d.h. mit wenigen Rechenoperationen, Funktionswerte $f(x)$ auszurechnen, ist es sinnvoll, den Term $x^3 - 8x + 3$ folgendermaßen zu schreiben:

$$f(x) = x^3 - 8x + 3 = 3 + x(-8 + x(0 + x(1)))$$

Nun gilt

$$f(0) = 3 + 0(-8 + 0(0 + 0(1))) = 3 > 0$$

$$f(1) = 3 + 1(-8 + 1(0 + 1(1))) = -4 < 0$$

$$f(-1) = 3 - 1(-8 - 1(0 - 1(1))) = 10 > 0$$

$$f(2) = 3 + 2(-8 + 2(0 + 2(1))) = -5 < 0$$

$$f(-2) = 3 - 2(-8 - 2(0 - 2(1))) = 11 > 0$$

$$f(3) = 3 + 3(-8 + 3(0 + 3(1))) = 6 > 0$$

$$f(-3) = 3 - 3(-8 - 3(0 - 3(1))) = 0.$$

Die Funktion f hat also die Nullstelle $x_0 = -3$. Da $f(0) > 0$ und $f(1) < 0$, befindet sich eine weitere Nullstelle im Intervall $(0, 1)$, weil f stetig ist. Da

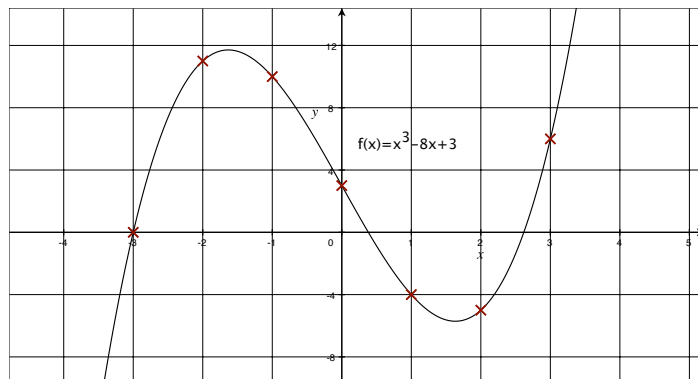


ABBILDUNG 3. Polynom dritten Grades

$f(2) < 0$ und $f(3) > 0$, befindet sich eine weitere Nullstelle im Intervall $(2, 3)$, weil f stetig ist. Ein reelles Polynom vom Grad 3 kann höchstens drei reelle Nullstellen besitzen, also $f(x) < 0$ für alle $x < -3$ und $f(x) > 0$ für alle $x > 3$. Anhand der berechneten Funktionswerte läßt sich der Graph der Funktion f skizzieren (siehe Abbildung 3).

Polynomdivision liefert $f(x) = (x - (-3))(x^2 - 3x + 1) = (x + 3)(x^2 - 3x + 1)$, denn

$$\begin{aligned} x^3 - 8x + 3 &= (x + 3)x^2 - 3x^2 - 8x + 3 \\ -3x^2 - 8x + 3 &= (x + 3)(-3x) + x + 3 \\ x + 3 &= (x + 3)1 + 0. \end{aligned}$$

Die Koeffizienten des Quotienten $q(x) := f(x) : (x + 3)$ erhält man auch bei der schrittweisen Auswertung der Klammern in $f(-3)$ von innen nach außen:

$$f(-3) = 3 - 3(-8 - 3(0 - 3(\mathbf{1}))) = 3 - 3(-8 - 3(-\mathbf{3})) = 3 - 3(\mathbf{1}) = 0.$$

Der Wert in der innersten Klammer ist der führende Koeffizient im Polynom $q(x)$. Die reellen Nullstellen von q sind $x_{1,2} = (3 \pm \sqrt{5})/2$. Daraus folgt

$$f(x) = x^3 - 8x + 3 = (x + 3) \left(x - \frac{3 + \sqrt{5}}{2} \right) \left(x - \frac{3 - \sqrt{5}}{2} \right).$$

Das Polynom $f(x)$ läßt sich als Produkt von Linearfaktoren schreiben.

SATZ 50. Wenn p ein reelles Polynom vom Grad $d > 0$ ist und $p(x_0) = 0$ für ein $x_0 \in \mathbb{R}$ gilt, dann existiert ein reelles Polynom q vom Grad $d - 1$ mit $p(x) = (x - x_0)q(x)$.

BEWEIS. Wenn $p(x) = \sum_{n=0}^d a_n x^n$ mit $a_n \in \mathbb{R}$ und $a_d \neq 0$, dann gilt

$$\begin{aligned} p(x) &= p(x) - 0 = p(x) - p(x_0) = \sum_{n=0}^d a_n x^n - \sum_{n=0}^d a_n x_0^n \\ &= \sum_{n=0}^d a_n (x^n - x_0^n) = \sum_{n=1}^d a_n (x^n - x_0^n), \end{aligned}$$

da $x^0 - x_0^0 = 1 - 1 = 0$.

Falls $x_0 = 0$, so gilt $p(x_0) = a_0 = 0$, also

$$p(x) = \sum_{n=1}^d a_n x^n = x \sum_{n=1}^d a_n x^{n-1} = x \sum_{n=0}^{d-1} a_{n+1} x^n = xq(x) \text{ mit } q(x) = \sum_{n=0}^{d-1} a_{n+1} x^n.$$

Falls $x_0 \neq 0$, so folgt aus der geometrischen Summenformel

$$\begin{aligned} \left(1 - \frac{x}{x_0}\right) \sum_{k=0}^{n-1} \left(\frac{x}{x_0}\right)^k &= 1 - \left(\frac{x}{x_0}\right)^n \\ (x_0 - x) \sum_{k=0}^{n-1} x^k x_0^{n-1-k} &= x_0^n - x^n \\ (x - x_0) \sum_{k=0}^{n-1} x^k x_0^{n-1-k} &= x^n - x_0^n \end{aligned}$$

und

$$\begin{aligned} p(x) &= \sum_{n=1}^d a_n (x - x_0) \sum_{k=0}^{n-1} x^k x_0^{n-1-k} = (x - x_0) \sum_{n=1}^d a_n \sum_{k=0}^{n-1} x^k x_0^{n-1-k} \\ &= (x - x_0) \sum_{k=0}^{d-1} x^k \left(\sum_{n=k+1}^d a_n x_0^{n-1-k} \right) = (x - x_0)q(x) \text{ mit} \\ q(x) &= \sum_{k=0}^{d-1} x^k \left(\sum_{n=k+1}^d a_n x_0^{n-1-k} \right). \end{aligned}$$

□

BEMERKUNG 41. Das Polynom $q(x)$ in Satz 50 berechnet man in konkreten Fällen durch Polynomdivision oder Berechnung der Klammern in $p(x_0)$ wie in Beispiel 59 und nicht mit der expliziten Formel. □

BEMERKUNG 42. Der Satz 50 gilt ähnlich auch für komplexe Nullstellen, also $z_0 \in \mathbb{C}$ mit $p(z_0) = 0$. Jedoch besitzt der Quotient $q(x) := f(x)/(x - z_0)$ dann komplexe Koeffizienten und ist kein reelles Polynom mehr. □

Da jedes nichtkonstante Polynom eine komplexe Nullstelle besitzt, existieren zu jedem Polynom $p(x)$ vom Grad d komplexe Zahlen $\lambda_1, \dots, \lambda_d \in \mathbb{C}$

mit $p(x) = (x - \lambda_1) \dots (x - \lambda_d)$. Hat das Polynom $p(x)$ nur reelle Koeffizienten, ist es also ein reelles Polynom, so treten die komplexen Nullstellen immer paarweise auf.

SATZ 51. *Wenn p ein reelles Polynom ist und $p(z_0) = 0$ für ein $z_0 \in \mathbb{C}$ gilt, dann ist auch $\bar{z}_0$ eine Nullstelle von p , also $p(\bar{z}_0) = 0$.*

BEWEIS. Wenn $p(x) = \sum_{n=0}^d a_n x^n$ mit $a_n \in \mathbb{R}$ und $p(z_0) = 0$, dann folgt $p(z_0) = 0 = \sum_{n=0}^d a_n z_0^n$ genau dann, wenn

$$\bar{0} = \overline{\sum_{n=0}^d a_n z_0^n} \Leftrightarrow 0 = \sum_{n=0}^d \overline{a_n z_0^n} \Leftrightarrow 0 = \sum_{n=0}^d \bar{a}_n \bar{z}_0^n \Leftrightarrow 0 = \sum_{n=0}^d a_n \bar{z}_0^n \Leftrightarrow 0 = p(\bar{z}_0).$$

Für diese Umformungen benötigt man die Rechenregeln $\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$ für alle $z_1, z_2 \in \mathbb{C}$, $\overline{z_1 \cdot z_2} = \bar{z}_1 \cdot \bar{z}_2$ für alle $z_1, z_2 \in \mathbb{C}$ und $\bar{\bar{x}} = x$ für alle $x \in \mathbb{R}$. $\square$

FOLGERUNG 16. *Für jedes reelle Polynom $p(x) = \sum_{n=0}^d a_n x^n$ mit $a_n \in \mathbb{R}$ und $a_d = 1$ existieren $x_1, \dots, x_r \in \mathbb{R}$ und $\lambda_1, \dots, \lambda_m \in \mathbb{C} \setminus \mathbb{R}$ mit $d = r + 2m$ und*

$$(52) \quad p(x) = \prod_{k=1}^r (x - x_k) \prod_{k=1}^m q_k(x) \text{ für } q_k(x) := x^2 - 2\Re(\lambda_k)x + |\lambda_k|^2.$$

BEWEIS. Ist $x_0 \in \mathbb{R}$ eine reelle Nullstelle, so gilt $p(x) = (x - x_0)q(x)$ für ein reelles Polynom q mit kleinerem Grad. Wiederholt man diese Abspaltung von Linearfaktoren $x - x_k$ so erhält man $p(x) = q(x) \prod_{k=1}^r (x - x_k)$ wobei $q(x)$ ein reelles Polynom ohne reelle Nullstellen ist. Die komplexen Nullstellen von q treten paarweise auf. Für alle $\lambda \in \mathbb{C}$ gilt

$$(x - \lambda)(x - \bar{\lambda}) = x^2 - x(\lambda + \bar{\lambda}) + \lambda\bar{\lambda} = x^2 - x2\Re(\lambda) + |\lambda|^2.$$

Für diese Umformungen braucht man $z + \bar{z} = 2\Re(z)$ und $z\bar{z} = |z|^2$ für alle $z \in \mathbb{C}$. $\square$

BEMERKUNG 43. Die Linearfaktoren $x - x_k$ und die quadratischen Faktoren $q_k(x)$ in Gleichung 52 sind bis auf Reihenfolge eindeutig bestimmt. Die reellen Nullstellen x_k und die quadratischen Faktoren q_k müssen nicht paarweise verschieden sein (siehe Beispiel 61). $\square$

BEISPIEL 60. Das Polynom $p(x) = x^4 + 1$ hat keine reellen Nullstellen, denn $x^4 + 1 \geq 1$ für alle $x \in \mathbb{R}$. Es existieren also zwei reelle, quadratische Polynome q_1 und q_2 mit $p = q_1 q_2$. Man kann diese quadratischen Faktoren mit einem Ansatz oder mit Hilfe der komplexen Nullstellen bestimmen:

- Mit dem Ansatz $q_1(x) = x^2 + ax + b$ und $q_2(x) = x^2 + cx + d$ für $a, b, c, d \in \mathbb{R}$ erhält man

$$\begin{aligned} x^4 + 1 &= (x^2 + ax + b)(x^2 + cx + d) \\ &= x^4 + (a + c)x^3 + (b + ac + d)x^2 + (ad + bc)x + bd, \end{aligned}$$

also $a + c = 0$, $b + ac + d = 0$, $ad + bc = 0$ und $bd = 1$. Aus der ersten und der letzten Bedingung folgt $c_0 a$ und $d = 1/b$. Die restlichen beiden Gleichungen werden dann zu

$$0 = b - a^2 + \frac{1}{b} \text{ und } 0 = a \left(\frac{1}{b} - b \right).$$

Da $b + 1/b \neq 0$ für alle $b \in \mathbb{R}$, gilt $a \neq 0$ und damit $1/b - b = 0$, also $b = \pm 1$. Wegen $a^2 = b + 1/b$, muss $b = 1$ und $a^2 = 2$ sein. Daraus folgt $b = d = 1$ und $a = \sqrt{2} = -c$ oder $a = -\sqrt{2} = -c$, also

$$x^4 + 1 = (x^2 + \sqrt{2}x + 1)(x^2 - \sqrt{2}x + 1).$$

- Die komplexen Nullstellen des Polynoms $x^4 + 1$ sind die Lösungen der Gleichung $z^4 = -1$. Da $-1 = e^{i\pi}$ folgt aus der Lösungsformel

$$\lambda_k = e^{i(\frac{\pi}{4} + k\frac{\pi}{2})} \quad k = 0, 1, 2, 3.$$

Es gilt

$$\overline{\lambda_0} = \overline{e^{i\pi/4}} = e^{-i\pi/4} = e^{i7\pi/4} = \lambda_3 \text{ und } \overline{\lambda_1} = \overline{e^{i3\pi/4}} = e^{-i3\pi/4} = e^{i5\pi/4} = \lambda_2,$$

denn für alle $\phi \in \mathbb{R}$ gilt

$$(53) \quad \overline{e^{i\phi}} = \overline{\cos \phi + i \sin \phi} = \cos \phi - i \sin \phi = \cos(-\phi) + i \sin(-\phi) = e^{-i\phi},$$

da $\cos$ eine gerade und $\sin$ eine ungerade Funktion ist. Nun folgt mit $\Re(\lambda_0) = \sqrt{2}/2$, $\Re(\lambda_1) = -\sqrt{2}/2$ und $|\lambda_k| = 1$ für $k = 0, 1, 2, 3$

$$x^4 + 1 = (x^2 - 2\Re(\lambda_0)x + |\lambda_0|^2)(x^2 - 2\Re(\lambda_1)x + |\lambda_1|^2) = (x^2 - \sqrt{2}x + 1)(x^2 + \sqrt{2}x + 1).$$

BEISPIEL 61. Das reelle Polynom $p(x) = x^4 + 2x^3 + 2x^2 + 2x + 1$ besitzt die reelle Nullstelle $x_1 = -1$, denn

$$\begin{aligned} f(x) &= 1 + x(2 + x(2 + x(2 + x(1)))) \\ f(-1) &= 1 - 1(2 - 1(2 - 1(2 - 1(\mathbf{1})))) \\ &= 1 - 1(2 - 1(2 - 1(\mathbf{1}))) = 1 - 1(2 - 1(\mathbf{1})) = 1 - 1(\mathbf{1}) = 0. \end{aligned}$$

Es folgt $p(x) = (x - (-1))(x^3 + x^2 + x + 1) = (x + 1)h(x)$ mit $h(x) = x^3 + x^2 + x + 1$ und

$$\begin{aligned} h(x) &= 1 + x(1 + x(1 + x(1))) \\ f(-1) &= 1 - 1(1 - 1(1 - 1(\mathbf{1}))) = 1 - 1(1 - 1(\mathbf{0})) = 1 - 1(\mathbf{1}) = 0, \end{aligned}$$

also $h(x) = (x + 1)(x^2 + 1)$. Da $x^2 + 1$ keine reellen Nullstellen hat und ein quadratisches Polynom ist, hat man die Zerlegung $p(x) = x^4 + 2x^3 + 2x^2 + 2x + 1 = (x + 1)(x + 1)(x^2 + 1)$ gefunden. Der Linearfaktor $x + 1$ tritt doppelt auf.

KAPITEL VII

Differenzierbarkeit

Eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ heißt im Punkt $x_0 \in (a, b)$ **differenzierbar**, falls ein $A \in \mathbb{R}$ existiert, so dass

$$\lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} = A =: f'(x_0).$$

Der Grenzwert heißt **Ableitung** von f in x_0 und wird mit $f'(x_0)$ bezeichnet.

Eine Funktion f ist genau dann in x_0 differenzierbar mit $f'(x_0) = A$, wenn für jede Folge $\{x_n\}$ mit $\lim_{n \rightarrow \infty} x_n = x_0$ gilt

$$\lim_{n \rightarrow \infty} \frac{f(x_n) - f(x_0)}{x_n - x_0} = A.$$

Der Grenzwert ist also von der gewählten, gegen x_0 konvergenten Folge $\{x_n\}$ unabhängig.

Eine Funktion f ist genau dann in x_0 differenzierbar mit $f'(x_0) = A$, wenn eine Funktion $r : (a, b) \rightarrow \mathbb{R}$ mit

$$f(x) = f(x_0) + A(x - x_0) + r(x) \text{ und } \lim_{x \rightarrow x_0} \frac{r(x)}{x - x_0} = 0$$

existiert. Für $x \approx x_0$ gilt also $f(x) \approx f(x_0) + A(x - x_0)$ und der Fehler $r(x)$, den man bei dieser Approximation macht ist so klein, dass er stärker gegen 0 konvergiert als die Differenz $x - x_0$ der Argumente.

Der Graph der Funktion $g(x) = f(x_0) + A(x - x_0)$ ist die **Tangente** an (den Graphen der Funktion) f im Punkt x_0 .

Abbildung 1 zeigt den Graph der Funktion $f(x) = x^2$ und die Tangente (blau) an diesen Graph im Punkt $x_0 = 2$. Man sieht deutlich, dass die Tangente für $x \approx 2$ fast mit dem Graph von f übereinstimmt.

BEMERKUNG 44. Anschaulich ist eine Funktion differenzierbar, wenn ihr Graph keine Knicke besitzt. Die Abbildungen 2 und 3 zeigen zwei Funktionsgraphen, die Knicke besitzen. Für manche Funktionen ist es aber auch unmöglich, den Graph zu skizzieren, so dass man nicht entscheiden kann, ob „Knicke“ vorhanden sind (siehe Beispiel 65). $\square$

BEISPIEL 62 (Ableitung konstanter Funktionen). Für jedes $c \in \mathbb{R}$ ist die konstante Funktion $f \equiv c$ für alle $x_0 \in \mathbb{R}$ differenzierbar und es gilt $f'(x_0) = 0$ für alle $x_0 \in \mathbb{R}$, denn

$$\frac{f(x) - f(x_0)}{x - x_0} = \frac{c - c}{x - x_0} = \frac{0}{x - x_0} = 0 \xrightarrow{x \rightarrow x_0} 0.$$

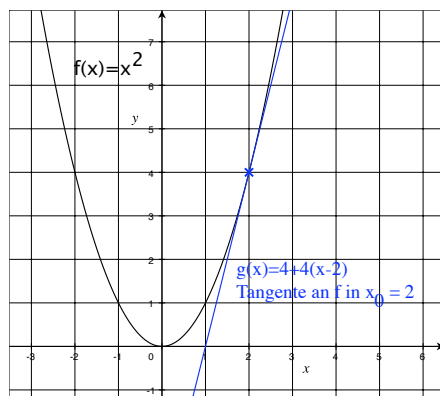
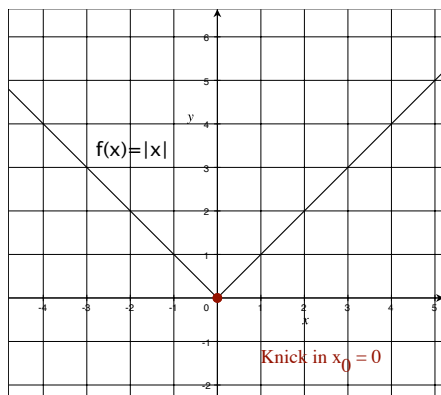
ABBILDUNG 1. $f(x) = x^2$ mit Tangente

ABBILDUNG 2

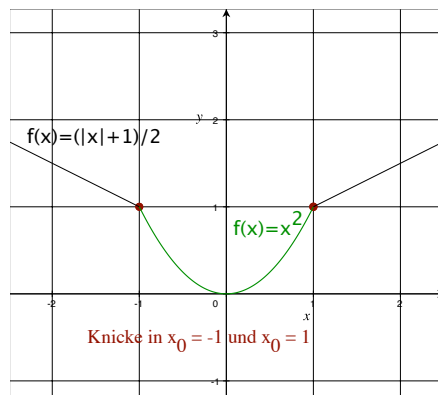


ABBILDUNG 3

BEISPIEL 63 (Ableitung von x). Die Funktion $f(x) = x$ für alle $x_0 \in \mathbb{R}$ differenzierbar und es gilt $f'(x_0) = 1$ für alle $x_0 \in \mathbb{R}$, denn

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

BEISPIEL 64 (Ableitung von x^2). Die Funktion $f(x) = x^2$ für alle $x_0 \in \mathbb{R}$ differenzierbar und es gilt $f'(x_0) = 2x_0$ für alle $x_0 \in \mathbb{R}$, denn

$$\frac{f(x) - f(x_0)}{x - x_0} = \frac{x^2 - x_0^2}{x - x_0} = \frac{(x - x_0)(x + x_0)}{x - x_0} = x + x_0 \xrightarrow{x \rightarrow x_0} x_0 + x_0 = 2x_0.$$

BEISPIEL 65. In Beispiel 52 wurde gezeigt, dass die Funktion

$$f(x) = \begin{cases} x \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

in $x_0 = 0$ stetig ist. Die Funktion f ist in $x_0 = 0$ aber nicht differenzierbar, denn

$$\frac{f(x) - f(x_0)}{x - x_0} = \frac{x \sin(1/x) - 0}{x - 0} = \frac{x \sin(1/x)}{x} = \sin(1/x)$$

und die Funktion $\sin(1/x)$ ist in $x_0 = 0$ nicht stetig (siehe Beispiel 51). Zum Beispiel gilt

$$\lim_{n \rightarrow \infty} \sin(1/x_n) = 0 \text{ für } x_n = \frac{1}{n\pi} \text{ und } \lim_{n \rightarrow \infty} \sin(1/x_n) = 1 \text{ für } x_n = \frac{1}{2n\pi + \pi/2}.$$

BEISPIEL 66 (Betragsfunktion). Die Funktion $f(x) = |x|$ ist in allen Punkten $x_0 \neq 0$ differenzierbar, denn $f|_{(-\infty, 0)}(x) = -x$ und $f|_{(0, \infty)}(x) = x$, also $f'(x) = -1$ für alle $x < 0$ und $f'(x) = 1$ für alle $x > 0$. Die Funktion $f(x) = |x|$ ist in $x_0 = 0$ nicht differenzierbar, denn

$$\lim_{n \rightarrow \infty} \frac{f(x_n) - f(0)}{x_n - 0} = \lim_{n \rightarrow \infty} \frac{|1/n|}{1/n} = 1 \text{ für } x_n = \frac{1}{n}$$

und

$$\lim_{n \rightarrow \infty} \frac{f(x_n) - f(0)}{x_n - 0} = \lim_{n \rightarrow \infty} \frac{|-1/n|}{-1/n} = \lim_{n \rightarrow \infty} \frac{1/n}{-1/n} = -1 \text{ für } x_n = -\frac{1}{n}.$$

SATZ 52 (Stetigkeit differenzierbarer Funktionen). Wenn f in $x_0 \in (a, b)$ differenzierbar ist, dann ist f in x_0 stetig.

BEWEIS. Für jede Folge $\{x_n\}$ mit $\lim_{n \rightarrow \infty} x_n = x_0$ gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} f(x_n) - f(x_0) &= \lim_{n \rightarrow \infty} \frac{f(x_n) - f(x_0)}{x - x_0} (x - x_0) \\ &= \lim_{n \rightarrow \infty} \frac{f(x_n) - f(x_0)}{x - x_0} \lim_{n \rightarrow \infty} (x - x_0) = f'(x_0) \cdot 0 = 0, \end{aligned}$$

also $\lim_{n \rightarrow \infty} f(x_n) = f(x_0)$. $\square$

BEMERKUNG 45. Die Eigenschaft „differenzierbar“ ist also stärker als die Eigenschaft „stetig“. Differenzierbarkeit impliziert Stetigkeit. $\square$

1. Ableitungsregeln

Nur für einfache Funktionen läßt sich die Ableitung mit Hilfe der Definition bestimmen. Die Berechnung der Ableitung komplizierter Funktionen führt man mit Hilfe der Ableitungsregeln auf Ableitungen der Grundfunktionen (siehe Beispiele 62 und 63) zurück.

SATZ 53. Wenn $f, g : (a, b) \rightarrow \mathbb{R}$ differenzierbar in $x_0 \in (a, b)$, so sind auch $f + g$ und fg in x_0 differenzierbar und es gilt

$$(54) \quad (f + g)'(x_0) = f'(x_0) + g'(x_0)$$

und

$$(55) \quad (fg)'(x_0) = f'(x_0)g(x_0) + f(x_0)g'(x_0) \quad (\text{Produktregel}).$$

Ist zusätzlich $g(x) \neq 0$ für alle $x \in (a, b)$, so gilt

$$(56) \quad \left(\frac{f}{g}\right)'(x_0) = \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g(x_0)^2} \quad (\text{Quotientenregel}).$$

BEWEIS. Es gilt

$$\begin{aligned} \frac{(f+g)(x) - (f+g)(x_0)}{x - x_0} &= \frac{f(x) + g(x) - (f(x_0) + g(x_0))}{x - x_0} \\ &= \frac{f(x) - f(x_0)}{x - x_0} + \frac{g(x) - g(x_0)}{x - x_0} \xrightarrow{x \rightarrow x_0} f'(x_0) + g'(x_0) \end{aligned}$$

$$\begin{aligned} \frac{(fg)(x) - (fg)(x_0)}{x - x_0} &= \frac{f(x)g(x) - f(x_0)g(x_0)}{x - x_0} \\ &= \frac{f(x)g(x) - f(x_0)g(x) + f(x_0)g(x) - f(x_0)g(x_0)}{x - x_0} \\ &= \frac{f(x) - f(x_0)}{x - x_0}g(x) + f(x_0)\frac{g(x) - g(x_0)}{x - x_0} \\ &\xrightarrow{x \rightarrow x_0} f'(x_0)g(x_0) + f(x_0)g'(x_0) \end{aligned}$$

$$\begin{aligned} \left(\frac{f}{g}\right)(x) - \left(\frac{f}{g}\right)(x_0) &= \frac{f(x)}{g(x)} - \frac{f(x_0)}{g(x_0)} = \frac{f(x)g(x_0) - g(x)f(x_0)}{g(x)g(x_0)} \\ &= \frac{f(x)g(x_0) - f(x_0)g(x_0) + f(x_0)g(x_0) - g(x)f(x_0)}{g(x)g(x_0)} \\ &= \frac{(f(x) - f(x_0))g(x_0) - f(x_0)(g(x) - g(x_0))}{g(x)g(x_0)} \\ \frac{(f/g)(x) - (f/g)(x_0)}{x - x_0} &= \frac{1}{g(x)g(x_0)} \left(\frac{f(x) - f(x_0)}{x - x_0}g(x_0) - f(x_0)\frac{g(x) - g(x_0)}{x - x_0} \right) \\ &\xrightarrow{x \rightarrow x_0} \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g(x_0)g(x_0)} \end{aligned}$$

□

BEISPIEL 67 (Ableitung von $1/x$). Die Funktion $h(x) = 1/x$ ist für alle $x_0 \neq 0$ differenzierbar, denn $h = f/g$ mit $f \equiv 1$ und $g(x) = x$. Aus der Quotientenregel folgt mit $f'(x_0) = 0$ und $g'(x_0) = 1$ (siehe Beispiele 62 und 63) für alle $x_0 \neq 0$

$$h'(x_0) = \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g(x_0)^2} = \frac{0 \cdot x_0 - 1 \cdot 1}{x_0^2} = -\frac{1}{x_0^2} = -x_0^{-2}.$$

BEISPIEL 68 (Ableitung von x^n). Für jedes $n \in \mathbb{N}$ ist die Funktion $h(x) = x^n$ für alle $x_0 \neq 0$ differenzierbar und es gilt $h'(x_0) = nx_0^{n-1}$, denn

$$\begin{aligned} \frac{h(x) - h(x_0)}{x - x_0} &= \frac{x^n - x_0^n}{x - x_0} = x^{n-1} \frac{1 - \left(\frac{x_0}{x}\right)^n}{1 - \frac{x_0}{x}} = x^{n-1} \sum_{k=0}^{n-1} \left(\frac{x_0}{x}\right)^k \\ &\xrightarrow{x \rightarrow x_0} x_0^{n-1} \sum_{k=0}^{n-1} (1)^k = nx_0^{n-1}. \end{aligned}$$

Man kann diese Aussage auch mit vollständiger Induktion beweisen. Für $n = 1$ ist $h(x) = x^n = x$ und $h'(x) = 1 = x^0$ (Induktionsanfang). Für den Induktionsschritt benutzt man die Produktregel $h = fg$ mit $f(x) = x$ und $g(x) = x^n$. Mit $f'(x_0) = 1$ und $g'(x_0) = nx_0^{n-1}$ folgt $h'(x_0) = f'(x_0)g(x_0) + f(x_0)g'(x_0) = 1 \cdot x_0^n + x_0 \cdot nx_0^{n-1} = (n+1)x_0^n$.

BEISPIEL 69 (Ableitung von $cf(x)$). Für jedes $c \in \mathbb{R}$ und jede Funktion f , die in x_0 differenzierbar ist, gilt $(cf)'(x_0) = cf'(x_0)$, denn mit $g(x) \equiv c$ und $g'(x_0) = 0$ folgt aus der Produktregel $(cf)'(x_0) = (gf)'(x_0) = f'(x_0)g(x_0) + f(x_0)g'(x_0) = cf'(x_0) + f(x_0) \cdot 0 = cf'(x_0)$.

BEISPIEL 70 (Ableitung von $1/x^n$ mit Quotientenregel). Die Funktion $h(x) = 1/x^n$ ist für alle $x_0 \neq 0$ differenzierbar, denn $h = f/g$ mit $g(x) = x^n$ und $f \equiv 1$. Aus der Quotientenregel folgt mit $f'(x_0) = 0$ und $g'(x_0) = nx_0^{n-1}$ für alle $x_0 \neq 0$

$$\begin{aligned} h'(x_0) &= \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g(x_0)^2} = \frac{0 \cdot x_0^n - 1 \cdot nx_0^{n-1}}{x_0^n x_0^n} \\ &= -n \frac{x_0^{n-1}}{x_0^{2n}} = -n \frac{1}{x_0^{n+1}} = -nx_0^{-n-1}. \end{aligned}$$

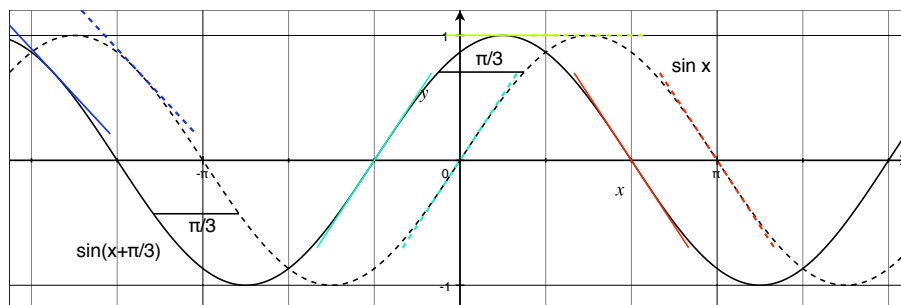
Mit Hilfe der Ableitungsregeln und der Beispiele 68 und 69 kann man beliebige Polynome differenzieren.

FOLGERUNG 17. Für alle $m, n \in \mathbb{Z}$ und alle $a_k \in \mathbb{R}$ ist die Funktion $f(x) = \sum_{k=m}^n a_k x^k$ für alle $x_0 \in \mathbb{R}$ (bzw. $x_0 \neq 0$, falls $a_k \neq 0$ für ein $k < 0$) differenzierbar und es gilt

$$(57) \quad f'(x_0) = \left(\sum_{k=m}^n a_k x^k \right)' \Big|_{x=x_0} = \sum_{k=m}^n k a_k x^{k-1}.$$

SATZ 54 (Kettenregel). Wenn g in x_0 und f in $g(x_0)$ differenzierbar ist, dann ist $f \circ g$ in x_0 differenzierbar und es gilt

$$(58) \quad (f \circ g)'(x_0) = f'(g(x_0))g'(x_0).$$

ABBILDUNG 4. Ableitungen von $f(x)$ und $f(x + b)$

BEWEIS. Da g in x_0 stetig ist, folgt

$$\begin{aligned} \frac{(f \circ g)(x) - (f \circ g)(x_0)}{x - x_0} &= \frac{f(g(x)) - f(g(x_0))}{x - x_0} \\ &= \frac{f(g(x)) - f(g(x_0))}{g(x) - g(x_0)} \cdot \frac{g(x) - g(x_0)}{x - x_0} \xrightarrow{x \rightarrow x_0} f'(g(x_0)) \cdot g'(x_0). \end{aligned}$$

□

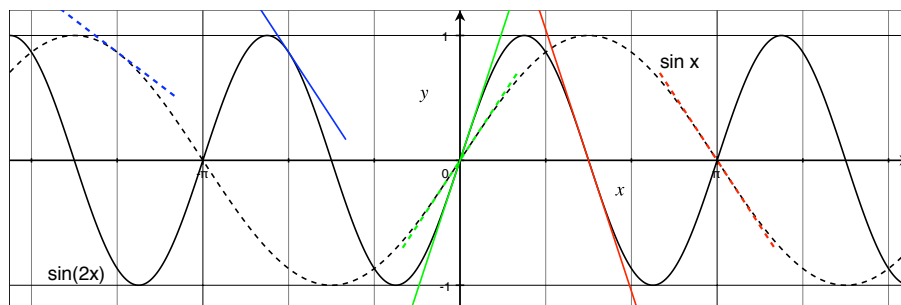
BEISPIEL 71 (Ableitung von $1/x^n$ mit Kettenregel). Die Funktion $h(x) = 1/x^n$ ist für alle $x_0 \neq 0$ differenzierbar, denn $h = f \circ g$ mit $g(x) = x^n$ und $f(x) = 1/x$. Aus der Kettenregel folgt mit $f'(x_0) = -1/x_0^2$ und $g'(x_0) = nx_0^{n-1}$ für alle $x_0 \neq 0$

$$h'(x_0) = f'(g(x_0))g'(x_0) = -\frac{1}{g(x_0)^2}nx_0^{n-1} = -\frac{1}{x_0^n x_0^n}nx_0^{n-1} = -n\frac{1}{x_0^{n+1}}.$$

BEISPIEL 72 (Ableitung von e^{-x^2} mit Kettenregel). Die Funktion $h(x) = e^{-x^2}$ ist für alle $x_0 \in \mathbb{R}$ differenzierbar, denn $h = f \circ g$ mit $g(x) = -x^2$ und $f(x) = e^x$. Aus der Kettenregel folgt mit $f'(x_0) = e^{x_0}$, denn mit dieser Eigenschaft wurde die Exponentialfunktion definiert, und $g'(x_0) = -2x_0$ für alle $x_0 \in \mathbb{R}$ folgt

$$h'(x_0) = f'(g(x_0))g'(x_0) = e^{g(x_0)}(-2x_0) = -2x_0e^{-x_0^2}.$$

Verkettungen einer Funktion f mit $g_b(x) = x + b$ oder $g_a(x) = ax$ mit $a, b \in \mathbb{R}$ mit $a \neq 0$, treten häufig auf. Es gilt $(f \circ g_b)(x) = f(x + b)$ und $(f \circ g_a)(x) = f(ax)$. Wegen $g'_b(x) = 1$ und $g'_a(x) = a$ folgt aus der Kettenregel $(f \circ g_b)'(x) = f'(x + b) \cdot 1 = f'(x + b)$ und $(f \circ g_a)'(x) = f'(ax) \cdot a = af'(ax)$. Abbildung 4 zeigt die Graphen der Funktionen $\sin(x + \pi/3)$ und $\sin x$ und die Tangenten in ausgewählten Punkten x_0 (an $\sin(x + \pi/3)$) und $x_0 + \pi/3$ (an $\sin x$). Abbildung 5 zeigt die Graphen der Funktionen $\sin(2x)$ und $\sin x$ und die Tangenten in ausgewählten Punkten x_0 (an $\sin(2x)$) und $2x_0$ (an $\sin x$). Für $g(x) := (g_b \circ g_a)(x) = ax + b$ gilt damit $g'(x) = a$, $(f \circ g)(x) = f(ax + b)$ und $(f \circ g)'(x) = af'(ax + b)$.

ABBILDUNG 5. Ableitungen von $f(x)$ und $f(ax)$

BEMERKUNG 46. Der Graph der Funktion $f \circ g_b$ ist der um b nach links verschobene Graph der Funktion f . Man erhält den Graph der Funktion $f \circ g_a$ durch Änderung der Skala auf der x -Achse im Graph der Funktion f . Die x -Werte werden a -mal schneller durchlaufen.

2. l'Hospitalsche Regel

SATZ 55. Wenn die Funktionen f und g in x_0 differenzierbar sind und $f(x_0) = g(x_0) = 0$ gilt, dann

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{f'(x_0)}{g'(x_0)},$$

falls $g'(x_0) \neq 0$.

BEWEIS. Da f und g in x_0 differenzierbar sind, existieren Funktionen r_f, r_g mit $f(x) = f(x_0) + f'(x_0)(x - x_0) + r_f(x)$, $g(x) = g(x_0) + g'(x_0)(x - x_0) + r_g(x)$ und

$$\lim_{x \rightarrow x_0} \frac{r_f(x)}{x - x_0} = \lim_{x \rightarrow x_0} \frac{r_g(x)}{x - x_0} = 0.$$

Da $f(x_0) = g(x_0) = 0$ folgt

$$\begin{aligned} \frac{f(x)}{g(x)} &= \frac{f(x_0) + f'(x_0)(x - x_0) + r_f(x)}{g(x_0) + g'(x_0)(x - x_0) + r_g(x)} = \frac{f'(x_0)(x - x_0) + r_f(x)}{g'(x_0)(x - x_0) + r_g(x)} \\ &= \frac{f'(x_0) + \frac{r_f(x)}{x - x_0}}{g'(x_0) + \frac{r_g(x)}{x - x_0}} \xrightarrow{x \rightarrow x_0} \frac{f'(x_0) + 0}{g'(x_0) + 0} = \frac{f'(x_0)}{g'(x_0)}, \end{aligned}$$

falls dieser Quotient existiert, also $g'(x_0) \neq 0$. □

BEISPIEL 73. Für $f(x) = e^x - 1$ und $g(x) = x$ gilt $f(0) = 0$, $g(0) = 0$, $f'(x) = e^x$, $f'(0) = 1$, $g'(x) = 1$ und $g'(0) = 1$, also

$$\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = \frac{f'(0)}{g'(0)} = \frac{1}{1} = 1.$$

BEMERKUNG 47. Wenn $f'(x_0) \neq 0$ und $g'(0) = 0$, so ist $\frac{f(x)}{g(x)}$ für $x \rightarrow x_0$ unbeschränkt. $\square$

BEMERKUNG 48. Mit Hilfe von Taylorpolynomen (siehe Abschnitt 2) läßt sich gut beschreiben, wie oft man die l'Hospitalsche Regel auf Ausdrücke f/g anwenden darf und muss, um den Grenzwert $\lim_{x \rightarrow \infty} f(x)/g(x)$ zu berechnen. $\square$

Ist eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ in allen $x \in (a, b)$ differenzierbar, so wird durch die Zuordnung $x \mapsto f'(x)$ eine Funktion $f' : (a, b) \rightarrow \mathbb{R}$ definiert. Sie heißt **Ableitung** der Funktion f . Ist die Funktion f' stetig auf (a, b) , so heißt f **stetig differenzierbar**.

BEISPIEL 74 (Nicht stetig differenzierbare Funktion). Die Funktion

$$f(x) = \begin{cases} x^2 \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

ist differenzierbar mit $f'(x) = 2x \sin(1/x) - \cos(1/x)$ für $x \neq 0$ und $f'(0) = \lim_{x \rightarrow 0} x \sin(1/x) = 0$. Da $\lim_{x \rightarrow 0} \cos(1/x)$ nicht existiert, ist f' in $x_0 = 0$ nicht stetig.

SATZ 56 (Mehrfache Anwendung der l'Hospitalschen Regel). *Wenn die Funktionen f und g stetig differenzierbar auf (a, b) sind und $f(x_0) = g(x_0) = 0$ für ein $x_0 \in (a, b)$ gilt, dann*

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)},$$

falls $g'(x) \neq 0$ für alle $x \neq x_0$ und der rechte Grenzwert existiert.

BEWEIS. Falls $g'(x_0) \neq 0$, so folgt die Behauptung aus der Stetigkeit der Funktionen f' und g' und Satz 55. Falls

$$\lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)} = a \in \mathbb{R},$$

so existieren zu jedem $\varepsilon > 0$ ein $\delta > 0$ mit $g'(x) \neq 0$ für alle $x \neq x_0$ mit $|x - x_0| < \delta$ und

$$\left| \frac{f'(x)}{g'(x)} - a \right| < \varepsilon \quad \forall |x - x_0| < \delta, \text{ also } |f'(x) - ag'(x)| < \varepsilon |g'(x)| \quad \forall |x - x_0| < \delta.$$

Da $f(x_0) = g(x_0) = 0$, folgt mit Hilfe des Hauptsatzes der Differential- und Integralrechnung (siehe)

$$\begin{aligned} \frac{f(x)}{g(x)} - a &= \frac{f(x) - ag(x)}{g(x)} = \frac{(f(x) - f(x_0)) - a(g(x) - g(x_0))}{g(x)} \\ &= \frac{1}{g(x)} \left(\int_{x_0}^x f'(\xi) d\xi - a \int_{x_0}^x g'(\xi) d\xi \right) = \frac{1}{g(x)} \int_{x_0}^x (f'(\xi) - ag'(\xi)) \\ \left| \frac{f(x)}{g(x)} - a \right| &\leq \frac{1}{|g(x)|} \int_{x_0}^x |f'(\xi) - ag'(\xi)| d\xi < \frac{\varepsilon}{|g(x)|} \int_{x_0}^x |g'(\xi)| d\xi \end{aligned}$$

Falls $g'(x) > 0$ für $x > x_0$, so $|g'(\xi)| = g'(\xi)$ und $g(x) \geq 0$ für alle $x_0 < \xi < x$ und

$$\left| \frac{f(x)}{g(x)} - a \right| < \frac{\varepsilon}{|g(x)|} (g(x) - g(x_0)) = \frac{\varepsilon}{|g(x)|} g(x) = \varepsilon.$$

Falls $g'(x) > 0$ für $x < x_0$, so $|g'(\xi)| = g'(\xi)$ und $g(x) \leq 0$ für alle $x < \xi < x_0$ und

$$\left| \frac{f(x)}{g(x)} - a \right| < \frac{\varepsilon}{|g(x)|} (-g(x) + g(x_0)) = \frac{\varepsilon}{|g(x)|} (-g(x)) = \varepsilon.$$

Falls $g'(x) < 0$ für $x > x_0$, so $|g'(\xi)| = -g'(\xi)$ und $g(x) \leq 0$ für alle $x_0 < \xi < x$ und

$$\left| \frac{f(x)}{g(x)} - a \right| < \frac{\varepsilon}{|g(x)|} (-g(x) + g(x_0)) = \frac{\varepsilon}{|g(x)|} (-g(x)) = \varepsilon.$$

Falls $g'(x) < 0$ für $x < x_0$, so $|g'(\xi)| = -g'(\xi)$ und $g(x) \geq 0$ für alle $x < \xi < x_0$ und

$$\left| \frac{f(x)}{g(x)} - a \right| < \frac{\varepsilon}{|g(x)|} (g(x) - g(x_0)) = \frac{\varepsilon}{|g(x)|} g(x) = \varepsilon.$$

□

BEMERKUNG 49. Der Beweis des Satzes 56 benutzt Ergebnisse aus Kapitel VIII. □

3. Monotonie

Das Vorzeichen der Ableitung $f'(x_0)$ einer Funktion f im Punkt x_0 beschreibt das Wachstumsverhalten der Funktion f in der Nähe des Punktes x_0 .

LEMMA 17. Wenn $f'(x_0) > 0$, dann existiert ein $\delta > 0$ mit $f(x_1) < f(x) < f(x_2)$ für alle $x - \delta < x_1 < x < x_2 < x + \delta$.

BEWEIS. Da f in x_0 differenzierbar ist, existiert eine Funktion r mit

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

Es existiert also ein $\delta > 0$ mit

$$\left| \frac{r(x)}{x - x_0} \right| < \frac{f'(x_0)}{2} \quad \forall |x - x_0| < \delta, \text{ also } |r(x)| < \frac{f'(x_0)}{2} |x - x_0| \quad \forall |x - x_0| < \delta.$$

Für $x_0 + \delta > x_2 > x_0$ folgt $r(x_2) > -\frac{1}{2}f'(x_0)(x_2 - x_0)$ und damit

$$\begin{aligned} f(x_2) &> f(x_0) + f'(x_0)(x_2 - x_0) - \frac{f'(x_0)(x_2 - x_0)}{2} \\ &= f(x_0) + \frac{f'(x_0)(x_2 - x_0)}{2} > f(x_0), \end{aligned}$$

da $x_2 - x_0 > 0$ und $f'(x_0) > 0$.

Für $x_0 - \delta < x_1 < x_0$ folgt $r(x_1) < -\frac{1}{2}f'(x_0)(x_1 - x_0)$ und damit

$$\begin{aligned} f(x_1) &> f(x_0) + f'(x_0)(x_1 - x_0) - \frac{f'(x_0)(x_1 - x_0)}{2} \\ &= f(x_0) + \frac{f'(x_0)(x_1 - x_0)}{2} > f(x_0), \end{aligned}$$

da $x_1 - x_0 < 0$ und $f'(x_0) > 0$. $\square$

FOLGERUNG 18. Wenn $f'(x_0) < 0$, dann existiert ein $\delta > 0$ mit $f(x_1) > f(x)$ für alle $x - \delta < x_1 < x < x_2 < x + \delta$.

BEWEIS. Für die Funktion $g(x) := -f(x)$ gilt $g'(x_0) = -f'(x_0) > 0$, $g(x) > g(x_0) \Leftrightarrow f(x) < f(x_0)$ und $g(x) < g(x_0) \Leftrightarrow f(x) > f(x_0)$. $\square$

Eine Funktion f heißt **streng monoton wachsend** auf dem Intervall (a, b) , falls $f(x_1) < f(x_2)$ für alle $x_1, x_2 \in (a, b)$ mit $x_1 < x_2$.

Eine Funktion f heißt **streng monoton fallend** auf dem Intervall (a, b) , falls $f(x_1) > f(x_2)$ für alle $x_1, x_2 \in (a, b)$ mit $x_1 < x_2$.

FOLGERUNG 19. Wenn $f'(x) > 0$ für alle $x \in (a, b)$, dann ist f streng monoton wachsend auf (a, b) .

Wenn $f'(x) < 0$ für alle $x \in (a, b)$, dann ist f streng monoton fallend auf (a, b) .

BEMERKUNG 50. Das Monotoniekriterium in Folgerung 19 ist nicht notwendig. Zum Beispiel ist die Funktion $f(x) = x^3$ streng monoton wachsend, aber $f'(x) = 3x^2$ und $f'(0) = 0$. $\square$

FOLGERUNG 20. Es sei $f : [a, b] \rightarrow \mathbb{R}$ stetig auf $[a, b]$.

Wenn $f'(x) > 0$ für alle $x \in (a, b)$, dann ist $f(a)$ das Minimum und $f(b)$ das Maximum von f auf $[a, b]$.

Wenn $f'(x) < 0$ für alle $x \in (a, b)$, dann ist $f(a)$ das Maximum und $f(b)$ das Minimum von f auf $[a, b]$.

BEISPIEL 75. Die Funktion $f(x) = x^2$ besitzt die Ableitung $f'(x) = 2x$. Für alle $x \in [1, 2]$ gilt $f'(x) > 0$. Darum ist $f(1) = 1$ das Minimum und $f(2) = 4$ das Maximum der Funktion f auf $[1, 2]$. Es gilt $f'(-1/2) = -1 < 0$ und $f(1/2) = 1 > 0$. Auf dem Intervall $[-1, 2]$ wird das Minimum der Funktion x^2 nicht am Rand angenommen.

BEMERKUNG 51. Eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ kann ihre Extremwerte auch am Rand annehmen, wenn sie nicht monoton ist. In Abschnitt 4 wird die Bestimmung der Extremwerte einer differenzierbaren Funktion besprochen. $\square$

FOLGERUNG 21. Wenn $f'(x) > 0$ für alle $x \in (a, b)$ oder $f'(x) < 0$ für alle $x \in (a, b)$, dann besitzt f auf (a, b) eine Umkehrfunktion.

BEMERKUNG 52. Das Kriterium für die Existenz einer Umkehrfunktion in Folgerung 21 ist nicht notwendig. Zum Beispiel besitzt die Funktion $f(x) = x^3$ die Umkehrfunktion $f^{-1}(x) = \sqrt[3]{x}$. Jedoch gilt $f'(0) = 0$. $\square$

BEMERKUNG 53. In Abschnitt 5 wird die Ableitung der Umkehrfunktion bestimmt. $\square$

4. Extremwerte

Sei $f : [a, b] \rightarrow \mathbb{R}$ eine stetige und auf (a, b) differenzierbare Funktion. Wenn $f(x_0)$ ein Maximum, also $f(x) \leq f(x_0)$ für alle $x \in [a, b]$, oder ein Minimum, also $f(x) \geq f(x_0)$ für alle $x \in [a, b]$, dann ist $f(x_0)$ auch ein Extremwert der Funktion f im Vergleich zu Funktionswerten $f(x)$ für Argumente x , die in der Nähe von x_0 liegen. Diese lokalen Extremstellen, lassen sich mit Hilfe der Ableitung einfach charakterisieren.

Der Funktionswert $f(x_0)$ heißt **lokales Maximum**, falls ein $\delta > 0$ mit der Eigenschaft $f(x) \leq f(x_0)$ für alle x mit $|x - x_0| < \delta$ existiert. Der Funktionswert $f(x_0)$ heißt **lokales Minimum**, falls ein $\delta > 0$ mit der Eigenschaft $f(x) \geq f(x_0)$ für alle x mit $|x - x_0| < \delta$ existiert. Wenn $f(x_0)$ ein lokales Minimum oder ein lokales Maximum ist, dann heißt das Argument x_0 **lokale Extremstelle**.

SATZ 57. Wenn $x_0 \in (a, b)$ eine lokale Extremstelle der auf (a, b) differenzierbaren und auf $[a, b]$ stetigen Funktion $f : [a, b] \rightarrow \mathbb{R}$ ist, dann $f'(x_0) = 0$.

BEWEIS. Die Behauptung folgt aus Lemma 17 und Folgerung 18. $\square$

Die Eigenschaft $f'(x_0) = 0$ ist nicht hinreichend dafür, dass x_0 eine lokale Extremstelle ist. Zum Beispiel besitzt $f(x) = x^3$ keine lokalen Extremstellen, aber $f'(0) = 0$. Betrachtet man aber das Vorzeichen der Funktion f' in der Nähe eines Punktes x_0 mit $f'(x_0) = 0$, so kann man entscheiden, ob x_0 eine lokale Extremstelle ist.

SATZ 58. Es sei $f'(x_0) = 0$.

Wenn ein $\delta > 0$ mit $f'(x) > 0$ für $x_0 - \delta < x < x_0$ und $f'(x) < 0$ für $x_0 < x < x_0 + \delta$, dann ist x_0 ein lokales Maximum.

Wenn ein $\delta > 0$ mit $f'(x) < 0$ für $x_0 - \delta < x < x_0$ und $f'(x) > 0$ für $x_0 < x < x_0 + \delta$, dann ist x_0 ein lokales Minimum.

BEWEIS. Die Behauptung folgt aus Lemma 17 und Folgerung 18. $\square$

Wenn $f(x)$ ein Extremwert der Funktion f auf $[a, b]$ ist, dann ist x eine lokale Extremstelle oder es gilt $x = a$ oder $x = b$.

BEISPIEL 76. Die Funktion $f(x) = 2x^3 + 3x^2 - 36x + 7$ ist auf dem Intervall $[-4, 4]$ stetig und auf $(-4, 4)$ differenzierbar. Wir bestimmen zuerst die lokalen Extremstellen und -werte und vergleichen diese dann mit den Funktionswerten am Rand des Intervalls, also mit $f(-4)$ und $f(4)$. Es gilt

$$f'(x) = 6x^2 + 6x - 36 = 6(x^2 + x - 6) = 6(x + 3)(x - 2),$$

also sind $x = -3$ und $x = 2$ die Nullstellen von $f'(x)$. Da $f'(x) > 0$ für $x < -3$, $f'(x) < 0$ für $-3 < x < 2$ und $f'(x) > 0$ für $x > 2$, sind $x = -3$

und $x = 2$ lokale Extremstellen, $f(-3)$ ein lokales Maximum und $f(2)$ ein lokales Minimum. Da

$$\begin{aligned} f(x) &= 7 + x(-36 + x(3 + x(2))) \\ f(-4) &= 7 - 4(-36 - 4(3 - 4(2))) = 71 \\ f(-3) &= 7 - 3(-36 - 3(3 - 3(2))) = 88 \\ f(2) &= 7 + 2(-36 + 2(3 + 2(2))) = -37 \\ f(4) &= 7 + 4(-36 + 4(3 + 4(2))) = 39 \end{aligned}$$

ist $f(2)$ das globale Minimum und $f(-3)$ das globale Maximum der Funktion f auf dem Intervall $[-4, 4]$.

5. Differenzierbarkeit der Umkehrfunktion

In Folgerung 21 wurde gezeigt, dass eine streng monotone Funktion eine Umkehrfunktion besitzt. In diesem Abschnitt berechnen wir die Ableitung der Umkehrfunktion f^{-1} mit Hilfe der Ableitung der Funktion f .

SATZ 59 (Ableitung der Umkehrfunktion). *Wenn $f'(x) > 0$ für alle $x \in (a, b)$ oder $f'(x) < 0$ für alle $x \in (a, b)$, dann ist die Umkehrfunktion für alle $y_0 \in f(a, b)$ differenzierbar und es gilt*

$$(59) \quad (f^{-1})'(y_0) = \frac{1}{f'(x_0)} \text{ für } y_0 = f(x_0).$$

BEWEIS. Es gilt

$$\begin{aligned} \frac{f^{-1}(y) - f^{-1}(y_0)}{y - y_0} &= \frac{f^{-1}(f(x)) - f^{-1}(f(x_0))}{f(x) - f(x_0)} = \frac{x - x_0}{f(x) - f(x_0)} \\ &= \frac{1}{\frac{f(x) - f(x_0)}{x - x_0}} \xrightarrow{y \rightarrow y_0} \frac{1}{f'(x_0)}, \end{aligned}$$

da f^{-1} eine stetige Funktion ist (siehe 48) und aus $y \rightarrow y_0$ auch $x \rightarrow x_0$ folgt. $\square$

BEISPIEL 77. Die Funktion $f(x) = x^k$ ist für alle $k \in \mathbb{N}$, $k > 0$, auf dem Intervall $(0, \infty)$ differenzierbar und streng monoton wachsend, denn $f'(x) = kx^{k-1} > 0$ für alle $x > 0$. Die Umkehrfunktion $f^{-1} = \sqrt[k]{} : (0, \infty) \rightarrow \mathbb{R}$ ist also auch differenzierbar und es gilt

$$(\sqrt[k]{y_0})' = \frac{1}{f'(x_0)} = \frac{1}{kx_0^{k-1}} = \frac{1}{k\sqrt[k]{y_0}^{k-1}} = \frac{1}{k} \cdot \frac{1}{\frac{y_0^{\frac{k-1}{k}}}{y_0^{\frac{1}{k}}}} = \frac{1}{k} y_0^{-\frac{k-1}{k}} = \frac{1}{k} y_0^{\frac{1}{k}-1},$$

also

$$(60) \quad \left(x^{\frac{1}{k}}\right)' = \frac{1}{k} x^{\frac{1}{k}-1} \text{ für alle } x > 0.$$

[Ableitung der Wurzelfunktionen]

FOLGERUNG 22. Für alle $n, m \in \mathbb{N}$ mit $m \neq 0$ gilt

$$(61) \quad \left(x^{\frac{n}{m}}\right)' = \frac{n}{m} x^{\frac{n}{m}-1} \text{ für alle } x > 0.$$

BEWEIS. Die Funktion $h(x) = x^{n/m}$ ist Verkettung $f \circ g$ der Funktionen f und g mit $g(x) = x^{1/m}$ und $f(x) = x^n$. Aus $f'(x) = nx^{n-1}$ und $g'(x) = \frac{1}{m}x^{-(m-1)/m}$ folgt mit Hilfe der Kettenregel

$$\begin{aligned} \left(x^{\frac{n}{m}}\right)' &= f'(g(x))g'(x) = n \left(x^{1/m}\right)^{n-1} \frac{1}{m} x^{-\frac{m-1}{m}} \\ &= \frac{n}{m} x^{\frac{n-1}{m}} x^{-\frac{m-1}{m}} = \frac{n}{m} x^{\frac{n-m}{m}} = \frac{n}{m} x^{\frac{n}{m}-1}. \end{aligned}$$

□

BEISPIEL 78 (Ableitung von arccos). Die Funktion $\cos$ ist auf dem Intervall $(0, \pi)$ differenzierbar und streng monoton fallend, also injektiv, denn $(\cos x)' = -\sin x < 0$ für alle $0 < x < \pi$. Die Umkehrfunktion $\arccos : (-1, 1) \rightarrow (0, \pi)$ ist differenzierbar. Da $\arccos(\cos x) = x$, folgt aus der Kettenregel

$$1 = \arccos(\cos(x)) = \arccos'(\cos x)(\cos x)' = -(\sin x) \arccos'(\cos x).$$

Wegen $\sin^2 x + \cos^2 x = 1$ folgt $\sin x = \sqrt{1 - \cos^2 x}$, da $\sin x > 0$ für $0 < x < \pi$. Das bedeutet

$$\arccos'(\cos x) = -\frac{1}{\sqrt{1 - \cos^2 x}},$$

also (mit $x = y = \cos x$)

$$(62) \quad (\arccos x)' = -\frac{1}{\sqrt{1 - x^2}} \text{ für } -1 < x < 1.$$

6. Mittelwertsatz

SATZ 60 (Satz von Rolle). Wenn f differenzierbar auf (a, b) und $f(x_1) = f(x_2)$ für $x_1, x_2 \in (a, b)$ mit $x_1 < x_2$, dann existiert ein $x \in (x_1, x_2)$ mit $f'(x) = 0$.

BEWEIS. Die Funktion f ist auf dem Intervall $[x_1, x_2]$ stetig und besitzt deshalb ein Maximum und ein Minimum auf $[a, b]$. Wenn das Maximum oder das Minimum von f auf $[x_1, x_2]$ in einem inneren Punkt angenommen wird, also für $x_0 \in (x_1, x_2)$, so ist x_0 eine lokale Extremstelle und es gilt $f'(x_0) = 0$. Wenn $f(x_1) = f(x_2)$ Minimum und Maximum der Funktion f , so gilt $f(x) = f(x_1)$ für alle $x \in (x_1, x_2)$, also $f'(x) = 0$ für alle $x \in (x_1, x_2)$. □

Abbildung 6 illustriert den Satz von Rolle, zwischen zwei Nullstellen einer differenzierbaren Funktion befindet sich (mindestens) ein lokales Extremum, also ein x_0 mit $f'(x_0) = 0$, d.h. die Tangente an den Graph von f im Punkt x_0 ist parallel zur x -Achse.

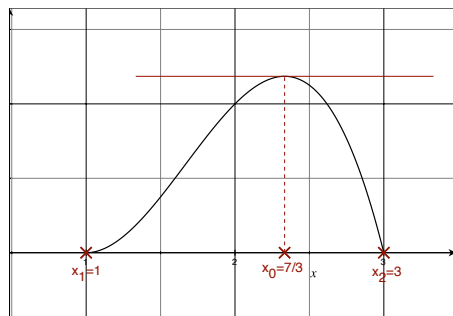


ABBILDUNG 6

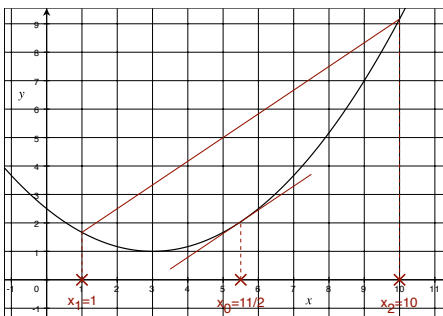


ABBILDUNG 7. Mittelwertsatz

SATZ 61 (Mittelwertsatz). Wenn f differenzierbar auf (a, b) und $x_1, x_2 \in (a, b)$ mit $x_1 < x_2$, dann existiert ein $x \in (x_1, x_2)$ mit

$$\frac{f(x_2) - f(x_1)}{x_2 - x_1} = f'(x).$$

BEWEIS. Die Funktion

$$h(x) := f(x) - \left(f(x_1) + (x - x_1) \frac{f(x_2) - f(x_1)}{x_2 - x_1} \right)$$

ist die Differenz der Funktion f und der Geraden $f(x_1) + \frac{f(x_2) - f(x_1)}{x_2 - x_1}(x - x_1)$, die durch die Punkte $(x_1, f(x_1))$ und $(x_2, f(x_2))$ verläuft. Da h eine stetige und differenzierbare Funktion ist und $h(x_1) = 0$ und $h(x_2) = 0$ gilt, existiert ein $x_0 \in (x_1, x_2)$ mit $h'(x_0) = 0$. Es gilt $h'(x) = f'(x) - \frac{f(x_2) - f(x_1)}{x_2 - x_1}$, also

$$0 = h'(x_0) = f'(x_0) - \frac{f(x_2) - f(x_1)}{x_2 - x_1}, \text{ d.h. } f'(x_0) = \frac{f(x_2) - f(x_1)}{x_2 - x_1}.$$

□

Abbildung 7 illustriert den Mittelwertsatz, zwischen zwei Argumenten $x_1 < x_2$ einer differenzierbaren Funktion befindet sich ein Punkt, dessen Ableitung mit dem Anstieg der Geraden durch $(x_1, f(x_1))$ und $(x_2, f(x_2))$ übereinstimmt.

FOLGERUNG 23. Wenn $f'(x) = 0$ für alle $x \in (a, b)$, dann ist f auf dem Intervall (a, b) konstant, d.h. es existiert ein $c \in \mathbb{R}$ mit $f \equiv c$ auf dem Intervall (a, b) .

BEWEIS. Für beliebige $x_1 < x_2$ existiert ein $x \in (x_1, x_2)$ mit

$$f'(x) = \frac{f(x_2) - f(x_1)}{x_2 - x_1}.$$

Aber es gilt $f'(x) = 0$. Daraus folgt $f(x_1) = f(x_2)$.

□

KAPITEL VIII

Das Riemann-Integral

Für eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ mit positiven Funktionswerten soll das Integral $\int_a^b f(x)dx$ den Inhalt A der Fläche messen, die durch den Graph der Funktion f , die x -Achse und die Geraden $x = a$ und $x = b$ begrenzt wird. Für einige Funktionen ist dieser Flächeninhalt einfach zu berechnen, z.B.

- $A = (b - a)c$ für jede konstante Funktion $f \equiv c \in \mathbb{R}$, da die Fläche ein Rechteck mit den Seitenlängen $b - a$ und c ist (siehe Abbildung 1).
- $A = \frac{1}{2}(b^2 - a^2)$ für die Funktion $f(x) = x$, da die Fläche ein Trapez ist (siehe), also den Flächeninhalt $\frac{1}{2}(b + a)(b - a)$ hat (siehe Abbildung 2).

Auf jeden Fall läßt sich der Flächeninhalt abschätzen. Wenn m das Minimum und M das Maximum der Funktion f auf dem Intervall $[a, b]$ ist, so gilt $U \leq A \leq O$ mit $U = (b - a)m$ und $O = (b - a)M$ für den Flächeninhalt A (siehe Abbildung 3). Diese Abschätzung kann sich verbessern, wenn man das Intervall $[a, b]$ in Teilintervalle zerlegt. Abbildung 4 zeigt die gleiche Funktion wie Abbildung 3, jedoch eine Zerlegung des Intervalls $[a, b]$ in drei Teilintervalle. Es gilt $U_1 + U_2 + U_3 \leq A \leq O_1 + O_2 + O_3$.

Hinter der Definition des Riemann-Integrals steckt die Idee, dass die Differenz der Flächeninhalte der den Funktionsgraph überdeckenden Rechtecke und der Rechtecke unter dem Funktionsgraph sehr klein wird, wenn man das Intervall $[a, b]$ nur in genügend viele sehr kurze Teilintervalle zerlegt. Um diese Idee in eine Definition zu verwandeln, muss man sicher sein

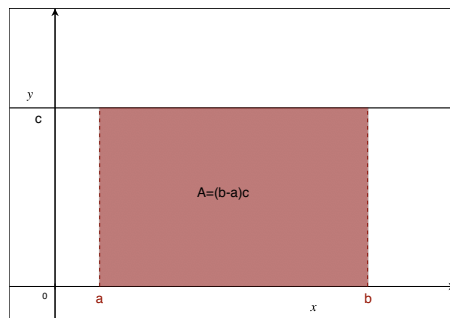


ABBILDUNG 1. $\int_a^b c dx$

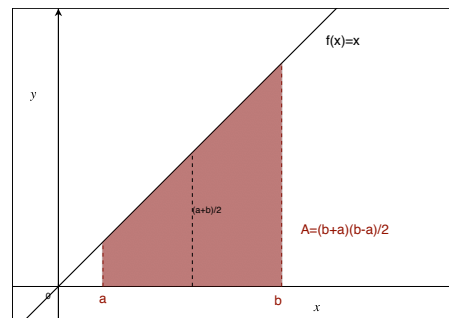


ABBILDUNG 2. $\int_a^b x dx$

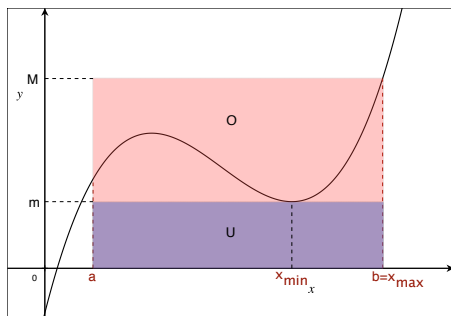


ABBILDUNG 3

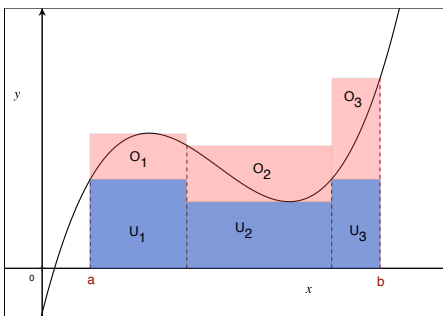


ABBILDUNG 4

können, dass es immer Rechtecke gibt, die den Graph überdecken bzw. sich unter dem Graph befinden.

1. Definition des Riemann-Integrals

1.1. Supremum und Infimum. Es sei $M \subset \mathbb{R}$ eine beschränkte Menge. Es existieren $s, S \in \mathbb{R}$ mit $S \geq y$ für alle $y \in M$ und $s \leq y$ für alle $y \in M$. Die Zahl S ist eine obere Schranke und die Zahl s eine untere Schranke für die Menge M . Wenn S obere Schranke und Element der Menge M ist, so ist S ein Maximum der Menge M . Wenn s untere Schranke und Element der Menge M ist, so ist s ein Minimum der Menge M .

Es gibt Mengen, z.B. $\{y \in \mathbb{R} : -1 < y < 3\}$, die zwar beschränkt sind, aber kein Minimum oder Maximum besitzen. Diese Lücke wird durch die Begriffe Infimum und Supremum gefüllt.

Das **Supremum** einer nichtleeren, nach oben beschränkten Menge M ist die kleinste obere Schranke der Menge. Sie wird mit $\sup M$ bezeichnet. Es gilt genau dann $S = \sup M$, wenn S eine obere Schranke der Menge M ist und jede kleinere Zahl als S keine obere Schranke der Menge M ist.

Jede nach oben beschränkte Teilmenge der reellen Zahlen besitzt ein Supremum. Ist die Menge M nicht nach oben beschränkt, so schreiben wir $\sup M = \infty$.

Das **Infimum** einer nichtleeren, nach unten beschränkten Menge M ist die größte untere Schranke der Menge. Sie wird mit $\inf M$ bezeichnet. Es gilt genau dann $s = \inf M$, wenn s eine untere Schranke der Menge M ist und jede größere Zahl als s keine untere Schranke der Menge M ist.

Jede nach unten beschränkte Teilmenge der reellen Zahlen besitzt ein Infimum. Ist die Menge M nicht nach unten beschränkt, so schreiben wir $\inf M = -\infty$.

BEISPIEL 79. Die Menge $M = \{x \in \mathbb{R} : -1 < x \leq 3\}$ ist beschränkt. Es gilt $\sup M = 3$ und $\inf M = -1$. Das Supremum ist gleichzeitig ein Maximum der Menge. Das Infimum ist kein Minimum der Menge, da $-1 \notin M$.

BEISPIEL 80. Die Menge $M = \{-1, 1\}$ besteht nur aus zwei Elementen. Es gilt $\sup M = \max M = 1$ und $\inf M = \min M = -1$. Die Folge $\{x_n\}_{n \in \mathbb{N}}$ mit $x_n = (-1)^n$ nimmt auch nur zwei Werte an. Sie konvergiert aber nicht, besitzt also keinen Grenzwert.

1.2. Zerlegung. Eine **Zerlegung** $\mathcal{Z}$ des Intervalls $[a, b]$ ist durch Zwischenpunkte $x_k \in [a, b]$ mit $a = x_0 < x_1 < \dots < x_n = b$ gegeben. Das Intervall $[a, b]$ ist damit in die n Teilintervalle $I_k := [x_{k-1}, x_k]$ mit $k = 1, \dots, n$ zerlegt.

Das Teilintervall I_k hat die Länge $|I_k| := x_k - x_{k-1}$. Das **Feinheitsmaß** $|\mathcal{Z}|$ der Zerlegung $\mathcal{Z}$ definiert als $|\mathcal{Z}| := \max\{|I_k| : k = 1, \dots, n\}$.

Eine Zerlegung $\mathcal{Z}_n$ des Intervalls $[a, b]$ in n Teilintervalle heißt **äquidistant**, falls alle n Teilintervalle gleich lang sind, also die Länge $(b - a)/n$ haben. D.h.

$$x_k = a + k \frac{b - a}{n} \text{ für } k = 1, \dots, n.$$

1.3. Ober- und Untersummen. Es sei $f : [a, b] \rightarrow \mathbb{R}$ eine beschränkte Funktion.

Für eine Zerlegung $\mathcal{Z}$ wie in Abschnitt 1.2 seien m_k das Infimum und M_k das Supremum der Funktion f auf dem Intervall I_k , also

$$m_k := \inf\{f(x) : x \in [x_{k-1}, x_k]\} \text{ und } M_k := \sup\{f(x) : x \in [x_{k-1}, x_k]\}.$$

Die **Obersumme** $O_{\mathcal{Z}}$ und **Untersumme** $U_{\mathcal{Z}}$ der Funktion f auf dem Intervall $[a, b]$ bezüglich der Zerlegung $\mathcal{Z}$ sind

$$(63) \quad O_{\mathcal{Z}} := \sum_{k=1}^n M_k (x_k - x_{k-1}), \quad U_{\mathcal{Z}} := \sum_{k=1}^n m_k (x_k - x_{k-1}).$$

Eine beschränkte Funktion $f : [a, b] \rightarrow \mathbb{R}$ heißt auf dem Intervall $[a, b]$ **Riemann-integrierbar**, wenn das Infimum aller Obersummen gleich dem Supremum aller Untersummen ist. Dieser Wert wird mit $\int_a^b f(x) dx$ bezeichnet und **bestimmtes Integral** der Funktion f auf dem Intervall $[a, b]$ genannt:

$$\int_a^b f(x) dx := \inf\{O_{\mathcal{Z}} : \mathcal{Z} \text{ Zerlegung}\} = \sup\{U_{\mathcal{Z}} : \mathcal{Z} \text{ Zerlegung}\}$$

BEISPIEL 81 (Integral der konstanten Funktion). Für jede konstante Funktion $f \equiv c \in \mathbb{R}$ und alle $a < b \in \mathbb{R}$ gilt $m_k = M_k = c$ für alle Zerlegungen $\mathcal{Z}$, also

$$U_{\mathcal{Z}} = \sum_{k=1}^n m_k (x_k - x_{k-1}) = \sum_{k=1}^n c (x_k - x_{k-1}) = c(b - a)$$

und

$$O_{\mathcal{Z}} = \sum_{k=1}^n M_k (x_k - x_{k-1}) = \sum_{k=1}^n c (x_k - x_{k-1}) = c(b - a).$$

Aus $\{O_Z : Z \text{ Zerlegung}\} = \{U_Z : Z \text{ Zerlegung}\} = \{c(b-a)\}$ folgt $\int_a^b c dx = c(b-a)$. Dies stimmt mit der Motivation des Integrals als Flächeninhalt (siehe Abbildung 1) überein.

BEMERKUNG 54. Das Integralzeichen $\int$ erinnert an den Buchstaben „S“ und damit daran, dass integrieren etwas mit „Summieren“ zu tun hat. Die Bezeichnung dx macht deutlich, dass man zur Definition des Integrals das Intervall in dem die unabhängige Variable x liegt, in immer kleinere Teilintervalle, Δx , zerlegt hat. $\square$

SATZ 62. *Es sei $f : [a, b] \rightarrow \mathbb{R}$ eine beschränkte Funktion. Die Funktion f ist genau dann auf $[a, b]$ Riemann-integrierbar, wenn es zu jedem $\varepsilon > 0$ eine Zerlegung Z mit $O_Z - U_Z < \varepsilon$ gibt.*

SATZ 63. *Wenn f auf $[a, b]$ Riemann-integrierbar ist, dann gilt*

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} U_{Z_n} = \lim_{n \rightarrow \infty} O_{Z_n}.$$

BEWEIS. Wir zeigen $\lim_{n \rightarrow \infty} U_{Z_n} = \int_a^b f(x) dx =: R$.

Es sei $\varepsilon > 0$. Es existieren eine Zerlegung $Z = (a = x_0, x_1, \dots, x_N = b)$ mit $O_Z - U_Z < \varepsilon/3$ und ein $n_0 \in \mathbb{N}$ mit

$$\frac{b-a}{n_0} \leq x_k - x_{k-1} \quad \forall k = 1, \dots, N \quad \text{und} \quad \frac{b-a}{n_0} N(M-m) < \frac{\varepsilon}{3}$$

für $m := \inf\{f(x) : x \in [a, b]\}$, $M := \sup\{f(x) : x \in [a, b]\}$.

Für alle $n \geq n_0$ sind höchstens N Teilintervalle der äquidistanten Zerlegung Z_n nicht vollständig in einem Teilintervall der Zerlegung Z enthalten. Damit gilt

$$R - U_{Z_n} \leq O_{Z_n} - U_{Z_n} \leq N \frac{b-a}{n} (M-n) + O_Z - U_Z \leq \frac{\varepsilon}{3} + \frac{\varepsilon}{3} < \varepsilon.$$

$\square$

1.4. Zwischensummen. Für eine Funktion $f : [a, b] \rightarrow \mathbb{R}$, eine Zerlegung Z wie in Abschnitt 1.2 und $\xi = (\xi_1, \dots, \xi_n)$ mit $\xi_k \in I_k$ heißt

$$S(Z, \xi) := \sum_{j=1}^n f(\xi_k)(x_k - x_{k-1})$$

Zwischensumme der Funktion f bezüglich der Zerlegung Z und Zwischenvektor ξ .

Eine Folge $(Z_n, \xi^{(n)})_{n \in \mathbb{N}}$ von Zerlegungen Z_n mit passenden Zwischenvektoren $\xi^{(n)}$ heißt **Zerlegungsnullfolge**, wenn $\lim_{n \rightarrow \infty} |Z_n| = 0$.

SATZ 64. *Es sei $f : [a, b] \rightarrow \mathbb{R}$.*

Die Funktion f ist genau dann auf $[a, b]$ beschränkt und Riemann-integrierbar, wenn für jede Zerlegungsnullfolge die dazugehörige Folge der Zwischensummen konvergiert.

1.5. Integral für stetige Funktionen.

SATZ 65. Wenn $f : [a, b] \rightarrow \mathbb{R}$ stetig, dann ist f auf $[a, b]$ Riemann-integrierbar.

BEISPIEL 82 (Integral der Funktion $f(x) = x$). Falls $f(x) = x$, so gilt $m_k = x_{k-1}$ und $M_k = x_k$ für alle Zerlegungen $\mathcal{Z}$. Für jede äquidistante Zerlegung $\mathcal{Z}_n$ folgt

$$\begin{aligned} O_{\mathcal{Z}_n} &= \sum_{k=1}^n x_k \frac{b-a}{n} = \frac{b-a}{n} \sum_{k=1}^n \left(a + k \frac{b-a}{n} \right) = \frac{b-a}{n} \left(na + \frac{b-a}{n} \sum_{k=1}^n k \right) \\ &= \frac{b-a}{n} \left(na + \frac{b-a}{n} \cdot \frac{n(n+1)}{2} \right) = (b-a) \left(a + \frac{b-a}{2} \cdot \frac{n+1}{n} \right) \\ &\xrightarrow{n \rightarrow \infty} (b-a) \left(a + \frac{b-a}{2} \right) = \frac{1}{2}(b-a)(b+a) = \frac{1}{2}(b^2 - a^2) = \int_a^b x dx. \end{aligned}$$

Dies stimmt mit der Interpretation des Integrals als Flächeninhalt überein (siehe Abbildung 2).

2. Eigenschaften des Integrals

Interpretiert man das bestimmte Integral als Flächeninhalt, so sind folgende Eigenschaften des Integrals anschaulich sofort klar.

SATZ 66. Wenn f auf den Intervallen $[a, c]$ und $[c, b]$ integrierbar ist, dann ist f auch auf dem Intervall $[a, b]$ integrierbar und es gilt

$$(64) \quad \int_a^b f(x) dx = \int_a^c f(x) dx + \int_c^b f(x) dx.$$

SATZ 67. Wenn f und g auf dem Intervall $[a, b]$ integrierbar sind, dann sind auch $f + g$ und cf für alle $c \in \mathbb{R}$ auf dem Intervall $[a, b]$ integrierbar und es gilt

$$(65) \quad \int_a^b (f+g)(x) dx = \int_a^b f(x) dx + \int_a^b g(x) dx \text{ und } \int_a^b (cf)(x) dx = c \int_a^b f(x) dx.$$

BEMERKUNG 55. Will man $\int_a^b f(x) dx$ auch für $a > b$ konsistent definieren, so folgt aus Satz 66

$$(66) \quad \int_b^a f(x) dx = - \int_a^b f(x) dx \text{ und } \int_a^a f(x) dx = 0 \quad \forall a, b \in \mathbb{R}. \quad \square$$

BEMERKUNG 56. Falls $f(x) < 0$ für alle $x \in [a, b]$, so gilt $-f(x) > 0$ für alle $x \in [a, b]$. Aus Satz 67 folgt $\int_a^b f(x) dx = - \int_a^b -f(x) dx < 0$ für $a < b$. Dies kann nur mit der Vorstellung des Integrals als Maß für den Flächeninhalt vereinbart werden, wenn man den Flächeninhalt der Teile „unter“ der x -Achse mit negativem Vorzeichen versieht. $\square$

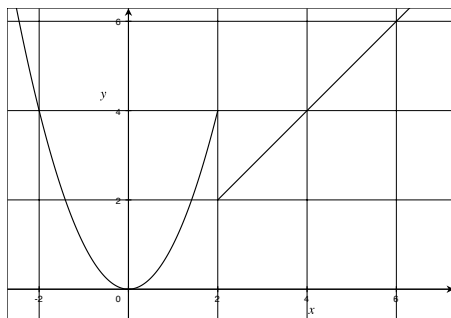


ABBILDUNG 5. Sprungstelle

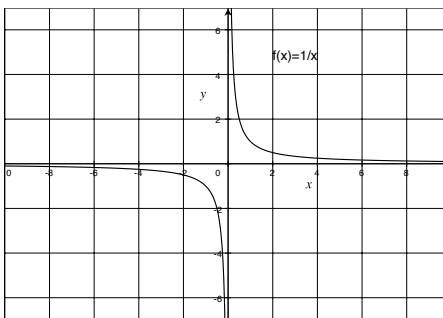


ABBILDUNG 6. Keine Sprungstelle

LEMMA 18. Wenn $f : [a, b] \rightarrow \mathbb{R}$ stetig auf dem offenen Intervall (a, b) ist und die Grenzwerte

$$\lim_{x \rightarrow a, x > a} f(x) \text{ und } \lim_{x \rightarrow b, x < b} f(x)$$

existieren, dann ist f auf $[a, b]$ Riemann-integrierbar und es gilt

$$\int_a^b f(x) dx = \int_a^b \tilde{f}(x) dx$$

für die Funktion $\tilde{f} : [a, b] \rightarrow \mathbb{R}$ mit $\tilde{f}(x) = f(x)$ für $a < x < b$,

$$\tilde{f}(a) = \lim_{x \rightarrow a, x > a} f(x) \text{ und } \tilde{f}(b) = \lim_{x \rightarrow b, x < b} f(x).$$

BEMERKUNG 57. Die Funktion $\tilde{f}$ in Lemma 18 ist die Fortsetzung der auf dem offenen Intervall (a, b) stetigen Funktion $f : (a, b) \rightarrow \mathbb{R}$ zu einer stetigen Funktion auf dem abgeschlossenen Intervall $[a, b]$. $\square$

Eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ heißt **stückweise stetig** auf dem Intervall $[a, b]$, wenn sie nur im Intervall $[a, b]$ nur endlich viele Unstetigkeitsstellen x_k besitzt und diese Unstetigkeitsstellen „gutartig“ sind, d.h. für alle x_k existieren die Grenzwerte

$$\lim_{x \rightarrow x_k, x > x_k} f(x) \text{ und } \lim_{x \rightarrow x_k, x < x_k} f(x)$$

und sind von $\pm\infty$ verschieden. Eine solche Unstetigkeitsstelle nennt man **Sprungstelle**.

LEMMA 19. Eine auf einem abgeschlossenen Intervall stückweise stetige Funktion ist auf diesem Intervall beschränkt.

BEISPIEL 83 (Sprungstelle). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$ für $x \leq 2$ und $f(x) = x$ für $x > 2$ ist stückweise stetig auf jedem abgeschlossenen Intervall, denn sie besitzt nur die Unstetigkeitsstelle $x_1 = 2$. Es gilt

$$\lim_{x \rightarrow 2, x < 2} f(x) = \lim_{x \rightarrow 2, x < 2} x^2 = 4 \text{ und } \lim_{x \rightarrow 2, x > 2} f(x) = \lim_{x \rightarrow 2, x > 2} x = 2.$$

Die beiden Grenzwerte existieren, aber sie sind nicht gleich. Abbildung 5 zeigt den Graph der Funktion.

BEISPIEL 84. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = 1/x$ für $x \neq 0$ und $f(0) = 0$ ist stetig in allen Punkten $x \neq 0$ und besitzt eine Unstetigkeitsstelle in $x = 0$, denn

$$\lim_{x \rightarrow 0, x < 0} f(x) = \lim_{x \rightarrow 0, x < 0} \frac{1}{x} = -\infty \text{ und } \lim_{x \rightarrow 0, x > 0} f(x) = \lim_{x \rightarrow 0, x > 0} \frac{1}{x} = \infty.$$

Die beiden Grenzwerte sind keine reellen Zahlen. D.h. $x = 0$ ist keine Sprungstelle. Abbildung 6 zeigt den Graph der Funktion.

BEISPIEL 85. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = \sin(1/x)$ für $x \neq 0$ und $f(0) = 0$ ist stetig in allen Punkten $x \neq 0$ und besitzt eine Unstetigkeitsstelle in $x = 0$ (siehe Beispiel 51). Die Grenzwerte

$$\lim_{x \rightarrow 0, x < 0} f(x) = \lim_{x \rightarrow 0, x < 0} \sin\left(\frac{1}{x}\right) \text{ und } \lim_{x \rightarrow 0, x > 0} f(x) = \lim_{x \rightarrow 0, x > 0} \sin\left(\frac{1}{x}\right)$$

existieren nicht. D.h. $x = 0$ ist keine Sprungstelle.

SATZ 68. Eine auf dem Intervall $[a, b]$ stückweise stetige Funktion f ist auf $[a, b]$ Riemann-integrierbar und es gilt

$$\int_a^b f(x) dx = \int_a^{x_1} f(x) dx + \int_{x_1}^{x_2} f(x) dx + \dots + \int_{x_m}^b f(x) dx,$$

falls $x_1 < x_2 < \dots < x_m$ die Sprungstellen von f auf dem Intervall $[a, b]$ sind.

BEWEIS. Die Behauptung folgt aus Satz 66 und aus Lemma 18. $\square$

Eine Menge $M \subset \mathbb{R}$ heißt **Nullmenge**, falls zu jedem $\varepsilon > 0$ Intervalle I_n , $n \in \mathbb{N}$, mit $M \subset \cup_n I_n$ und $\sum_n |I_n| < \varepsilon$ existieren.

Eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ heißt **fast überall stetig**, falls die Menge der Unstetigkeitsstellen der Funktion f eine Nullmenge ist.

SATZ 69. Eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ ist genau dann Riemann-integrierbar auf $[a, b]$, falls f beschränkt und fast überall stetig ist.

FOLGERUNG 24. Wenn f und g auf $[a, b]$ Riemann-integrierbar sind, so sind auch die Funktionen

$$|f|, f^+ := \frac{1}{2}(f + |f|), f^- := f - f^+, \max(f, g), \min(f, g) \text{ und } fg$$

Riemann-integrierbar auf $[a, b]$.

3. Uneigentliche Integrale

Eine Funktion f heißt **stückweise stetig** auf dem offenen Intervall (a, b) , wenn sie auf jedem abgeschlossenen Intervall, das in (a, b) enthalten ist, stückweise stetig ist. Zum Beispiel ist die Funktion $f(x) = 1/x$ auf den Intervallen $(0, \infty)$ und $(-\infty, 0)$ stetig und damit auch stückweise stetig.

Das **uneigentliche Integral** einer auf $[a, \infty)$ stückweise stetigen Funktion f ist

$$\int_a^\infty f(x)dx := \lim_{b \rightarrow \infty} \int_a^b f(x)dx,$$

falls dieser Grenzwert existiert und verschieden von $\pm\infty$ ist.

Das **uneigentliche Integral** einer auf $(-\infty, b]$ stückweise stetigen Funktion f ist

$$\int_{-\infty}^b f(x)dx := \lim_{a \rightarrow -\infty} \int_a^b f(x)dx,$$

falls dieser Grenzwert existiert und verschieden von $\pm\infty$ ist.

Das **uneigentliche Integral** einer auf $[a, b_0)$ stückweise stetigen, unbeschränkten Funktion f ist, falls dieser Grenzwert existiert und verschieden von $\pm\infty$ ist,

$$\int_a^{b_0} f(x)dx := \lim_{b \rightarrow b_0} \int_a^b f(x)dx.$$

Das **uneigentliche Integral** einer auf $(a_0, b]$ stückweise stetigen, unbeschränkten Funktion f ist, falls dieser Grenzwert existiert und verschieden von $\pm\infty$ ist,

$$\int_{a_0}^b f(x)dx := \lim_{a \rightarrow a_0} \int_a^b f(x)dx.$$

BEMERKUNG 58. Das uneigentliche Integrale misst den Flächeninhalt einer unbeschränkten Fläche (siehe Beispiele in Abschnitt 4). $\square$

4. Hauptsatz der Differential- und Integralrechnung

Ist eine Funktion f auf dem Intervall $[a, b]$ Riemann-integrierbar und $x_0 \in [a, b]$, dann wird durch

$$(67) \quad F(x) := \int_{x_0}^x f(t) dt$$

eine Funktion $F : [a, b] \rightarrow \mathbb{R}$ definiert. Die Funktion F beschreibt, wie sich der Flächeninhalt ändert, wenn man die obere Intervallgrenze des Integrals ändert.

BEISPIEL 86. Für $f \equiv c$ und $x_0 = 0$ gilt $F(x) = \int_0^x c dt = cx$.

BEISPIEL 87. Für $f(x) = x$ und $x_0 = 1$ gilt $F(x) = \int_1^x t dt = \frac{1}{2}(x^2 - 1^2) = \frac{1}{2}(x^2 - 1)$.

BEISPIEL 88 (Integration glättet Knick). Für $f(x) = |x|$ und $x_0 = 0$ berechnen wir die Stammfunktion wie in Gleichung 67. Für $t \geq 0$ gilt $|t| = t$, also gilt für $x \geq 0$

$$F(x) = \int_0^x |t| dt = \int_0^x t dt = \frac{1}{2}x^2.$$

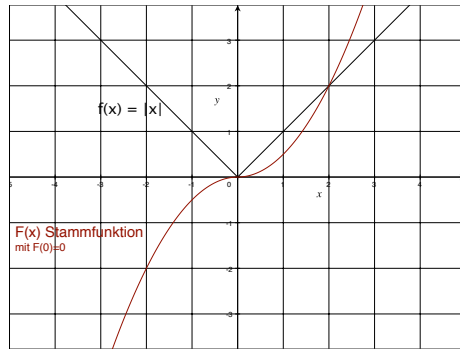


ABBILDUNG 7

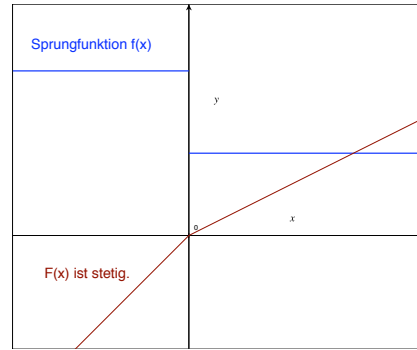


ABBILDUNG 8

Für $t \leq 0$ gilt $|t| = -t$, also gilt für $x < 0$

$$F(x) = \int_0^x |t| dt = \int_0^x -t dt = -\frac{1}{2}x^2.$$

Die Funktion $f(x) = |x|$ ist stetig für alle $x \in \mathbb{R}$, aber nicht differenzierbar in $x = 0$. Die Stammfunktion $F(x)$ ist für alle $x \in \mathbb{R}$ differenzierbar. Abbildung 7 zeigt die Graphen der Funktionen $f(x) = |x|$ und der Stammfunktion $F(x)$.

BEISPIEL 89 (Integration beseitigt Sprungstelle). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = 1$ für alle $x \leq 0$ und $f(x) = 1/2$ für alle $x > 0$ ist stückweise stetig. Sie besitzt eine Unstetigkeit in $x = 0$. Wir berechnen $F(x)$ wie in Gleichung 67. Für $t \leq 0$ gilt $f(t) = 1$, also gilt für $x \leq 0$

$$F(x) = \int_0^x f(t) dt = \int_0^x 1 dt = x.$$

Für $t \geq 0$ gilt $f(t) = 1/2$, also gilt für $x > 0$

$$F(x) = \int_0^x f(t) dt = \int_0^x \frac{1}{2} dt = \frac{1}{2}x.$$

Die Funktion $F(x)$ ist für alle $x \in \mathbb{R}$ stetig. Sie ist in $x = 0$ nicht differenzierbar. Abbildung 8 zeigt die Graphen der unstetigen Funktion f und der stetigen „Stammfunktion“ F .

SATZ 70. Wenn f auf dem Intervall (a, b) stetig ist und $x_0 \in [a, b]$, dann ist

$$F(x) := \int_{x_0}^x f(t) dt$$

auf dem Intervall (a, b) differenzierbar und es gilt $F'(x) = f(x)$ für alle $x \in (a, b)$.

Eine **Stammfunktion** der Funktion $f : (a, b) \rightarrow \mathbb{R}$ ist eine differenzierbare Funktion $F : (a, b) \rightarrow \mathbb{R}$ mit $F'(x) = f(x)$ für alle $x \in (a, b)$. Eine Stammfunktion der Funktion f wird mit $\int f(x)dx$ bezeichnet und auch **unbestimmtes Integral** genannt.

Ist F eine Stammfunktion, so für jedes $c \in \mathbb{R}$ auch $F + c$ eine Stammfunktion, denn $(F + c)'(x) = F'(x)$.

LEMMA 20. Sind F und G Stammfunktionen einer Funktion $f : (a, b) \rightarrow \mathbb{R}$, so existiert eine Konstante $c \in \mathbb{R}$ mit $F(x) = G(x) + c$ für alle $x \in (a, b)$.

BEWEIS. Die Funktion $H := F - G$ ist differenzierbar. Da $H'(x) = F'(x) - G'(x) = f(x) - f(x) = 0$ für alle $x \in (a, b)$, liefert Folgerung 23 $H \equiv c$. $\square$

Für jede stetige Funktion f liefert Gleichung 67 eine Stammfunktion. Diese hängt von der Wahl des Startpunktes x_0 ab.

SATZ 71 (Hauptsatz der Differential- und Integralrechnung). Wenn F ein Stammfunktion von f auf einem Intervall I ist, dann gilt

$$(68) \quad \int_a^b f(x)dx = F(b) - F(a) \text{ für alle Intervalle } [a, b] \subset I.$$

BEWEIS. Die Funktion $G(x) := \int_a^x f(t) dt$ ist auch eine Stammfunktion von f . Darum gilt $F = G + c$. Daraus folgt, wegen $G(a) = 0$,

$$\int_a^b f(x)dx = G(b) = G(b) - G(a) = F(b) - c - (F(a) - c) = F(b) - F(a).$$

$\square$

FOLGERUNG 25 (Mittelwertsatz der Integralrechnung). Wenn $f : [a, b] \rightarrow \mathbb{R}$ stetig ist, dann existiert ein $\xi \in [a, b]$ mit

$$(69) \quad \int_a^b f(x)dx = f(\xi)(b - a).$$

BEWEIS. Die Funktion f besitzt eine Stammfunktion F . Da F auf dem Intervall $[a, b]$ differenzierbar ist, existiert ein $\xi \in [a, b]$ mit $(b - a)F'(\xi) = F(b) - F(a)$ (Mittelwertsatz der Differentialrechnung). Die Behauptung folgt wegen $F'(\xi) = f(\xi)$ und $F(b) - F(a) = \int_a^b f(x)dx$. $\square$

BEISPIEL 90 (Stammfunktion von Monomen). Da $(x^{n+1})' = (n + 1)x^n$ für alle $n \in \mathbb{R}$, folgt für alle $n \neq -1$

$$(70) \quad \int_a^b x^n dx = \left[\frac{1}{n+1} x^{n+1} \right]_a^b = \frac{1}{n+1} (b^{n+1} - a^{n+1}).$$

BEISPIEL 91 (Stammfunktion von e^{cx+d}). Da $(e^{cx+d})' = ce^{cx+d}$ für alle $c, d \in \mathbb{R}$, folgt für alle $c \neq 0$

$$(71) \quad \int_a^b e^{cx+d} dx = \left[\frac{1}{c} e^{cx+d} \right]_a^b = \frac{1}{c} (e^{cb+d} - e^{ca+d}).$$

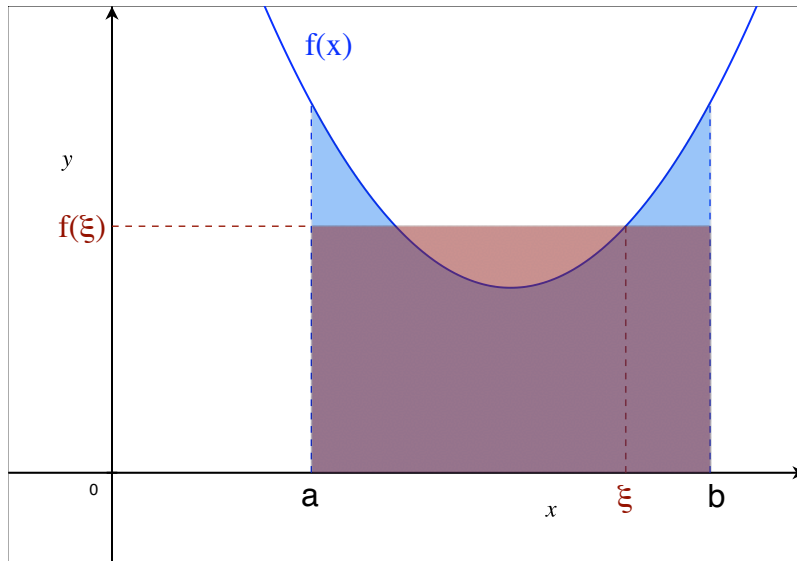


ABBILDUNG 9. Mittelwertsatz der Integralrechnung

BEISPIEL 92 (Stammfunktion von $\cos(cx + d)$). Da $(\sin(cx + d))' = c \cos(cx + d)$ für alle $c, d \in \mathbb{R}$, folgt für alle $c \neq 0$

$$(72) \quad \int_a^b \cos(cx + d) dx = \left[\frac{1}{c} \sin(cx + d) \right]_a^b = \frac{1}{c} (\sin(cb + d) - \sin(ca + d)).$$

BEISPIEL 93 (Stammfunktion von $\sin(cx + d)$). Da $(\cos(cx + d))' = -c \sin(cx + d)$ für alle $c, d \in \mathbb{R}$, folgt für alle $c \neq 0$

$$(73) \quad \int_a^b \sin(cx + d) dx = \left[-\frac{1}{c} \cos(cx + d) \right]_a^b = \frac{1}{c} (-\cos(cb + d) + \cos(ca + d)).$$

BEISPIEL 94 (Stammfunktion von $\frac{1}{1+x^2}$). Da $(\arctan x)' = \frac{1}{1+x^2}$ folgt

$$(74) \quad \int_a^b \frac{1}{1+x^2} dx = [\arctan x]_a^b = \arctan b - \arctan a.$$

Wir benutzen die Grundintegrale aus den obigen Beispielen, um uneigentliche Integrale zu untersuchen.

BEISPIEL 95. Für $b > 1$ gilt

$$\int_1^b \frac{1}{x^2} dx = \int_1^b x^{-2} dx = [-x^{-1}]_1^b = -b^{-1} + 1^{-1} = 1 - \frac{1}{b}.$$

Daraus folgt

$$\int_1^\infty \frac{1}{x^2} dx = \lim_{b \rightarrow \infty} \int_1^b \frac{1}{x^2} dx = \lim_{b \rightarrow \infty} \left(1 - \frac{1}{b} \right) = 1.$$

BEISPIEL 96. Für $b > 0$ gilt

$$\int_0^b \sin(2x) dx = \left[-\frac{1}{2} \cos(2x) \right]_1^b = -\frac{\cos(2b)}{2} + \frac{\cos(0)}{2} = \frac{1}{2}(1 - \cos(2b)).$$

Da sich $\cos(2b)$ für $b \rightarrow \infty$ nicht einem Grenzwert nähert, existiert auch das uneigentliche Integral $\int_0^\infty \sin(2x) dx$ nicht.

BEISPIEL 97. Für $a > 0$ gilt

$$\int_a^1 x^{-1/3} dx = \left[\frac{3}{2} x^{2/3} \right]_a^1 = \frac{3}{2} (1^{2/3} - a^{2/3}) = \frac{3}{2} (1 - a^{2/3}).$$

Daraus folgt

$$\int_0^1 x^{-1/3} dx = \frac{3}{2}, \text{ da } \lim_{a \rightarrow 0} a^{2/3} = 0.$$

Viele Stammfunktionen lassen sich explizit mit Hilfe der Ableitungsregeln und den Sätzen in den Abschnitten 6 und 7 finden. Jedoch kann man die Stammfunktionen mancher stetiger Funktionen wie z.B. e^{-x^2} nicht in einer geschlossenen Formel angeben, obwohl diese Stammfunktionen natürlich existieren und differenzierbar sind. Die Funktionswerte einer solchen Stammfunktion lassen sich dann numerisch, z.B. durch Berechnung von Ober- bzw. Untersummen, bestimmen.

5. Der Logarithmus und die Exponentialfunktion

In diesem Abschnitt definieren wir die Funktion $\ln x$ als eine Stammfunktion der Funktion $1/x$ und die Exponentialfunktion als Umkehrfunktion des natürlichen Logarithmus.

5.1. Definition und Grundeigenschaften des Logarithmus. Die Funktion $f(x) = \frac{1}{x}$ ist für $x > 0$ stetig. Für alle $b > 0$ existiert das bestimmte Integral $\int_1^b \frac{1}{x} dx$. Wir definieren den natürlichen **Logarithmus** als

$$(75) \quad \ln x := \int_1^x \frac{1}{t} dt \text{ für } x > 0.$$

Aus dieser Definition erhält man folgende Eigenschaften

- $\ln 1 = 0$, denn $\ln 1 = \int_1^1 \frac{1}{x} dx = 0$.
- $(\ln x)' = \frac{1}{x}$
- $\ln x$ ist streng monoton wachsend, denn $(\ln x)' = \frac{1}{x} > 0$ für alle $x > 0$.
- $\ln x > 0$ für $x > 1$, denn $\frac{1}{t} > 0$ für alle $t > 0$.
- $\ln x < 0$ für alle $x < 1$, denn $\int_1^x f(t) dt = -\int_x^1 f(t) dt$.

BEISPIEL 98. Aus der Kettenregel folgt $(\ln(cx + d))' = \frac{c}{cx+d}$ für alle $c, d \in \mathbb{R}$. Aus dem Hauptsatz der Differential- und Integralrechnung erhält

man für $c \neq 0$, falls $cx + d \neq 0$ für alle $x \in [a, b]$

$$(76) \quad \int_a^b \frac{1}{cx+d} dx = \frac{1}{c} [\ln(cx+d)]_a^b = \frac{1}{c} (\ln(ac+d) - \ln(bc+d)) \\ = \frac{1}{c} \ln \frac{ac+d}{bc+d}.$$

Für die letzte Umformung benutzt man die Logarithmusgesetze aus Abschnitt 5.4.

BEISPIEL 99. Für $x > 0$ gilt $(\ln x)' = 1/x$. Aus der Kettenregel folgt für $x < 0$

$$(\ln(-x))' = \frac{1}{-x} \cdot (-1) = \frac{1}{x}, \text{ also } \int \frac{1}{x} dx = \ln|x| \text{ für } x \neq 0.$$

BEISPIEL 100. Für $x \neq 0, 1$ gilt

$$\frac{1}{x} - \frac{1}{x+1} = \frac{1}{x(x+1)}, \text{ also ist } \int \frac{1}{x(x+1)} dx = \ln x - \ln(x+1) = \ln \frac{x}{x+1}$$

eine Stammfunktion für $x > 0$. Für die letzte Umformung benutzt man die Logarithmusgesetze aus Abschnitt 5.4.

5.2. $\ln x$ für $x \rightarrow 0$ und $x \rightarrow \infty$, Harmonische Reihe. Wir berechnen die Untersumme des Integrals $\int_1^n 1/x dx$ bezüglich der äquidistanten Zerlegung $\mathcal{Z}$ des Intervalls $[1, n]$ in $n-1$ Teilintervalle der Länge 1 für alle $n \in \mathbb{N}$ mit $n > 0$. Da die Funktion $1/x$ streng monoton fallend ist, nimmt sie auf jedem Intervall $[x_{k-1}, x_k]$ ihr Minimum in x_k an. Es gilt $x_k = 1+k$, $x_k - x_{k-1} = 1$ und

$$\ln n = \int_1^n \frac{1}{x} dx \geq U_{\mathcal{Z}} = \sum_{k=1}^{n-1} \frac{1}{1+k} = \sum_{k=2}^n \frac{1}{k}.$$

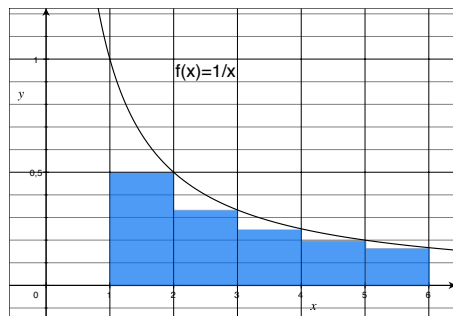
Wir berechnen die Untersumme des Integrals $\int_{1/n}^1 1/x dx$ bezüglich der äquidistanten Zerlegung $\mathcal{Z}$ des Intervalls $[1/n, 1]$ in $n-1$ Teilintervalle der Länge $1/n$ für alle $n \in \mathbb{N}$ mit $n > 0$. Da die Funktion $1/x$ streng monoton fallend ist, nimmt sie auf jedem Intervall $[x_{k-1}, x_k]$ ihr Maximum in x_{k-1} an. Es gilt $x_k = (k+1)/n$, $x_k - x_{k-1} = 1/n$ und

$$\ln \frac{1}{n} = \int_{1/n}^1 \frac{1}{x} dx = - \int_{1/n}^1 \frac{1}{x} dx \leq -U_{\mathcal{Z}} = - \sum_{k=1}^{n-1} \frac{1}{\frac{1+k}{n}} \frac{1}{n} \\ = - \sum_{k=1}^{n-1} \frac{1}{1+k} = - \sum_{k=2}^n \frac{1}{k}.$$

In beiden Fällen spielt die Summenfolge $\{s_n\}_{n \in \mathbb{N}}$ mit

$$(77) \quad s_n = \sum_{k=1}^n \frac{1}{k}$$

eine besondere Rolle. Sie heißt **harmonische Reihe**.



ABBILDUNG

$$10. \ln n \leq \sum_{k=2}^n \frac{1}{k}$$

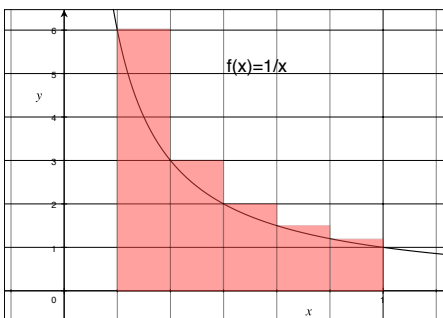


ABBILDUNG 11

LEMMA 21. Die harmonische Reihe divergiert, d.h. sie ist streng monoton wachsend und nach oben unbeschränkt.

BEWEIS. Für $n = 2^m$ gilt

$$\begin{aligned} \sum_{k=1}^{2^m} \frac{1}{k} &= 1 + \sum_{k=2}^{2^m} \frac{1}{k} = 1 + \sum_{l=1}^m \sum_{k=2^{l-1}+1}^{2^l} \frac{1}{k} \\ &\geq 1 + \sum_{l=1}^m \sum_{k=2^{l-1}+1}^{2^l} \frac{1}{2^l} = 1 + \sum_{l=1}^m \frac{2^{l-1}}{2^l} = 1 + \frac{m}{2}. \end{aligned}$$

□

FOLGERUNG 26. Es gilt $\lim_{x \rightarrow 0} \ln x = -\infty$ und $\lim_{x \rightarrow \infty} \ln x = \infty$.

BEISPIEL 101. Das uneigentliche Integral $\int_0^1 \frac{1}{x} dx$ existiert nicht, denn für $a > 0$ gilt

$$\int_a^1 \frac{1}{x} dx = - \int_1^a \frac{1}{x} dx = -[\ln x]_1^a = -\ln a + \ln 1 = -\ln a \xrightarrow{a \rightarrow 0} \infty.$$

5.3. Die Exponentialfunktion als Umkehrfunktion des Logarithmus. Da $(\ln x)' = 1/x > 0$ für alle $x > 0$, ist $\ln x$ streng monoton wachsend und besitzt eine Umkehrfunktion. Diese wird mit e^x bezeichnet. Da $\{\ln x : x > 0\} = \mathbb{R}$ ist die **Exponentialfunktion** e^x für alle $x \in \mathbb{R}$ definiert. Aus $\ln(e^x) = x$ folgt durch Ableitung der Gleichung mit Hilfe der Kettenregel

$$1 = (\ln(e^x))' = \frac{1}{e^x} (e^x)', \text{ also } (e^x)' = e^x.$$

Aus $e^{\ln x} = x$ folgt $e^0 = 1$ und $e^x > 0$ für alle $x \in \mathbb{R}$.

5.4. Die Logarithmus- und Potenzgesetze.

SATZ 72. Für alle $x, y \in \mathbb{R}$ mit $x, y > 0$ gilt

$$(78) \quad \ln(xy) = \ln x + \ln y.$$

BEWEIS. Für jedes $y_0 > 0$ definiert man die Funktion $f(x) := \ln(xy_0) - \ln x - \ln y_0$. Für alle $x > 0$ gilt

$$f'(x) = \frac{1}{xy_0}y_0 - \frac{1}{x} = 0.$$

Daraus folgt, dass f eine konstante Funktion ist. Da $f(1) = \ln(y_0) - \ln 1 - \ln y_0 = 0$ ist, gilt $f \equiv 0$. $\square$

SATZ 73. Für alle $x \in \mathbb{R}$ mit $x > 0$ gilt

$$(79) \quad \ln \frac{1}{x} = -\ln x.$$

BEWEIS. Für die Funktion $f(x) := \ln(1/x) + \ln x$ folgt aus der Kettenregel für alle $x > 0$

$$f'(x) = \frac{1}{\frac{1}{x}} \cdot \left(-\frac{1}{x^2}\right) + \frac{1}{x} = 0.$$

Also ist f eine konstante Funktion. Da $f(1) = \ln 1 - \ln 1 = 0$ ist, gilt $f \equiv 0$. $\square$

FOLGERUNG 27. Für alle $x, y \in \mathbb{R}$ gilt $e^{x+y} = e^x e^y$.

BEWEIS. Da $\ln x$ die Umkehrfunktion zu e^x ist, folgt aus den Logarithmusgesetzen $\ln(e^x e^y) = \ln(e^x) + \ln(e^y) = x + y$. Wendet man auf diese Gleichung die Exponentialfunktion an, so erhält man $e^x e^y = e^{x+y}$. $\square$

5.5. Die Funktionen a^x und $\log_a x$. Für alle $a \in \mathbb{R}$ mit $a > 0$ gilt $a = e^{\ln a}$. Für alle $n \in \mathbb{N}$ folgt $a^n = e^{n \ln a}$. Wir definieren die **Potenzen**

$$(80) \quad a^x := e^{x \ln a} \text{ für alle } x \in \mathbb{R}.$$

Es gelten die Potenzgesetze, denn

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

und die Ableitung der Funktion $f(x) = a^x$ lässt sich leicht mit Hilfe der Kettenregel bestimmen:

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

Für $0 < a < 1$ ist a^x streng monoton fallend, denn $\ln a < 0$. Für $a > 1$ ist a^x streng monoton wachsend, denn $\ln a > 0$. Für alle $a > 0$ mit $a \neq 1$ gilt $\{a^x : x \in \mathbb{R}\} = (0, \infty)$.

Die Umkehrfunktion der Funktion a^x heißt **Logarithmus zur Basis a** und wird mit $\log_a x$ bezeichnet. Da $x = a^{\log_a x}$ für alle $x > 0$, folgt aus der Kettenregel

$$1 = a^{\log_a x} \ln a (\log_a x)' = x \ln a (\log_a x)', \text{ also } (\log_a x)' = \frac{1}{x \ln a}.$$

6. Kettenregel und partielle Integration

Die Kettenregel und die Produktregel für die Ableitung lassen sich in Regeln für die Bestimmung von Integralen übersetzen. Allerdings geben diese Integrationsregeln keine exakten Anweisungen für die Zerlegung einer zu integrierenden Funktion in passende Faktoren.

Aus der Kettenregel folgt, dass $f \circ g$ eine Stammfunktion von $f'(g(x))g'(x)$ ist.

BEISPIEL 102. Gesucht ist eine Stammfunktion von xe^{-x^2} . Es gilt

$$xe^{-x^2} = -\frac{1}{2}e^{-x^2}(-2x), \text{ also } \int xe^{-x^2} dx = -\frac{1}{2} \int e^{-x^2}(-2x) dx.$$

Wählen wir $g(x) = -x^2$ und $g'(x) = -2x$, so muss $f'(x) = e^x$ sein. Daraus folgt $f(x) = e^x$ und

$$\int xe^{-x^2} dx = -\frac{1}{2}f(g(x)) = -\frac{1}{2}e^{-x^2}.$$

BEISPIEL 103. Gesucht ist eine Stammfunktion von $\frac{x}{x^2+1}$. Es gilt

$$\frac{x}{x^2+1} = \frac{1}{2} \frac{2x}{x^2+1}, \text{ also } \int \frac{x}{x^2+1} dx = \frac{1}{2} \int \frac{2x}{x^2+1} dx = \frac{1}{2} \int \frac{1}{x^2+1} \cdot (2x) dx.$$

Wählen wir $g(x) = x^2 + 1$ und $g'(x) = 2x$, so muss $f'(x) = 1/x$ sein. Daraus folgt $f(x) = \ln x$ und

$$\int \frac{x}{x^2+1} dx = \frac{1}{2}f(g(x)) = \frac{1}{2}\ln(x^2+1).$$

BEISPIEL 104. Wir berechnen das bestimmte Integral

$$\int_{\pi/3}^{\pi/2} \frac{\cos x}{\sin x} dx.$$

Mit $g'(x) = \cos x$, $g(x) = \sin x$, $f'(x) = 1/x$ und $f(x) = \ln x$, da $\sin x > 0$ für alle $\pi/3 \leq x \leq \pi/2$, erhält man die Stammfunktion $\ln(\sin x)$ und

$$\begin{aligned} \int_{\pi/3}^{\pi/2} \frac{\cos x}{\sin x} dx &= [\ln(\sin x)]_{\pi/3}^{\pi/2} = \ln\left(\sin\left(\frac{\pi}{2}\right)\right) - \ln\left(\sin\left(\frac{\pi}{3}\right)\right) = \ln 1 - \ln \frac{\sqrt{3}}{2} \\ &= -\ln \frac{\sqrt{3}}{2} = \ln \frac{2}{\sqrt{3}}. \end{aligned}$$

SATZ 74 (Partielle Integration). Sind f, g stetig differenzierbar, so gilt

$$(81) \quad \int f'(x)g(x) dx = f(x)g(x) - \int f(x)g'(x) dx.$$

BEWEIS. Aus der Produktregel $(fg)' = f'g + g'f$ folgt $f'g = (fg)' - g'f$ und durch Integration der Gleichung

$$\int f'(x)g(x) dx = \int (fg)'(x) - g'(x)f(x) dx = (fg)(x) - \int g'(x)f(x) dx.$$

□

BEISPIEL 105. Mit $f'(x) = e^x$ und $g(x)$ erhält man $f(x) = e^x$, $g'(x) = 1$ und

$$\int x e^x dx = e^x x - \int e^x \cdot 1 dx = e^x x - \int e^x dx = e^x x - e^x = e^x(x - 1).$$

BEISPIEL 106 (Rekursionsformel für $\int x^n e^x dx$). Für $f(x) = f'(x) = e^x$, $g(x) = x^n$ und $g'(x) = n x^{n-1}$ erhält man mit partieller Integration

$$(82) \quad \int x^n e^x dx = x^n e^x - \int n x^{n-1} e^x dx = x^n e^x - n \int x^{n-1} e^x dx \quad \forall n \in \mathbb{N}.$$

Für $n = 2$ ergibt sich

$$(83) \quad \int x^2 e^x dx = x^2 e^x - 2 \int x e^x dx = x^2 e^x - 2 e^x(x - 1) = e^x(x^2 - 2x + 2).$$

BEISPIEL 107 (Stammfunktion von $\ln x$). Mit $f'(x) = 1$, $f(x) = x$, $g(x) = \ln x$ und $g'(x) = 1/x$ erhält man

$$\begin{aligned} \int \ln x dx &= \int 1 \cdot \ln x dx = x \ln x - \int x \frac{1}{x} dx = x \ln x - \int 1 dx = x \ln x - x \\ &= x(\ln x - 1). \end{aligned}$$

BEISPIEL 108 (Stammfunktion von $x^n \ln x$). Für $f'(x) = x^n$, $f(x) = \frac{1}{n+1} x^{n+1}$, $g(x) = \ln x$ und $g'(x) = 1/x$ erhält man mit partieller Integration für alle $n \in \mathbb{N}$ die explizite Formel

$$\int x^n \ln x dx = \frac{x^{n+1}}{n+1} \ln x - \int \frac{x^{n+1}}{n+1} \frac{1}{x} dx = \frac{x^{n+1}}{n+1} \ln x - \int \frac{x^n}{n+1} dx,$$

also

$$(84) \quad \int x^n \ln x dx = \frac{x^{n+1}}{n+1} \ln x - \frac{x^{n+1}}{(n+1)^2}.$$

BEISPIEL 109 (Stammfunktion von $\sin^2 x$). Für $f'(x) = \sin x$, $f(x) = -\cos x$, $g(x) = \sin x$ und $g'(x) = \cos x$ erhält man mit partieller Integration und der Identität $1 = \sin^2 x + \cos^2 x$

$$\begin{aligned} \int \sin^2 x dx &= \int (\sin x)(\sin x) dx = -\cos x \sin x - \int -\cos x \cos x dx \\ &= -\cos x \sin x + \int \cos^2 x dx = -\cos x \sin x + \int (1 - \sin^2 x) dx \\ \int \sin^2 x dx &= -\cos x \sin x + \int 1 dx - \int \sin^2 x dx \\ 2 \int \sin^2 x dx &= -\cos x \sin x + x \\ \int \sin^2 x dx &= -\frac{1}{2} \cos x \sin x + \frac{x}{2}. \end{aligned}$$

BEISPIEL 110 (Stammfunktion von $\cos^2 x$). Die Stammfunktion von $\cos^2 x$ kann man ähnlich wie in Beispiel 109 mit partieller Integration bestimmen. Man kann aber auch die Identität $\sin^2 x + \cos^2 x = 1$ und die Stammfunktion von $\sin^2 x$ benutzen:

$$\int \cos^2 x \, dx = \int (1 - \sin^2 x) \, dx = x - \left(-\frac{1}{2} \cos x \sin x + \frac{x}{2} \right) = \frac{x}{2} + \frac{1}{2} \cos x \sin x.$$

BEISPIEL 111 (Rekursionformel für $\int \sin^n x \, dx$). Für $f'(x) = \sin x$, $f(x) = -\cos x$, $g(x) = \sin^{n-1} x$ und $g'(x) = (n-1) \sin^{n-2} x \cos x$ erhält man mit partieller Integration und der Identität $1 = \sin^2 x + \cos^2 x$

$$\begin{aligned} \int \sin^n x \, dx &= \int \sin x (\sin x)^{n-1} \, dx \\ &= -\cos x (\sin x)^{n-1} - \int -\cos x (n-1) (\sin x)^{n-2} \cos x \, dx \\ &= -\cos x (\sin x)^{n-1} + (n-1) \int \cos^2 x (\sin x)^{n-2} \, dx \\ &= -\cos x (\sin x)^{n-1} + (n-1) \int (1 - \sin^2 x) (\sin x)^{n-2} \, dx \\ n \int \sin^n x \, dx &= -\cos x (\sin x)^{n-1} + (n-1) \int (\sin x)^{n-2} \, dx, \end{aligned}$$

also

$$(85) \quad \int (\sin x)^n \, dx = -\frac{1}{n} \cos x (\sin x)^{n-1} + \frac{n-1}{n} \int (\sin x)^{n-2} \, dx.$$

Für $n = 2$ erhält man die Formel aus Beispiel 109. Bei der Anwendung dieser Rekursionsformel wird der Parameter n in jedem Schritt um 2 kleiner. Man kann rekursiv alle Stammfunktionen $\int \sin^n x \, dx$ berechnen, da man die Integrale für $n = 0$ und $n = 1$, also $\int 1 \, dx$ und $\int \sin x \, dx$, kennt.

BEISPIEL 112 (Rekursionformel für $\int \cos^n x \, dx$). Eine Rekursionsformel für das Integral $\int \cos^n x \, dx$ findet man ähnlich wie in Beispiel 111. Für $f'(x) = \cos x$, $f(x) = \sin x$, $g(x) = \cos^{n-1} x$ und $g'(x) = (n-1) \cos^{n-2} x (-\sin x)$ erhält man mit partieller Integration und der Identität $1 = \sin^2 x + \cos^2 x$

$$\begin{aligned} \int \cos^n x \, dx &= \int \cos x (\cos x)^{n-1} \, dx \\ &= \sin x (\cos x)^{n-1} - \int \sin x (n-1) (\cos x)^{n-2} (-\sin x) \, dx \\ &= \sin x (\cos x)^{n-1} + (n-1) \int \sin^2 x (\cos x)^{n-2} \, dx \\ &= \sin x (\cos x)^{n-1} + (n-1) \int (1 - \cos^2 x) (\cos x)^{n-2} \, dx \\ n \int \cos^n x \, dx &= \sin x (\cos x)^{n-1} + (n-1) \int (\cos x)^{n-2} \, dx, \end{aligned}$$

also

$$(86) \quad \int (\cos x)^n dx = \frac{1}{n} \sin x (\cos x)^{n-1} + \frac{n-1}{n} \int (\cos x)^{n-2} dx.$$

Für $n = 2$ erhält man die Formel aus Beispiel 110.

BEISPIEL 113 (Rekursionsformel für $\int (1+x^2)^{-n} dx$). Für alle $n \in \mathbb{N}$ mit $n > 0$ gilt wegen $((x^2+1)^{-n})' = -n(x^2+1)^{-n-1}2x$

$$\begin{aligned} \int \frac{1}{(1+x^2)^n} dx &= \int 1 \cdot \frac{1}{(1+x^2)^n} dx = x \frac{1}{(1+x^2)^n} - \int x \frac{-n2x}{(x^2+1)^{n+1}} dx \\ &= \frac{x}{(1+x^2)^n} + 2n \int \frac{x^2}{(x^2+1)^{n+1}} dx \\ &= \frac{x}{(1+x^2)^n} + 2n \int \frac{x^2+1-1}{(x^2+1)^{n+1}} dx \\ \int \frac{1}{(1+x^2)^n} dx &= \frac{x}{(1+x^2)^n} + 2n \int \frac{1}{(x^2+1)^n} dx - 2n \int \frac{1}{(x^2+1)^{n+1}} dx \\ 2n \int \frac{1}{(x^2+1)^{n+1}} dx &= \frac{x}{(1+x^2)^n} + (2n-1) \int \frac{1}{(x^2+1)^n} dx. \end{aligned}$$

Daraus folgt die Rekursionsformel

$$(87) \quad \int \frac{1}{(x^2+1)^{n+1}} dx = \frac{1}{2n} \cdot \frac{x}{(1+x^2)^n} + \frac{2n-1}{2n} \int \frac{1}{(x^2+1)^n} dx \quad \forall n \in \mathbb{N}, n > 0.$$

Bei der Anwendung dieser Rekursionsformel verringert sich der Parameter n in jedem Schritt um 1. Man kann rekursiv alle Integrale $\int (1+x^2)^{-n} dx$ für $n \in \mathbb{N}$ berechnen, da man $\int (1+x^2)^{-1} dx = \arctan x$ kennt.

7. Substitutionsregel

SATZ 75. Wenn f stetig auf dem Intervall $[a, b]$ ist, g stetig differenzierbar und streng monoton auf dem Intervall $[c, d]$ ist und $g([c, d]) = [a, b]$ gilt, dann

$$(88) \quad \int_a^b f(x) dx = \int_{g^{-1}(a)}^{g^{-1}(b)} f(g(t)) g'(t) dt.$$

BEWEIS. Da g streng monoton ist, besitzt g eine Umkehrfunktion g^{-1} . Ist g streng monoton wachsend, so gilt $g(c) = a$ und $g(d) = b$. Ist g streng monoton fallend, so gilt $g(c) = b$ und $g(d) = a$.

Da f stetig ist, existiert eine differenzierbare Funktion F mit $F'(x) = f(x)$ und $\int_a^b f(x) dx = F(b) - F(a)$. Nun gilt $(F \circ g)'(t) = F'(g(t)) g'(t) =$

$f(g(t))g'(t)$. Daraus folgt

$$\begin{aligned} \int_{g^{-1}(a)}^{g^{-1}(b)} f(g(t))g'(t) dt &= [(F \circ g)(t)]_{g^{-1}(a)}^{g^{-1}(b)} = F(g(g^{-1}(b))) - F(g(g^{-1}(a))) \\ &= F(b) - F(a) = \int_a^b f(x) dx. \end{aligned}$$

□

BEISPIEL 114. Die Funktion $g(t) = \sin t$ ist auf dem Intervall $[-\pi/2, \pi/2]$ streng monoton wachsend und stetig differenzierbar. Da $g'(t) = \cos t$, $g(0) = 0$ und $g(\pi/6) = 1/2$, gilt

$$\begin{aligned} \int_0^{1/2} \frac{1}{\sqrt{1-x^2}} dx &= \int_0^{\pi/6} \frac{1}{\sqrt{1-\sin^2 t}} \cos t dt = \int_0^{\pi/6} \frac{1}{\sqrt{\cos^2 t}} \cos t dt \\ &= \int_0^{\pi/6} \frac{1}{\cos t} \cos t dt = \int_0^{\pi/6} 1 dt = \frac{\pi}{6}. \end{aligned}$$

BEISPIEL 115 (Flächeninhalt eines Halbkreises). Es sei $r > 0$ fest. Die Funktion $f: [-r, r] \rightarrow \mathbb{R}$ mit $f(x) = \sqrt{r^2 - x^2}$ ist stetig. Ihr Graph und die x -Achse begrenzen einen Halbkreis mit Radius r , denn $f(x)^2 + x^2 = r^2$. Wir berechnen den Flächeninhalt dieses Halbkreises mit Hilfe der Substitution $x = g(t) = r \sin t$ mit $t \in [-\pi/2, \pi/2]$, also $g(-\pi/2) = -r$, $g(\pi/2) = r$ und $g'(t) = r \cos t$, und der Stammfunktion von $\cos^2 t$ (siehe Beispiel 110):

$$\begin{aligned} \int_{-r}^r f(x) dx &= \int_{-r}^r \sqrt{r^2 - x^2} dx = \int_{-\pi/2}^{\pi/2} \sqrt{r^2 - (r \sin t)^2} (r \cos t) dt \\ &= \int_{-\pi/2}^{\pi/2} \sqrt{r^2(1 - \sin^2 t)} r \cos t dt = \int_{-\pi/2}^{\pi/2} r \sqrt{\cos^2 t} r \cos t dt \\ &= r^2 \int_{-\pi/2}^{\pi/2} \cos^2 t dt = r^2 \frac{1}{2} [t + \sin t \cos t]_{-\pi/2}^{\pi/2} = \frac{1}{2} r^2 \pi. \end{aligned}$$

BEISPIEL 116. Für beliebige $a, b \in \mathbb{R}$ mit $b \neq 0$ gilt

$$\int \frac{1}{(x+a)^2 + b^2} dx = \frac{1}{b^2} \int \frac{1}{\left(\frac{x+a}{b}\right)^2 + 1} dx.$$

Mit der Substitution $t = g^{-1}(x) = (x+a)/b$, also $g(t) = bt - a$ und $g'(t) = b$, folgt

$$\begin{aligned} \int \frac{1}{(x+a)^2 + b^2} dx &= \frac{1}{b^2} \int \frac{1}{t^2 + 1} b dt = \frac{1}{b} \int \frac{1}{t^2 + 1} dt = \frac{1}{b} \arctan t \\ &= \frac{1}{b} \arctan \frac{x+a}{b}. \end{aligned}$$

8. Partialbruchzerlegung

Für beliebige reelle Polynome p, q mit $q \not\equiv 0$ ist die Funktion $f := p/q$ für alle x mit $q(x) \neq 0$ definiert und stetig. Die Funktion f besitzt auf jedem Intervall, das keine Nullstelle von q enthält, eine Stammfunktion. In einigen Fällen ist diese Stammfunktion leicht zu bestimmen, z.B.

$$\int \frac{1}{(x + \alpha)^n} dx = \begin{cases} \ln|x + \alpha| & n = 1 \\ -\frac{1}{n-1}(x + \alpha)^{-(n-1)} & n > 1 \end{cases} \quad \text{für alle } x \neq \alpha,$$

mit Hilfe der Substitution $y = (x + \alpha)/\beta$, z.B.

$$\int \frac{1}{((x + \alpha)^2 + \beta^2)^n} dx = \frac{1}{b^{2n-1}} \int \frac{1}{(y^2 + 1)^n} dy,$$

vereinfachen und dann direkt mit Hilfe der Kettenregel

$$\int \frac{x}{(x^2 + 1)^n} dx = \begin{cases} \frac{1}{2} \ln|x^2 + 1| & n = 1 \\ -\frac{1}{2(n-1)}(x^2 + 1)^{n-1} & n > 1 \end{cases}$$

oder mit der Rekursionsformel 87 berechnen.

Kennt man alle (reellen und komplexen) Nullstellen des reellen Polynoms q , also die Zerlegung des Polynoms q in lineare und quadratische Faktoren, so kann man die Stammfunktion der Funktion p/q mit Hilfe der **Partialbruchzerlegung** bestimmen.

SATZ 76 (Partialbruchzerlegung). *Es seien p, q reelle Polynome mit $q \not\equiv 0$. Wenn*

$$q(x) = c \prod_{j=1}^R (x - \alpha_j)^{m_j} \prod_{j=1}^K q_j(x)^{n_j}$$

für paarweise verschiedene Nullstellen $\alpha_j \in \mathbb{R}$ von q , paarweise verschiedene quadratische Polynome q_j ohne reelle Nullstelle, $n_j, m_j \in \mathbb{N}$ und $c \in \mathbb{R}$, dann existieren ein Polynom h und $A_{jl}, B_{jl}, C_{jl} \in \mathbb{R}$ mit

$$(89) \quad \frac{p(x)}{q(x)} = h(x) + \sum_{j=1}^R \frac{A_{j1}}{x - \alpha_j} + \frac{A_{j2}}{(x - \alpha_j)^2} + \dots + \frac{A_{jm_j}}{(x - \alpha_j)^{m_j}} \\ + \sum_{j=1}^K \frac{B_{j1}x + C_{j1}}{q_j(x)} + \frac{B_{j2}x + C_{j2}}{q_j(x)^2} + \dots + \frac{B_{jn_j}x + C_{jn_j}}{q_j(x)^{n_j}}.$$

Im einfachsten Fall, wenn das Polynom q nur einfache reelle Nullstellen besitzt, also $K = 0$ und $m_j = 1$ für alle j , kann man die Koeffizienten $A_1, \dots, A_R$ leicht bestimmen. Der Ansatz

$$\frac{p(x)}{q(x)} = \frac{A_1}{x - \alpha_1} + \dots + \frac{A_R}{x - \alpha_R}$$

führt durch Multiplikation mit dem Nenner $q(x) = c \prod_{j=1}^R (x - \alpha_j)$ zu

$$p(x) = c \sum_{j=1}^R A_j \prod_{k \neq j} (x - \alpha_k).$$

Wertet man diese Gleichung für $x = \alpha_j$ aus, so erhält man für alle $j = 1, \dots, R$

$$(90) \quad A_j = \frac{p(\alpha_j)}{c \prod_{k \neq j} (\alpha_j - \alpha_k)}.$$

BEISPIEL 117. Wir suchen die Partialbruchzerlegung von

$$\frac{x}{(x+1)(x-1)(x-2)},$$

also $p(x) = x$ und $q(x) = (x+1)(x-1)(x-2)$ hat drei einfache reelle Nullstellen. Der Ansatz

$$\frac{x}{(x+1)(x-1)(x-2)} = \frac{A_1}{x+1} + \frac{A_2}{x-1} + \frac{A_3}{x-2}$$

mit Koeffizienten $A_j \in \mathbb{R}$ führt zu

$$x = A_1(x-1)(x-2) + A_2(x+1)(x-2) + A_3(x+1)(x-1)$$

$$-1 = A_1 \cdot (-2) \cdot (-3) \quad \Rightarrow A_1 = -\frac{1}{6}$$

$$1 = A_2 \cdot 2 \cdot (-1) \quad \Rightarrow A_2 = -\frac{1}{2}$$

$$2 = A_3 \cdot 3 \cdot 1 \quad \Rightarrow A_3 = \frac{2}{3},$$

also

$$\frac{x}{(x+1)(x-1)(x-2)} = -\frac{1}{6} \cdot \frac{1}{x+1} - \frac{1}{2} \cdot \frac{1}{x-1} + \frac{2}{3} \cdot \frac{1}{x-2}.$$

Treten mehrfache reelle Nullstellen auf, so könnte man Residuen benutzen, um die Koeffizienten A_{jl} zu berechnen. In jedem Fall kann man die Gleichung 89 mit $q(x)$ multiplizieren und das durch Koeffizientenvergleich der Polynome entstehende lineare Gleichungssystem lösen.

BEISPIEL 118. Für $p(x) = x^3$ und $q(x) = (x^2 + 1)^2$ gilt $\deg p < \deg q$. Der Ansatz

$$\frac{x^3}{(x^2 + 1)^2} = \frac{B_1 x + C_1}{x^2 + 1} + \frac{B_2 x + C_2}{(x^2 + 1)^2}$$

mit $B_1, B_2, C_1, C_2 \in \mathbb{R}$ führt durch Multiplikation mit $q(x)$ zu

$$x^3 = (B_1 x + C_1)(x^2 + 1) + B_2 x + C_2 = B_1 x^3 + C_1 x^2 + (B_1 + B_2)x + (C_1 + C_2),$$

also $B_1 = 1$, $C_1 = 0$, $B_2 = -B_1 = -1$ und $C_1 = -C_2 = 0$. Damit

$$\int \frac{x^3}{(x^2 + 1)^2} dx = \int \left(\frac{x}{x^2 + 1} - \frac{x}{(x^2 + 1)^2} \right) dx = \frac{1}{2} \ln(x^2 + 1) + \frac{1}{2} (x^2 + 1)^{-1}.$$

8.1. Größter gemeinsamer Teiler für natürliche Zahlen. Der **größte gemeinsame Teiler** zweier natürlicher Zahlen $a, b \in \mathbb{N}$ mit $a \neq 0$ ist die größte natürliche Zahl, die a und b teilt. Sie wird mit $\text{ggT}(a, b)$ bezeichnet.

$$\text{ggT}(a, b) := \max\{n \in \mathbb{N} : n|a \text{ und } n|b\}.$$

Für alle $n \in \mathbb{N}$ gilt $\text{ggT}(n, 0) = n$ und $\text{ggT}(n, 1) = 1$. Wenn $a, b \in \mathbb{N}$ mit $b \geq a$, dann existieren eindeutige $m, r \in \mathbb{N}$ mit $b = ma + r$ und $0 \leq r < a$ (Division mit Rest).

LEMMA 22. Wenn $b = ma + r$ für $a, b, m, r \in \mathbb{N}$, dann $\text{ggT}(a, b) = \text{ggT}(r, a)$.

Der **Euklidische Algorithmus** führt die Berechnung des $\text{ggT}(a, b)$ schrittweise durch wiederholte Division mit Rest auf $\text{ggT}(a_n, 0) = a_n$ zurück. Für $a, b \in \mathbb{N}$ mit $b \geq a$ definiert man $a_0 := a$ und rekursiv $a_{k+1}, m_k \in \mathbb{N}$ durch

$$\begin{aligned} b &= m_0 a_0 + a_1 & \text{mit } 0 \leq a_1 < a_0 \\ a_{k-1} &= m_k a_k + a_{k+1} & \text{mit } 0 \leq a_{k+1} < a_k \end{aligned}$$

bis $a_{n+1} = 0$. Dann gilt $\text{ggT}(a, b) = \text{ggT}(a_n, a_{n+1}) = a_n$.

BEISPIEL 119. Für $b = 71$ und $a = 32$ erhält man $a_0 = 32$ und

$$\begin{aligned} 71 &= 2 \cdot 32 + 7, & a_1 &= 7, & m_0 &= 2 \\ 32 &= 4 \cdot 7 + 4, & a_2 &= 4, & m_1 &= 4 \\ 7 &= 1 \cdot 4 + 3, & a_3 &= 3, & m_2 &= 1 \\ 4 &= 1 \cdot 3 + 1, & a_4 &= 1, & m_3 &= 1 \\ 3 &= 3 \cdot 1 + 0, & a_5 &= 0, & m_4 &= 3, \end{aligned}$$

also $\text{ggT}(71, 32) = a_4 = 1$.

FOLGERUNG 28. Es seien $a, b \in \mathbb{N}$. Es existieren $\alpha, \beta \in \mathbb{Z}$ mit $\text{ggT}(a, b) = \alpha a + \beta b$.

BEISPIEL 120. Aus der Durchführung des Euklidischen Algorithmus in Beispiel 119 folgt für $a = 32$ und $b = 71$

$$\begin{aligned} \text{ggT}(32, 71) &= a_4 = a_2 - m_3 a_3 = 4 - 1 \cdot 3 \\ &= a_2 - m_3 \cdot (a_1 - m_2 a_2) = 2 \cdot a_2 - 1 \cdot a_1 = 2 \cdot 4 - 1 \cdot 7 \\ &= 2 \cdot (a_0 - m_1 a_1) - 1 \cdot a_1 = 2 \cdot a_0 - 9 \cdot a_1 = 2 \cdot 32 - 9 \cdot 7 \\ &= 2 \cdot a_0 - 9 \cdot (b - m_0 a_0) = 20 \cdot a_0 - 9 \cdot b \end{aligned}$$

$$\text{ggT}(32, 71) = 1 = 20 \cdot 32 - 9 \cdot 71,$$

also $\alpha = 20$ und $\beta = -9$.

BEMERKUNG 59. Kennt man die Primzahlzerlegung zweier natürlicher Zahlen, so kann man den größten gemeinsamen Teiler sofort ablesen, er ist das Produkt aller Faktoren, die in beiden Zerlegungen vorkommen. Aber die Bestimmung der Primzahlzerlegung ist viel aufwendiger als Durchführung des Euklidischen Algorithmus.

8.2. Größter gemeinsamer Teiler für reelle Polynome. Ein reelles Polynom h heißt Teiler des reellen Polynoms p , falls ein reelles Polynom q mit $hq = p$ existiert. Man schreibt dann $h|p$. Wenn $h(x)$ Teiler des Polynoms $p(x)$, so ist für alle $c \in \mathbb{R}$ mit $c \neq 0$ auch $ch(x)$ Teiler des Polynoms $p(x)$.

Für reelle Polynome p, q mit $q \neq 0$ ist ein **größter gemeinsamer Teiler** ein Polynom mit maximalem Grad in der Menge der reellen Polynome, die q und p teilen. Man schreibt $\text{ggT}(p, q)$. Es gilt $\text{ggT}(q, 0) = q$ für alle Polynome $q \neq 0$. Ist h ein größter gemeinsamer Teiler der Polynome p und q , so ist für alle $c \in \mathbb{R}$ mit $c \neq 0$ auch $ch(x)$ ein größter gemeinsamer Teiler.

Der **Euklidische Algorithmus** führt die Berechnung eines $\text{ggT}(q, p)$ schrittweise durch wiederholte Polynomdivision mit Rest auf $\text{ggT}(q_n, 0) = q_n$ zurück. Für reelle Polynome p, q mit $\deg p \geq \deg q$ und $q \neq 0$ definiert man $q_0 := a$ und rekursiv Polynome q_{k+1}, m_k durch

$$\begin{aligned} p(x) &= m_0(x)q_0(x) + q_1(x) \quad \text{mit } 0 \leq \deg q_1(x) < \deg q_0(x) \\ q_{k-1}(x) &= m_k(x)q_k(x) + q_{k+1}(x) \quad \text{mit } 0 \leq \deg q_{k+1}(x) < \deg q_k(x) \end{aligned}$$

bis $q_{n+1} \equiv 0$. Dann gilt $\text{ggT}(p, q) = \text{ggT}(q_n, q_{n+1}) = q_n$.

BEISPIEL 121. Für $p(x) = x^2 + 1$ und $q(x) = x^2 + 2x + 1$ erhält man $q_0(x) = x^2 + 2x + 1$ und

$$\begin{aligned} p(x) &= x^2 + 1 = 1 \cdot (x^2 + 2x + 1) - 2x, & q_1(x) &= -2x, & m_0(x) &\equiv 1 \\ q_0(x) &= x^2 + 2x + 1 = \left(-\frac{x}{2} - 1\right)(-2x) + 1, & q_2(x) &\equiv 1, & m_1(x) &= -\frac{x}{2} - 1 \\ q_1(x) &= -2x = (-2x) \cdot 1 + 0, & q_3(x) &\equiv 0, & m_2(x) &= -2x \end{aligned}$$

also $\text{ggT}(x^2+1, x^2+2x+1) = q_2(x) \equiv 1$. Es gilt $p(x) = x^2+1 = (x-i)(x+i)$ und $q(x) = x^2+2x+1 = (x+1)^2$. Die Polynome p und q haben keine gemeinsame Nullstelle.

FOLGERUNG 29. *Es seien $p, q \neq 0$ reelle Polynome. Es existieren reelle Polynome α, β mit $\text{ggT}(p, q) = \alpha q + \beta p$.*

BEISPIEL 122. Aus den Gleichungen in Beispiel 121 erhält man für $p(x) = x^2 + 1$ und $q(x) = x^2 + 2x + 1$ mit $\text{ggT}(p, q) = 1$

$$\begin{aligned} \text{ggT}(p, q) &= 1 = q_0(x) - m_1(x)q_1(x) = x^2 + 2x + 1 + \left(\frac{x}{2} + 1\right)(-2x) \\ &= q_0(x) - m_1(x)(p(x) - m_0(x)q_0(x)) = -\frac{x}{2}q_0(x) + \left(\frac{x}{2} + 1\right)p(x) \\ 1 &= -\frac{x}{2}q(x) + \left(\frac{x}{2} + 1\right)p(x), \end{aligned}$$

also $\alpha(x) = -x/2$ und $\beta(x) = 1 + x/2$.

LEMMA 23. *Der größte gemeinsame Teiler der reellen Polynome $p, q \neq 0$ ist genau dann eine Konstante, wenn p und q keine gemeinsamen Nullstellen besitzen.*

BEMERKUNG 60. Kennt man die Zerlegung der Polynome p und q in lineare und quadratische Faktoren, also alle reellen und komplexen Nullstellen, so kann man den größten gemeinsamen Teiler der Polynome ablesen. Die exakte Bestimmung der Nullstellen der Polynome ist vielleicht gar nicht möglich. Der Euklidische Algorithmus aber ist immer durchführbar.

8.3. Beweis der Existenz der Partialbruchzerlegung (Gleichung 89). Wenn p, q reelle Polynome sind und $q(x) = q_1(x)q_2(x)$ für zwei reelle Polynome q_1, q_2 ohne gemeinsame Nullstelle gilt, dann existieren wegen $\text{ggT}(q_1, q_2) = 1$ reelle Polynome α_1, β_1 mit $1 = \alpha_1(x)q_1(x) + \beta_1(x)q_2(x)$. Daraus folgt

$$\frac{1}{q_1(x)q_2(x)} = \frac{\alpha_1(x)}{q_2(x)} + \frac{\beta_1(x)}{q_1(x)}, \text{ also } \frac{p(x)}{q_1(x)q_2(x)} = \frac{\alpha_1(x)p(x)}{q_2(x)} + \frac{\beta_1(x)p(x)}{q_1(x)}.$$

Wendet man diese Methode auf alle Faktoren $(x - \alpha_j)^{m_j}$ und $q_j(x)^{n_j}$ der Zerlegung des Polynoms $q(x)$ in lineare und quadratische Faktoren an, so erhält man Polynome p_{Rj} und p_{Kj} mit

$$\frac{p(x)}{q(x)} = \sum_{j=1}^R \frac{p_{Rj}(x)}{(x - \alpha_j)^{m_j}} + \sum_{j=1}^K \frac{p_{Kj}(x)}{q_j(x)^{n_j}}.$$

Jeden einzelnen Summanden kann man jetzt mit Hilfe der Polynomdivision weiter vereinfachen, denn $p_{Rj}(x) = m_j(x)(x - \alpha_j) + a_j$ für ein Polynom m_j und eine reelle Zahl $a_j \in \mathbb{R}$ und $p_{Kj}(x) = n_j(x)q_j(x) + b_jx + c_j$ für ein Polynom n_j und eine reelle Zahlen $b_j, c_j \in \mathbb{R}$.

9. Die Gammafunktion

Wir haben bereits gesehen, dass man zu einer stetigen Funktion f durch Bildung des bestimmten Integrals $\int_{x_0}^x f(t) dt$ mit variabler oberer Intervallgrenze eine Stammfunktion $F(x)$ definieren kann.

Hängt eine Funktion $f(t, x)$ von zwei reellen Variablen ab und ist für jedes $x_0 \in I$ die Funktion $f(t, x_0)$ auf $[a, b]$ stetig, wird durch $F(x_0) := \int_a^b f(t, x_0) dt$ für jedes $x_0 \in I$ eine Funktion $F : I \rightarrow \mathbb{R}$ definiert.

Dies kann auch für unbeschränkte Integrationsintervalle, also $a = -\infty$ oder $b = \infty$, oder zwar auf (a, b) stetige aber unbeschränkte Funktionen $f(t, x_0)$ funktionieren, falls die auftretenden uneigentlichen Integrale existieren. Ein Beispiel dafür ist die **Gamma-Funktion**

$$(91) \quad \Gamma : \mathbb{R}^{>0} \rightarrow \mathbb{R}, \quad \Gamma(x) := \int_0^\infty e^{-t} t^{x-1} dt.$$

Wir müssen nachweisen, dass die in Gleichung 91 auftretenden uneigentlichen Integrale existieren, d.h. die Gammafunktion für alle $x > 0$ definiert ist. Einige der uneigentlichen Integrale kann man exakt bestimmen.

BEISPIEL 123 ($\Gamma(1) = 1$). Für $x = 1$ erhält man

$$\begin{aligned}\Gamma(1) &= \int_0^\infty e^{-t} t^{1-1} dt = \int_0^\infty e^{-t} dt = \lim_{b \rightarrow \infty} \int_0^b e^{-t} dt = \lim_{b \rightarrow \infty} [-e^{-t}]_0^b \\ &= \lim_{b \rightarrow \infty} (-e^{-b} + e^{-0}) = 1 - \lim_{b \rightarrow \infty} \frac{1}{e^b} = 1,\end{aligned}$$

da $\lim_{b \rightarrow \infty} e^b = \infty$.

BEISPIEL 124 ($\Gamma(2) = 1$). Für $x = 2$ erhält man mit Hilfe der partiellen Integration

$$\begin{aligned}\Gamma(2) &= \int_0^\infty e^{-t} t^{2-1} dt = \int_0^\infty e^{-t} t dt = \lim_{b \rightarrow \infty} \int_0^b e^{-t} t dt \\ \int e^{-t} t dt &= -e^{-t} t - \int (-e^{-t}) dt = -e^{-t} t - e^{-t} = -e^{-t}(t+1) \\ \Gamma(2) &= \lim_{b \rightarrow \infty} [-e^{-t}(t+1)]_0^b = -\lim_{b \rightarrow \infty} \frac{b+1}{e^b} + \frac{1+0}{e^0} = 1,\end{aligned}$$

da die Exponentialfunktion stärker wächst als jedes Polynom.

BEISPIEL 125 ($\Gamma(3) = 2$). Für $x = 3$ erhält man mit Hilfe der partiellen Integration und Beispiel 124

$$\begin{aligned}\Gamma(3) &= \int_0^\infty e^{-t} t^{3-1} dt = \int_0^\infty e^{-t} t^2 dt = \lim_{b \rightarrow \infty} \int_0^b e^{-t} t^2 dt \\ \int e^{-t} t^2 dt &= -e^{-t} t^2 - \int (-e^{-t} 2t) dt = -e^{-t} t^2 + 2 \int e^{-t} t dt \\ &= -e^{-t}(t^2 + 2t + 2) \\ \Gamma(3) &= \lim_{b \rightarrow \infty} [-e^{-t}(t^2 + 2t + 2)]_0^b \\ &= -\lim_{b \rightarrow \infty} \frac{b^2 + 2b + 2 + 1}{e^b} + \frac{0 + 0 + 2}{e^0} = 2,\end{aligned}$$

da die Exponentialfunktion stärker wächst als jedes Polynom.

9.1. Existenz der Gamma-Funktion für $x \geq 1$. Für $x \geq 1$ gilt $x-1 \geq 0$ und die Funktion $f(t) := e^{-t} t^{x-1}$ ist als Funktion von t stetig auf $\mathbb{R}$ und damit auf jedem Intervall $[0, b]$ integrierbar. Da die Exponentialfunktion schneller wächst als jedes Polynom, existiert eine Konstante $t_0 \in \mathbb{R}$ mit $t^{x+1} \leq e^t$ für alle $t \geq t_0$. Da $f(t) \geq 0$ für alle $t \geq 0$, gilt

$$F(b) := \int_0^b f(t) dt = \int_0^a f(t) dt + \int_a^b f(t) dt \leq \int_0^a f(t) dt + \int_a^b \frac{1}{t^2} dt$$

für alle $a \geq t_0$. Wir können die zu integrierende Funktion $f(t)$ gegen die größere Funktion t^{-2} abschätzen. Da das uneigentliche Integral

$$\int_a^\infty \frac{1}{t^2} dt = \lim_{b \rightarrow \infty} \left[-\frac{1}{t} \right]_a^b = -\lim_{b \rightarrow \infty} \frac{1}{b} + \frac{1}{a} = \frac{1}{a}$$

für alle $a > 0$ existiert und, wegen der expliziten Stammfunktion, einfach zu berechnen ist, existiert auch das uneigentliche Integral $\int_0^\infty f(t)dt$, also $\Gamma(x)$ für $x \geq 1$ (siehe Satz 77).

SATZ 77 (Vergleichskriterium für uneigentliche Integrale). *Wenn $f : [a, \infty)$ stetig, $|f(t)| \leq g(t)$ für alle $t \geq a$ und das uneigentliche Integral $\int_a^\infty g(t) dt$ existiert, dann existiert auch das uneigentliche Integral $\int_a^\infty f(t) dt$.*

BEMERKUNG 61. Satz 77 sichert nur die Existenz des uneigentlichen Integrals. Er liefert nicht den Grenzwert $\lim_{b \rightarrow \infty} \int_a^b f(t) dt$. $\square$

9.2. Existenz der Gamma-Funktion für $0 < x < 1$. Für $0 < x < 1$ gilt $x - 1 < 0$ und die Funktion $f(t) := e^{-t}t^{x-1}$ ist für $t = 0$ nicht definiert, da $\lim_{t \rightarrow 0} |f(t)| = \infty$. Das Integral $\int_0^\infty f(t)dt$ ist also „doppelt“ uneigentlich und muss zerlegt werden. Da f auf dem offenen Intervall $(0, \infty)$ stetig ist, existiert $\int_0^\infty f(t) dt$ genau dann, wenn für ein $a > 0$, z.B. $a = 1$, beide uneigentlichen Integrale $\int_0^a f(t) dt$ und $\int_a^\infty f(t) dt$ existieren.

Da $x - 1 < 0$, gilt $t^{x-1} \leq 1$ für alle $t \geq 1$. Daraus folgt $|f(t)| \leq e^{-t}$ für alle $t \geq 1$ und das uneigentliche Integral $\int_1^\infty f(t) dt$ existiert, weil $\int_1^\infty e^{-t} dt = e^{-1}$ existiert (siehe Beispiel 123 und Satz 77).

LEMMA 24. *Das Integral $\int_1^\infty x^\alpha dx$ existiert genau dann, wenn $\alpha < -1$.*

BEWEIS. Für $a \neq -1$ gilt

$$\int_1^\infty x^\alpha dx = \frac{1}{\alpha + 1} \left(\lim_{b \rightarrow \infty} b^{\alpha+1} - 1 \right).$$

Der Grenzwert $\lim_{b \rightarrow \infty} b^{\alpha+1}$ existiert genau dann, wenn $\alpha + 1 < 0$. Für $a = -1$ gilt $\int_1^\infty x^{-1} dx = \lim_{b \rightarrow \infty} \ln b = \infty$. $\square$

Durch die Substitution $t = s^{-1}$, also $dt = -s^{-2}ds$, ergibt sich

$$\int_a^1 e^{-t}t^{x-1} dt = \int_{1/a}^1 e^{-1/s} s^{1-x} (-s^{-2}) ds = \int_1^{1/a} e^{-1/s} s^{-x-1} ds.$$

Da $|e^{-1/s} s^{-x-1}| \leq s^{-x-1}$ für alle $s \geq 1$ und das uneigentliche Integral $\int_1^\infty s^{-x-1} ds$ für $-x - 1 < -1$ konvergiert, ist auch die Existenz des uneigentlichen Integrals $\int_0^1 f(t) dt$, also der Gamma-Funktion für $0 < x < 1$, gesichert.

Wenn f stetig und unbeschränkt auf $(a, b]$, dann kann man mit Hilfe der Transformation $s = (t - a)^{-1}$, also $t = s^{-1} + a$ und $dt = -s^{-2}ds$,

$$\int_a^b f(t) dt = \int_\infty^{(b-a)^{-1}} f(s^{-1} + a)(-s^{-2}) ds = \int_{(b-a)^{-1}}^\infty f(s^{-1} + a)s^{-2} ds$$

statt des $\int_a^b f(t)dt$ ein uneigentliches Integral mit oberer Grenze ∞ betrachten, wobei die zu integrierende Funktion $f(s^{-1} + a)$ für $s \geq (b - a)^{-1}$ stetig ist. Aus Satz 77 ergibt sich ein Kriterium, bei dem man sich die Ausführung der Transformation spart:

FOLGERUNG 30. Wenn $f : (a, b]$ stetig, $|f(t)| \leq g(t)$ für alle $a < t \leq b$ und das uneigentliche Integral $\int_a^b g(t) dt$ existiert, dann existiert auch $\int_a^b f(t) dt$.

BEISPIEL 126. Für $0 < t \leq 1$ gilt $e^{-t} \leq 1$ und $|e^{-t}t^{x-1}| \leq t^{x-1}$. Da für $1 > x > 0$ das uneigentliche Integral

$$\int_0^1 t^{x-1} dx = \lim_{a \rightarrow 0} \left[\frac{1}{x} t^x \right]_a^1 = \frac{1}{x} - \lim_{a \rightarrow 0} \frac{a}{x} = \frac{1}{x}$$

existiert, existiert auch das uneigentliche Integral $\int_0^1 f(t) dt$ mit $f(t) = e^{-t}t^{x-1}$, also die Gamma-Funktion für $0 < x < 1$.

9.3. Rekursionsformel.

SATZ 78. Für alle $x > 0$ gilt

$$(92) \quad \Gamma(x+1) = x\Gamma(x).$$

BEWEIS. Mit $f'(t) = e^{-t}$, $f(t) = -e^{-t}$, $g(t) = t^x$ und $g'(t) = xt^{x-1}$ und partieller Integration erhält man, da die Existenz der uneigentlichen Integrale schon gesichert ist,

$$\begin{aligned} \int e^{-t}t^x dt &= -e^{-t}t^x - \int -e^{-t}xt^{x-1} dt \\ \Gamma(x+1) &= \int_0^\infty e^{-t}t^x dt = [-e^{-t}t^x]_0^\infty - \int -e^{-t}xt^{x-1} dt \\ &= -\lim_{b \rightarrow \infty} \frac{b^x}{e^b} + e^{-0}0^x + x \int e^{-t}t^{x-1} dt = x\Gamma(x), \end{aligned}$$

da die Exponentialfunktion schneller wächst als jedes Polynom. $\square$

FOLGERUNG 31. Für alle $n \in \mathbb{N}$ gilt $\Gamma(n+1) = n!$

BEWEIS. Es gilt $\Gamma(1) = 1 = 0!$, $\Gamma(2) = 1 = 1!$ und $\Gamma(3) = 2 = 2!$ (siehe Beispiele 123, 124 und 125). Die Behauptung folgt mit vollständiger Induktion aus der Rekursionsformel in Satz 78. $\square$

Die Gamma-Funktion interpoliert also $n!$ für $n \in \mathbb{N}$. Sie ist eine kontinuierliche Version des Terms $n!$, also könnte man $\Gamma(x) = x!$ für $x \in \mathbb{R}$ schreiben. Wir werden später sehen, dass die Gamma-Funktion differenzierbar ist.

KAPITEL IX

Ableitungen höherer Ordnung

Ist eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ in jedem Punkt $x \in (a, b)$ differenzierbar, so definiert die Zuordnung $x \mapsto f'(x)$ eine Funktion $f' : (a, b) \rightarrow \mathbb{R}$. Die Funktion f' heißt Ableitung der Funktion f . Ist f' eine stetige Funktion, so heißt f stetig differenzierbar.

Für $n \in \mathbb{N}$ definieren wir die **n -te Ableitung** einer Funktion rekursiv. Eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ ist n -mal differenzierbar, falls die Funktion $(n-1)$ -mal stetig differenzierbar ist und die Funktion $f^{(n-1)} : (a, b) \rightarrow \mathbb{R}$ in allen Punkten $x \in (a, b)$ differenzierbar ist. Die n -te Ableitung der Funktion f wird mit $f^{(n)}$ bezeichnet.

Wir bezeichnen die Menge aller auf dem Intervall (a, b) n -mal stetig differenzierbaren Funktionen mit $\mathcal{C}^n((a, b))$. Dabei ist $\mathcal{C}^0((a, b))$ die Menge der auf dem Intervall (a, b) stetigen Funktionen. Die Menge aller auf dem Intervall (a, b) beliebig oft differenzierbaren Funktionen wird mit $\mathcal{C}^\infty((a, b))$ bezeichnet.

Verkettungen von Polynomen, e^x , $\sin x$ und $\cos x$ sind Elemente der Menge $\mathcal{C}^\infty(\mathbb{R})$. Verkettungen, Summen, Produkte, Quotienten und Umkehrfunktionen beliebig oft differenzierbarer Funktionen sind auf offenen Intervallen in ihrem Definitionsgebiet beliebig oft differenzierbar.

Für eine auf einem abgeschlossenen Intervall definierte Funktion $f : [a, b] \rightarrow \mathbb{R}$ existieren die n -ten Ableitungen in den Randpunkten a und b , falls die rechts- bzw. linksseitigen Grenzwerte der Differenzenquotienten existieren:

$$f^{(n)}(a) := \lim_{x \rightarrow a, x > a} \frac{f^{(n-1)}(x) - f^{(n-1)}(a)}{x - a}$$

$$f^{(n)}(b) := \lim_{x \rightarrow b, x < b} \frac{f^{(n-1)}(x) - f^{(n-1)}(b)}{x - b}$$

BEISPIEL 127 ($f \in \mathcal{C}^1(\mathbb{R})$ und $f \notin \mathcal{C}^2(\mathbb{R})$). Die Stammfunktion

$$f(x) := \int_0^x |t| dt = \begin{cases} \frac{1}{2}x^2 & , \text{ falls } x \geq 0 \\ -\frac{1}{2}x^2 & , \text{ falls } x < 0 \end{cases}$$

ist als Stammfunktion einer stetigen Funktion differenzierbar. Es gilt $f'(x) = |x|$ (siehe Beispiel 88). Die Ableitung f' ist stetig aber in $x = 0$ nicht differenzierbar.

BEISPIEL 128 (f differenzierbar aber $\notin \mathcal{C}^1(\mathbb{R})$). Die Funktion

$$f(x) = \begin{cases} x^2 \sin(1/x) & , \text{ falls } x \neq 0 \\ 0 & , \text{ falls } x = 0 \end{cases}$$

ist für alle $x \neq 0$ differenzierbar. Aus der Kettenregel folgt für $x \neq 0$

$$f'(x) = 2x \sin(1/x) + x^2 \cos(1/x)(-x^{-2}) = 2x \sin(1/x) - 2 \cos(1/x).$$

Auch die Ableitung der Funktion in $x = 0$ existiert, denn

$$\lim_{x \rightarrow 0} \frac{f(x) - f(0)}{x - 0} = \lim_{x \rightarrow 0} \frac{x^2 \sin(1/x)}{x} = \lim_{x \rightarrow 0} x \sin(1/x) = 0 = f'(0).$$

Die Ableitung f' ist in $x = 0$ nicht stetig, denn der Grenzwert $\lim_{x \rightarrow 0} \cos(1/x)$ existiert nicht.

1. Bedeutung der zweiten Ableitung

Das Vorzeichen der zweiten Ableitung verrät, auf welcher Seite der Tangente an einen Funktionsgraphen der Graph verläuft.

Wenn $f''(x) > 0$ für alle $x \in (a, b)$, dann liegt der Graph der Funktion $f : (a, b) \rightarrow \mathbb{R}$ immer oberhalb der Tangenten an den Graph, z.B. für $f(x) = x^2$ mit $f''(x) \equiv 2 > 0$ (siehe Abbildung 1).

LEMMA 25. *Es sei $f : (a, b) \rightarrow \mathbb{R}$ eine zweimal differenzierbare Funktion. Wenn $f''(x) > 0$ für alle $x \in (a, b)$, dann gilt $f(x) > f(x_0) + f'(x_0)(x - x_0)$ für alle $x_0 \neq x \in (a, b)$.*

Wenn $f''(x) < 0$ für alle $x \in (a, b)$, dann gilt $f(x) < f(x_0) + f'(x_0)(x - x_0)$ für alle $x_0 \neq x \in (a, b)$.

BEWEIS. Für jedes $x_0 \in (a, b)$ besitzt die Funktion

$$h(x) := f(x) - f(x_0) - f'(x_0)(x - x_0)$$

in x_0 ein lokales Minimum, weil $h'(x) = f'(x) - f'(x_0)$, $h'(x_0) = 0$, $h''(x) = f''(x)$ und $h''(x_0) > 0$. Da $h''(x) = f''(x) > 0$ für alle $x \in (a, b)$, ist h' streng monoton wachsend und hat deshalb nur eine Nullstelle. Daraus folgt, dass h nur eine lokale Extremstelle x_0 hat und diese lokale Extremstelle auch eine globale Extremstelle ist, also $h(x) \geq 0 = h(x_0)$ für alle $x \in (a, b)$.

Die zweite Behauptung folgt aus der ersten, da $(-f)'' = -f''$. $\square$

Wenn $f''(x) < 0$ für alle $x \in (a, b)$, dann liegt der Graph der Funktion $f : (a, b) \rightarrow \mathbb{R}$ immer unterhalb der Tangenten an den Graph, z.B. für $f(x) = \ln x$ mit $f''(x) = -1/x^2 < 0$ für alle $x > 0$ (siehe Abbildung 2).

Für die Funktion $f(x) = \arctan x$ gilt

$$f'(x) = (1 + x^2)^{-1} \text{ und } f''(x) = -2x(1 + x^2)^{-1}.$$

Die zweite Ableitung wechselt in $x = 0$ das Vorzeichen. Es gilt $f''(0) = 0$, $f''(x) < 0$ für alle $x > 0$ und $f''(x) > 0$ für alle $x < 0$. Die Tangente an den Graph in $f(0)$ schneidet den Graph (siehe Abbildung 3).

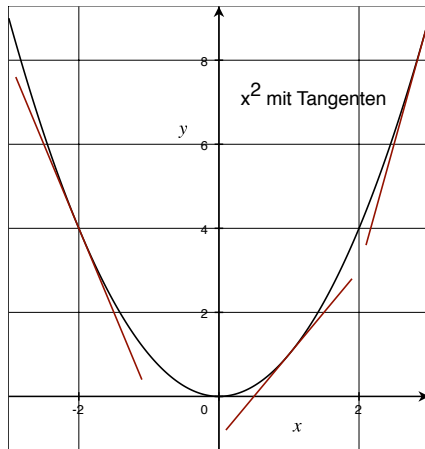


ABBILDUNG 1

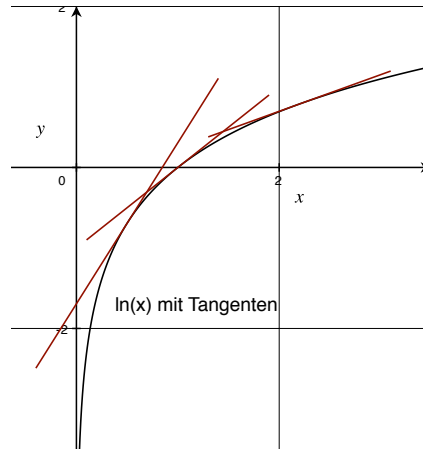


ABBILDUNG 2

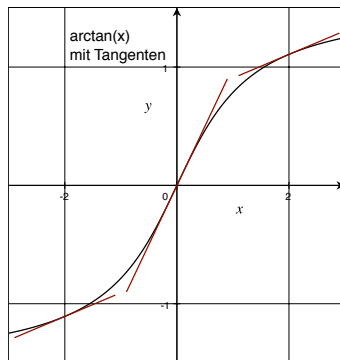
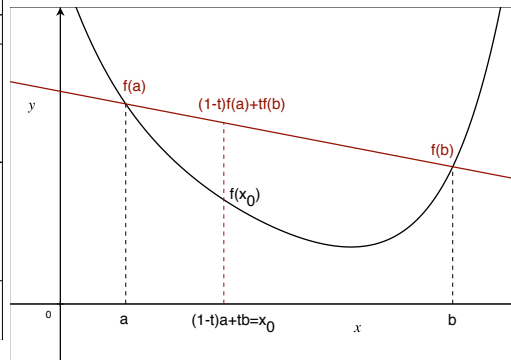
ABBILDUNG
3

ABBILDUNG 4

BEMERKUNG 62. Das Vorzeichen und das Wachstum der zweiten Ableitung einer Funktion kann man am Graph der Funktion erkennen. Jedoch ist der Wert $f''(x)$ nicht direkt ablesbar. Abbildung 7 zeigt den Graph der Stammfunktion $f(x) = \int_0^x |t| dt$, die in $x = 0$ nicht zweimal differenzierbar ist. Auch die Tatsache, dass $f''(x)$ nicht existiert, ist im Graph der Funktion unsichtbar. $\square$

Man kann den Funktionswerte $f(x)$ für $x \in [a, b]$ auch mit den Punkten auf der Verbindungsstrecke zwischen $f(a)$ und $f(b)$ vergleichen (siehe Abbildung 4).

SATZ 79. Wenn $f''(x) > 0$ für alle $x \in [a, b]$, dann gilt für alle $0 \leq t \leq 1$

$$(1-t)f(a) + tf(b) \geq f((1-t)a + tb).$$

Wenn $f''(x) < 0$ für alle $x \in [a, b]$, dann gilt für alle $0 \leq t \leq 1$

$$(1-t)f(a) + tf(b) \leq f((1-t)a + tb).$$

BEWEIS. Für $x_0 = a + t(b - a) = (1 - t)a + tb$ folgt aus dem Zwischenwertsatz

$$f(x_0) - f(a) = f'(\xi)(x_0 - a) \text{ und } f(b) - f(x_0) = f'(\eta)(b - x_0)$$

für ein $\xi \in [a, x_0]$ und $\eta \in [x_0, b]$. Wenn $f''(x) > 0$ für alle $x \in [a, b]$, dann ist f' streng monoton wachsend, also $f'(\xi) \leq f'(\eta)$. Daraus folgt, da $x_0 - a, b - x_0 \geq 0$

$$\begin{aligned} \frac{f(x_0) - f(a)}{x_0 - a} &\leq \frac{f(b) - f(x_0)}{b - x_0} \\ (b - x_0)(f(x_0) - f(a)) &\leq (f(b) - f(x_0))(x_0 - a) \\ (b - x_0 + x_0 - a)f(x_0) &\leq f(b)(x_0 - a) + (b - x_0)f(a) \\ (b - a)f(x_0) &\leq f(b)t(b - a) + (b - a)(1 - t)f(a) \\ f(x_0) &\leq (1 - t)f(a) + tf(b). \end{aligned}$$

Die zweite Behauptung folgt aus der ersten, da $(-f)'' = -f''$. $\square$

FOLGERUNG 32. Wenn $f''(x) > 0$ für alle $x \in (a, b)$, $x_k \in (a, b)$, $\lambda_k \in [0, 1]$ und $\sum_{k=1}^n \lambda_k = 1$, dann gilt

$$(93) \quad \sum_{k=1}^n \lambda_k f(x_k) \geq f\left(\sum_{k=1}^n \lambda_k x_k\right).$$

Wenn $f''(x) < 0$ für alle $x \in (a, b)$, $x_k \in (a, b)$, $\lambda_k \in [0, 1]$ und $\sum_{k=1}^n \lambda_k = 1$, dann gilt

$$(94) \quad \sum_{k=1}^n \lambda_k f(x_k) \leq f\left(\sum_{k=1}^n \lambda_k x_k\right).$$

BEWEIS. Wir beweisen die erste Ungleichung mit vollständiger Induktion. Die zweite Ungleichung folgt aus der ersten, da $(-f)'' = -f''$. Die Ungleichung gilt für $n = 1$ und $n = 2$. Man kann $\lambda_k > 0$ für alle $k = 1, \dots, n+1$ annehmen. Mit $\lambda := \sum_{k=1}^n \lambda_k < 1$ erhält man aus der Induktionsvoraussetzung und der Ungleichung für $n = 2$

$$\begin{aligned} \sum_{k=1}^n \frac{\lambda_k}{\lambda} f(x_k) &\geq f\left(\sum_{k=1}^n \frac{\lambda_k}{\lambda} x_k\right) \\ \sum_{k=1}^n \lambda_k f(x_k) &\geq \lambda f\left(\frac{1}{\lambda} \sum_{k=1}^n \lambda_k x_k\right) \\ \sum_{k=1}^n \lambda_k f(x_k) + \lambda_{n+1} f(x_{n+1}) &\geq \lambda f\left(\frac{1}{\lambda} \sum_{k=1}^n \lambda_k x_k\right) + \lambda_{n+1} f(x_{n+1}) \\ &\geq f\left(\sum_{k=1}^n \lambda_k x_k + \lambda_{n+1} x_{n+1}\right), \end{aligned}$$

da $\lambda + \lambda_{n+1} = 1$. $\square$

Die Terme $\sum_{k=1}^n \lambda_k x_k$ und $\sum_{k=1}^n \lambda_k f(x_k)$ sind **gewichtete Mittelwerte** der Argumente x_k bzw. der Funktionswerte $f(x_k)$. Falls $\lambda_1 = \dots = \lambda_n = 1/n$, so erhält man die ungewichteten Mittelwerte.

BEISPIEL 129. Für die Funktion $f(x) = \ln x$ gilt $f''(x) = -1/x^2 < 0$ für alle $x > 0$. Daraus folgt, falls $x_k > 0$, $\lambda_k \in [0, 1]$ für alle k und $\sum_{k=1}^n \lambda_k = 1$,

$$(95) \quad \lambda_1 \ln x_1 + \dots + \lambda_n \ln x_n \leq \ln(\lambda_1 x_1 + \dots + \lambda_n x_n).$$

BEISPIEL 130. Für die Funktion $f(x) = e^x$ gilt $f''(x) = e^x > 0$ für alle $x \in \mathbb{R}$. Daraus folgt, falls $\lambda_k \in [0, 1]$ für alle k und $\sum_{k=1}^n \lambda_k = 1$,

$$(96) \quad \lambda_1 e^{x_1} + \dots + \lambda_n e^{x_n} \geq e^{\lambda_1 x_1 + \dots + \lambda_n x_n}.$$

Für $a_j > 0$ folgt mit $x_j := \ln a_j$ und $\lambda_j = 1/n$ die allgemeine Ungleichung zwischen arithmetischem und geometrischem Mittel (siehe Satz 12):

$$\frac{1}{n} \sum_{j=1}^n a_j = \frac{1}{n} \sum_{j=1}^n e^{\ln a_j} \geq e^{\frac{1}{n} \sum_{j=1}^n \ln a_j} = \left(e^{\sum_{j=1}^n \ln a_j} \right)^{1/n} = \sqrt[n]{\prod_{j=1}^n e^{\ln a_j}} = \sqrt[n]{\prod_{j=1}^n a_j}.$$

2. Taylorpolynome

Das **Taylorpolynom** $T_{f,x_0,n}(x)$ vom Grad $n \in \mathbb{N}$ einer n -mal differenzierbaren Funktion $f : (a, b) \rightarrow \mathbb{R}$ um den Punkt $x_0 \in (a, b)$ ist

$$(97) \quad T_{f,x_0,n}(x) := \sum_{k=0}^n \frac{1}{k!} f^{(k)}(x_0) (x - x_0)^k.$$

Das Taylorpolynom hängt nur von den Ableitungen der Funktion f im Punkt x_0 ab. Es ist das eindeutige Polynom p mit den Eigenschaften $\deg p \leq n$ und $p^{(k)}(x_0) = f^{(k)}(x_0)$ für alle $k = 0, 1, \dots, n$.

BEISPIEL 131 (Taylorpolynom der Exponentialfunktion). Für $f(x) = e^x$ und $x_0 = 0$ gilt $f^{(k)}(x) = e^x$ und $f^{(k)}(x_0) = e^0 = 1$ für alle $k \in \mathbb{N}$. Daraus folgt

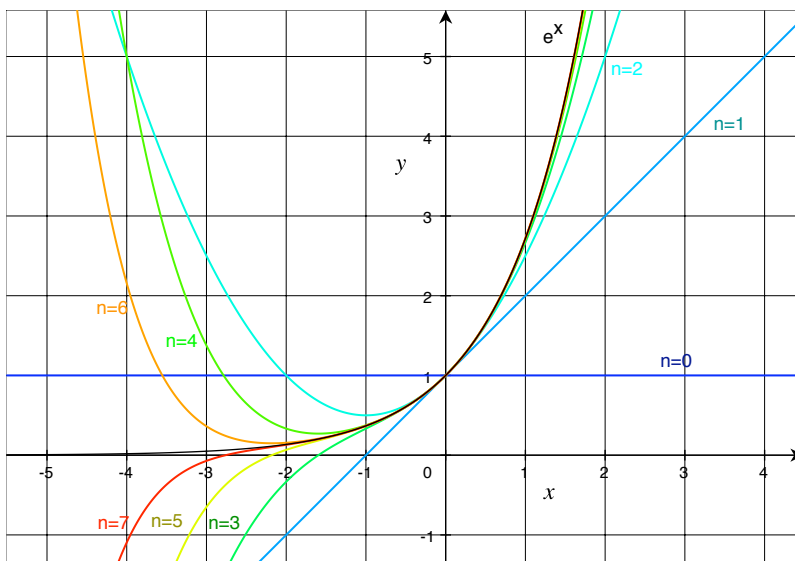
$$T_{e^x, x_0=0, n}(x) = 1 + x + \frac{x^2}{2} + \frac{x^3}{3!} + \dots + \frac{x^n}{n!} = \sum_{k=0}^n \frac{x^k}{k!}.$$

Abbildung 5 zeigt die Taylorpolynome von e^x um $x_0 = 0$ vom Grad $0, 1, \dots, 7$.

BEISPIEL 132 (Taylorpolynom der Winkelfunktionen). Für $f(x) = \sin x$ gilt $f'(x) = \cos x$, $f''(x) = -\sin x$, also $f^{(2m)}(x) = (-1)^m \sin x$ und $f^{(2m+1)}(x) = (-1)^m \cos x$ für alle $m \in \mathbb{N}$.

Für $x_0 = 0$ gilt und $f^{(2m)}(x_0) = 0$ und $f^{(2m+1)}(x_0) = (-1)^m$ für alle $m \in \mathbb{N}$. Daraus folgt $T_{\sin x, x_0=0, 2n+1} = T_{\sin x, x_0=0, 2n+2}$ und

$$T_{\sin x, x_0=0, 2n+1}(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \dots + (-1)^n \frac{x^{2n+1}}{(2n+1)!} = \sum_{m=0}^n \frac{(-1)^m}{(2m+1)!} x^{2m+1}.$$

ABBILDUNG 5. Taylorpolynome von e^x um $x_0 = 0$

Für $x_0 = \pi/2$ gilt und $f^{(2m)}(x_0) = (-1)^m$ und $f^{(2m+1)}(x_0) = 0$ für alle $m \in \mathbb{N}$. Daraus folgt $T_{\sin x, x_0=\pi/2, 2n} = T_{\sin x, x_0=\pi/2, 2n+1}$ und

$$\begin{aligned} T_{\sin x, x_0=\pi/2, 2n}(x) &= 1 - \frac{(x - \pi/2)^2}{2!} + \frac{(x - \pi/2)^4}{4!} + \dots + (-1)^n \frac{(x - \pi/2)^{2n}}{(2n)!} \\ &= \sum_{m=0}^n \frac{(-1)^m}{(2m)!} (x - \pi/2)^{2m}. \end{aligned}$$

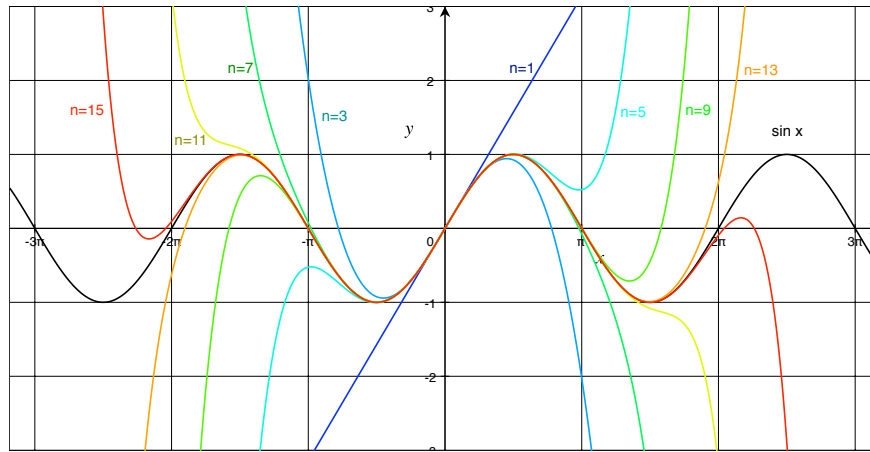
Wegen $\cos x = \sin(x - \pi/2)$ erhält man das Taylorpolynom von $\cos x$ um $x_0 = 0$

$$T_{\cos x, x_0=0, 2n}(x) = \sum_{m=0}^n \frac{(-1)^m}{(2m)!} x^{2m}.$$

Abbildung 6 zeigt die Taylorpolynome der Funktion $\sin x$ um den Punkt $x_0 = 0$ vom Grad 1, 3, 5, ..., 15.

BEISPIEL 133 (Taylorpolynom eines Polynoms). Für $f(x) = x^3 + 2x^2 - 4x + 6$ und $x_0 = 1$ gilt $f(1) = 5$, $f'(x) = 3x^2 + 4x - 4$, $f'(1) = 3$, $f''(x) = 6x + 4$, $f''(1) = 10$, $f^{(3)}(x) \equiv 6$ und $f^{(3)}(1) = 6$. Daraus folgt

$$\begin{aligned} T_{f, x_0=1, 0}(x) &\equiv 5 \\ T_{f, x_0=1, 1}(x) &= 5 + 3(x - 1) \\ T_{f, x_0=1, 2}(x) &= 5 + 3(x - 1) + 5(x - 1)^2 \\ T_{f, x_0=1, 3}(x) &= 5 + 3(x - 1) + 5(x - 1)^2 + (x - 1)^3 \end{aligned}$$

ABBILDUNG 6. Taylorpolynome von $\sin x$ um $x_0 = 0$

und $T_{f,x_0=1,n}(x) = T_{f,x_0=1,3}(x)$ für alle $n \geq 3$. Es gilt $f(x) = T_{f,x_0=1,3}(x)$. Abbildung 7 zeigt die Taylorpolynome des Polynoms $f(x)$ um $x_0 = 1$ vom Grad 0, 1, 2 und 3.

BEISPIEL 134 (Taylorpolynom des Logarithmus um $x_0 > 0$). Für $f(x) = \ln x$ und $x_0 > 0$ gilt $f(x_0) = \ln x_0$, $f'(x) = x^{-1}$ und $f^{(k)}(x) = (-1)^{k-1}(k-1)!x^{-k}$ für alle $k \in \mathbb{N}$ mit $k > 0$. Daraus folgt

$$\begin{aligned} T_{\ln x, x_0, n}(x) &= \ln x_0 + \frac{1}{x_0}(x - x_0) - \frac{1}{2x_0^2}(x - x_0)^2 + \frac{2!}{3!x_0^3}(x - x_0)^3 - \dots \\ &\quad + \frac{(-1)^{n-1}(n-1)!}{n!x_0^n}(x - x_0)^n \\ &= \ln x_0 + \sum_{k=1}^n \frac{(-1)^{k-1}(k-1)!}{k!x_0^k}(x - x_0)^k = \ln x_0 - \sum_{k=1}^n \frac{(-1)^k}{kx_0^k}(x - x_0)^k. \end{aligned}$$

2.1. Abschätzung des Restgliedes. Ist die Funktion f sogar $n+1$ -mal stetig differenzierbar, so kann man den Unterschied zwischen $f(x)$ und dem Wert des Taylorpolynoms in x abschätzen.

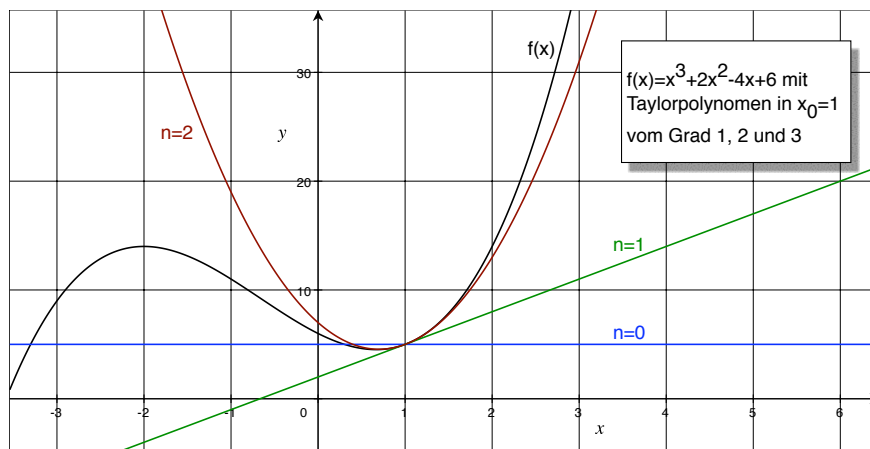
SATZ 80. Wenn $f \in \mathcal{C}^{n+1}(a, b)$ und $x_0, x \in (a, b)$, dann gilt

$$(98) \quad f(x) - T_{f, x_0, n}(x) = \frac{f^{(n+1)}(\xi)}{(n+1)!}(x - x_0)^{n+1}$$

für ein ξ zwischen x_0 und x .

BEWEIS. Für ein festes $x \in (a, b)$ betrachten wir das Taylorpolynom von f vom Grad n um einen variablen Punkt t . Für die Funktion

$$g(t) := f(x) - T_{f, t, n}(x) \text{ mit } T_{f, t, n}(x) = \sum_{k=0}^n \frac{1}{k!} f^{(k)}(t)(x - t)^k$$

ABBILDUNG 7. Taylorpolynome eines Polynoms um $x_0 = 1$

gilt $g(x) = 0$ und $g(x_0) = f(x) - T_{f,x_0,n}(x) =: R(x)$. Da

$$g'(t) = - \sum_{k=0}^n \frac{1}{k!} f^{(k+1)}(t)(x-t)^k - \frac{k}{k!} f^{(k)}(t)(x-t)^{k-1} = - \frac{1}{k!} f^{(n+1)}(t)(x-t)^n,$$

gilt $G(x_0) = G(x) = g(x) = 0$ für

$$G(t) := g(t) - g(x_0) \left(\frac{x-t}{x-x_0} \right)^{n+1}.$$

Aus dem Zwischenwertsatz folgt, dass $G'(\xi) = 0$ für ein ξ zwischen x_0 und x , also

$$\begin{aligned} 0 &= G'(\xi) = g'(\xi) + (n+1)g(x_0) \frac{(x-\xi)^n}{(x-x_0)^{n+1}} \\ &= -\frac{1}{k!} f^{(n+1)}(\xi)(x-\xi)^n + (n+1)R(x) \frac{(x-\xi)^n}{(x-x_0)^{n+1}} \\ R(x) &= \frac{1}{(n+1)!} f^{(n+1)}(\xi)(x-x_0)^{n+1}. \end{aligned}$$

□

BEMERKUNG 63. Für $n = 1$ liefert Satz 80 gerade den Zwischenwertsatz.

□

FOLGERUNG 33. Wenn $f \in \mathcal{C}^{n+1}(a, b)$ und $x_0 \in (a, b)$ und $M \geq |f^{(n+1)}(x)|$ für alle $x \in (a, b)$, dann gilt für alle $x \in (a, b)$ die Abschätzung

$$(99) \quad |f(x) - T_{f,x_0,n}(x)| \leq \frac{M}{(n+1)!} (x-x_0)^{n+1}$$

BEMERKUNG 64. Falls eine Abschätzung der Ableitungen wie in Folgerung 33 existiert, dann approximiert das Taylorpolynom die Funktion f

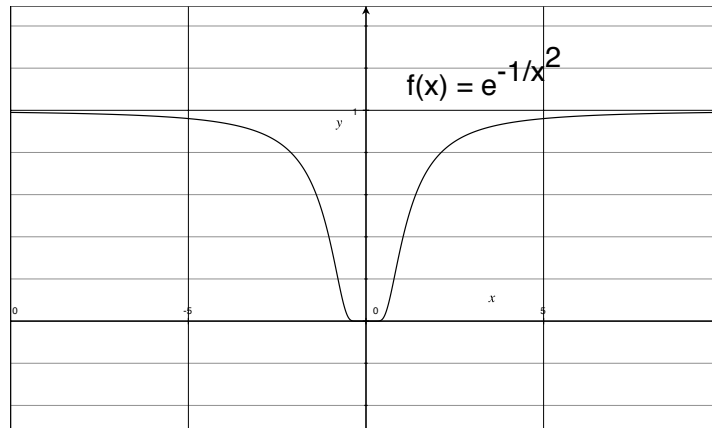


ABBILDUNG 8

besonders gut, wenn n groß ist und x nahe an x_0 liegt, die Differenz $|x - x_0|$ also klein ist. $\square$

BEISPIEL 135. Für $|x| \leq 1$ gilt $0 < e^x < 3$, also $|f^{(3)}(x)| < 3$. Daraus folgt

$$|e^x - T_{e^x, x_0=0, 2}(x)| = |e^x - (1 + x + x^2/2)| \leq \frac{x^3}{3!} \cdot 3 = \frac{x^3}{2}.$$

Für $|x| < 0,1$ ist der Unterschied zwischen e^x und $1 + x + x^2/2$ kleiner als 0,0005.

BEISPIEL 136. Für die Funktion $f(x) = \sin x$ gilt $|f^{(n)}(x)| \leq 1$, da $|\sin x| \leq 1$ und $|\cos x| \leq 1$ für alle $x \in \mathbb{R}$. Daraus folgt

$$|\sin x - T_{\sin x, x_0=0, n}(x)| \leq \frac{x^{n+1}}{(n+1)!},$$

also kann man für ein festes $x \in \mathbb{R}$ die reelle Zahl $\sin x$ beliebig gut durch den Wert des Taylorpolynoms $T_{\sin x, x_0=0, n}(x)$ approximieren, wenn man n sehr groß wählt.

2.2. Taylorpolynom der Funktion e^{-1/x^2} . Die Funktion $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$(100) \quad \varphi(x) = e^{-1/x^2} \text{ für } x \neq 0 \text{ und } \varphi(0) = 0$$

ist stetig, denn $\lim_{x \rightarrow 0} -1/x^2 = -\infty$ und $\lim_{x \rightarrow -\infty} e^x = 0$.

Da $-1/x^2 < 0$ für alle $x \in \mathbb{R}$, gilt $0 \leq f(x) < 1$ für alle $x \in \mathbb{R}$. Abbildung 8 zeigt den Graph der Funktion φ . Die Funktion φ ist für $x \neq 0$ beliebig oft differenzierbar, z.B.

$$\varphi'(x) = e^{-1/x^2} (-x^{-2})' = 2x^{-3} e^{-1/x^2}.$$

Es existieren Polynome p_n mit $\varphi^{(n)}(x) = p_n(1/x)e^{-1/x^2}$ für $x \neq 0$. Da für alle $k \in \mathbb{N}$ gilt

$$\lim_{x \rightarrow 0} x^{-k} e^{-1/x^2} = \lim_{x \rightarrow 0} \frac{x^{-k}}{e^{(x^{-2})}} = \lim_{u=1/x \rightarrow \pm\infty} \frac{u^k}{e^{2k}} = 0,$$

folgt $\varphi^{(n)}(0) = 0$ für alle $n \in \mathbb{N}$. Damit sind alle Taylorpolynome von φ um $x_0 = 0$ trivial, also $T_{\varphi, x_0=0, n}(x) \equiv 0$ für alle $n \in \mathbb{N}$.

3. Newtonverfahren

Das Newtonverfahren ist eine Methode zur näherungsweisen Bestimmung der Nullstellen einer differenzierbaren Funktion f . Falls $f(a) < 0$ und $f(b) > 0$, so befindet sich zwischen a und b eine Nullstelle c der Funktion f , da f stetig ist.

Man könnte diese Nullstelle mit Hilfe des Halbierungsverfahrens (siehe Abschnitt 4.2) suchen. Dies funktioniert immer. Allerdings muß man viele Iterationsschritte durchführen, um eine hohe Genauigkeit zu erhalten.

3.1. Idee. Die Idee des Newtonverfahrens besteht nun darin, dass man die Funktion f durch eine Tangente $g(x) := f(x_n) + f'(x_n)(x - x_n)$ an f ersetzt. Falls x_n schon nahe an der gesuchten Nullstelle c liegt, so ist g eine gute Näherung für die Funktion f für Argumente, die nahe an x_n und damit nahe an der Nullstelle c liegen. Die Nullstelle der linearen Funktion g ist leicht zu berechnen. Es gilt

$$g(x) = 0 \Leftrightarrow x = x_n - \frac{f(x_n)}{f'(x_n)}, \text{ falls } f'(x_n) \neq 0.$$

Man hofft, dass die Nullstelle der Funktion g noch näher an der Nullstelle von f liegt als x_n . Die Tangente $g(x)$ besitzt nur dann eine Nullstelle, wenn $f'(x_n) \neq 0$ oder $f(x_n) = 0$.

3.2. Rekursionsformel. Wenn $f : (a, b) \rightarrow \mathbb{R}$ differenzierbar ist und eine Nullstelle im Intervall (a, b) besitzt, dann wählt man ein $x_1 \in (a, b)$ (möglichst nahe an der Nullstelle) und berechnet rekursiv für $n \in \mathbb{N}$

$$(101) \quad x_{n+1} := x_n - \frac{f(x_n)}{f'(x_n)}.$$

BEISPIEL 137 (Newtonverfahren für $\sqrt{2}$). Die irrationale Zahl $\sqrt{2}$ ist die positive Nullstelle der Funktion $f(x) = x^2 - 2$. Da $f(1) = -1 < 0$ und $f(2) = 2 > 0$, liegt $\sqrt{2}$ im Intervall $[1, 2]$. Es gilt $f'(x) = 2x$. Wir starten mit $x_1 = 2$. Die Rekursionsformel lautet

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

Wir erhalten auf 14 Nachkommastellen gerundet

$$\begin{aligned}x_2 &= 1,5 \\x_3 &= 1,41666666666667 \\x_4 &= 1,41421568627451 \\x_5 &= 1,41421356237469 \\x_6 &= 1,41421356237309\end{aligned}$$

und danach keine Änderung der angezeigten Nachkommastellen mehr.

3.3. Güte der Approximation. Wenn $f(c) = 0$, dann folgt aus dem Zwischenwertsatz $f(x_n) = f(x_n) - f(c) = f'(\xi)(x_n - c)$ für ein ξ zwischen x_n und c . Wenn $M \leq |f'(x)|$ für alle x zwischen x_n und c , dann folgt

$$|x_n - c| \leq \frac{|f(x_n)|}{M}.$$

Wenn die Ableitung der Funktion in der Nähe der Nullstelle c sehr groß ist oder $f(x_n)$ schon sehr klein ist, liegt x_n nahe an der gesuchte Nullstelle.

3.4. Konvergenzkriterien. In bestimmten Fällen liefert die Iteration des Newtonverfahrens streng monotone und beschränkte Folgen, die dann konvergieren und der Grenzwert ist die gesuchte Nullstelle.

SATZ 81. Wenn $f''(x) > 0$ für alle $x \in [a, b]$, $f(a) < 0$ und $f(b) > 0$, dann liefert das Newtonverfahren mit dem Startwert $x_1 = b$ eine streng monoton fallende Folge, die gegen die Nullstelle von f im Intervall (a, b) konvergiert.

Wenn $f''(x) > 0$ für alle $x \in [a, b]$, $f(a) > 0$ und $f(b) < 0$, dann liefert das Newtonverfahren mit dem Startwert $x_1 = a$ eine streng monoton wachsende Folge, die gegen die Nullstelle von f im Intervall (a, b) konvergiert.

Wenn $f''(x) < 0$ für alle $x \in [a, b]$, $f(a) < 0$ und $f(b) > 0$, dann liefert das Newtonverfahren mit dem Startwert $x_1 = a$ eine streng monoton wachsende Folge, die gegen die Nullstelle von f im Intervall (a, b) konvergiert.

Wenn $f''(x) < 0$ für alle $x \in [a, b]$, $f(a) > 0$ und $f(b) < 0$, dann liefert das Newtonverfahren mit dem Startwert $x_1 = b$ eine streng monoton fallende Folge, die gegen die Nullstelle von f im Intervall (a, b) konvergiert.

BEWEIS. Wenn $f''(x) > 0$ für alle $x \in [a, b]$, $f(a) < 0$ und $f(b) > 0$, dann hat f im Intervall (a, b) genau eine Nullstelle, denn wenn $f(c_1) = f(c_2) = 0$, $f(x) < 0$ für alle $x < c_1$, dann $f'(x) > 0$ für alle $x \geq c_1$, was der Existenz eines lokalen Maximums zwischen c_1 und c_2 widerspricht.

Wenn $f''(x) > 0$ für alle $x \in [a, b]$ und $f(x_n) > 0$, dann liegt die Tangente $f(x_n) + f'(x_n)(x - x_n)$ immer unterhalb des Graphen der Funktion f . Daraus folgt $f(x_{n+1}) > 0$. Da $f'(x) > 0$ für alle $x > c$, folgt $c < x_{n+1} < x_n$ für alle n . $\square$

BEISPIEL 138. Für die Funktion $f(x) = x^2 - 2$ gilt $f'(x) = 2x > 0$ für alle $x \in [1, 2]$, $f(1) = -1 < 0$ und $f(2) = 2 > 0$. Also liefert das

Newtonverfahren mit Startwert $x_1 = 2$ wie in Beispiel 137 eine monoton fallende Folge, die gegen $\sqrt{2}$ konvergiert.

3.5. Abbruch, Divergenz oder Konvergenz gegen falsche Nullstelle. Sind die Voraussetzungen des Satzes 81 nicht erfüllt, so kann es passieren, dass nicht alle Folgenglieder definiert sind, die Folge nicht konvergiert oder gegen eine Nullstelle konvergiert, die man nicht approximieren wollte.

Die Rekursionsformel des Newtonverfahrens ist nur anwendbar, wenn $f'(x_n) \neq 0$ für alle $n \in \mathbb{N}$. Dies kann man für den Startwert, eventuell durch kleine Änderungen, absichern. Besitzt f' Nullstellen, so kann man jedoch nicht abschätzen, ob $f'(x_n) \neq 0$ für alle in der Folge auftretenden x_n gilt.

Ist die Funktion f nicht für alle $x \in \mathbb{R}$ definiert, so kann es passieren, dass ein Folgenglied x_n nicht mehr im Definitionsgebiet der Funktion liegt und die Rekursion damit abbricht.

Auch wenn die Rekursionsformel eine Folge $\{x_n\}$ erzeugt, so ist es möglich, dass diese Folge nicht gegen eine reelle Zahl konvergiert. Besitzt eine Funktion zwei verschiedene Nullstellen und startet das Newtonverfahren sehr nah an der einen Nullstelle, so kann die Folge auch gegen die andere Nullstelle konvergieren.

KAPITEL X

Reihen

Eine **Reihe** $\sum_{n=0}^{\infty} a_n$ ist die Folge $\{s_n\}_{n \in \mathbb{N}}$ der Partialsummen

$$(102) \quad s_n := \sum_{k=0}^n a_k = a_0 + a_1 + a_2 + \dots + a_n$$

der Folge $\{a_n\}_{n \in \mathbb{N}}$. Eine Reihe ist also eine Summe, die aus unendlich vielen Summanden, den Folgengliedern a_n , besteht. Addiert man endlich viele reelle Zahlen, so erhält man immer eine eindeutige reelle Zahl, z.B. die Partialsummen s_n . Für unendlich viele Summanden kann man nicht sicher sein, dass die „Summe“ existiert. Eine Reihe $\sum_{n=0}^{\infty} a_n$ heißt **konvergent**, falls die dazugehörige Folge der Partialsummen $\{s_n\}$ mit $s_n = \sum_{k=0}^n a_k$ gegen eine reelle Zahl konvergiert. Dieser Grenzwert der Partialsummenfolge heißt dann **Wert** bzw. Summe der Reihe,

$$\sum_{n=0}^{\infty} a_n = \lim_{n \rightarrow \infty} \sum_{k=0}^n a_k = \lim_{n \rightarrow \infty} s_n.$$

Sind alle Summanden positiv, also $a_n \geq 0$ für alle $n \in \mathbb{N}$, so ist die Partialsummenfolge $\{s_n\}$ monoton wachsend. In dem Fall konvergiert $\{s_n\}$ oder $\lim_{n \rightarrow \infty} s_n = \infty$.

1. Die geometrische Reihe

Falls $a_n = q^n$, so existiert eine geschlossene Formel für die Partialsummen s_n (siehe Kapitel II). Für alle $n \in \mathbb{N}$ gilt, falls $q \neq 1$,

$$s_n = 1 + q + q^2 + q^3 + \dots + q^n = \frac{1 - q^{n+1}}{1 - q}.$$

- Wenn $q > 1$, dann gilt $\lim_{n \rightarrow \infty} q^{n+1} = \infty$. Die Partialsummenfolge und damit auch die Reihe $\sum_{n=0}^{\infty} q^n$ konvergieren nicht.
- Wenn $-1 < q < 1$, dann gilt $\lim_{n \rightarrow \infty} q^{n+1} = 0$. Die Partialsummenfolge und damit auch die Reihe $\sum_{n=0}^{\infty} q^n$ konvergieren, es gilt

$$(103) \quad \sum_{n=0}^{\infty} q^n = 1 + q + q^2 + q^3 + \dots = \frac{1}{1 - q} \quad \text{für } -1 < q < 1.$$

- Wenn $q \leq -1$, dann gilt konvergiert die Folge $\{q^{n+1}\}$ nicht. Also konvergieren auch die Partialsummenfolge und die Reihe $\sum_{n=0}^{\infty} q^n$ nicht.

- Wenn $a_n = 1^n = 1$ für alle n , also $q = 1$, gilt $s_n = (n + 1)$ und $\lim_{n \rightarrow \infty} s_n = \infty$. Die Partialsummenfolge und damit auch die Reihe $\sum_{n=0}^{\infty} 1$ konvergieren nicht.

BEISPIEL 139. Für $q = -1/2$ erhält man

$$\sum_{n=0}^{\infty} \left(-\frac{1}{2}\right)^n = \sum_{n=0}^{\infty} \frac{1}{(-2)^n} = 1 - \frac{1}{2} + \frac{1}{4} - \frac{1}{8} + \dots = \frac{1}{1 - (-\frac{1}{2})} = \frac{2}{2+1} = \frac{2}{3}.$$

BEISPIEL 140. Mit der Summenformel für die geometrische Reihe kann man auch Summen der Form $\sum_{n=n_0}^{\infty} q^n$ berechnen. Zum Beispiel erhält man mit Hilfe einer Indexverschiebung

$$\begin{aligned} \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^n &= \left(\frac{2}{3}\right)^2 \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^{n-2} = \left(\frac{2}{3}\right)^2 \sum_{n=0}^{\infty} \left(\frac{2}{3}\right)^n = \left(\frac{2}{3}\right)^2 \cdot \frac{1}{1 - \frac{2}{3}} \\ &= \left(\frac{2}{3}\right)^2 \cdot \frac{3}{3-2} = \frac{4}{3} \end{aligned}$$

oder durch Hinzufügen und Abziehen der fehlenden Summanden

$$\begin{aligned} \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^n &= -1 - \frac{2}{3} + 1 + \frac{2}{3} + \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^n = -\frac{5}{3} + \sum_{n=0}^{\infty} \left(\frac{2}{3}\right)^n = -\frac{5}{3} + \frac{1}{1 - \frac{2}{3}} \\ &= -\frac{5}{3} + \frac{3}{3-2} = -\frac{5}{3} + 3 = \frac{4}{3}. \end{aligned}$$

2. Leibnizkriterium

LEMMA 26. Wenn die Reihe $\sum_{n=0}^{\infty} a_n$ konvergiert, dann bilden die Summanden a_n eine Nullfolge, also $\lim_{n \rightarrow \infty} a_n = 0$.

Lemma 26 liefert ein notwendiges Kriterium für die Konvergenz einer Reihe. Da für $|q| \geq 1$ die Folgen $\{q^n\}_{n \in \mathbb{N}}$ keine Nullfolgen sind, konvergieren die Reihen $\sum_{n=0}^{\infty} q^n$ für $|q| \geq 1$ nicht. Die Summanden der harmonischen Reihe $\sum_{n=1}^{\infty} \frac{1}{n}$ bilden eine Nullfolge, aber die harmonische Reihe divergiert (siehe Abschnitt 5.2).

SATZ 82. Wenn $a_n \geq 0$ für alle $n \in \mathbb{N}$, $\{a_n\}_{n \in \mathbb{N}}$ monoton fallend und $\lim_{n \rightarrow \infty} a_n = 0$, dann konvergiert die alternierende Reihe $\sum_{n=0}^{\infty} (-1)^n a_n$.

BEISPIEL 141 (Alternierende harmonische Reihe). Die Folge $\{\frac{1}{n+1}\}_{n \in \mathbb{N}}$ ist eine Nullfolge. Also konvergiert die alternierende harmonische Reihe. Mit Hilfe der Resultate aus Abschnitt 4 zeigt man, dass die Summe der alternierenden harmonischen Reihe $\ln 2$ ist:

$$\sum_{n=0}^{\infty} (-1)^n \frac{1}{n+1} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \dots = \ln 2.$$

3. Absolute Konvergenz

Eine Reihe $\sum_{n=0}^{\infty} a_n$ heißt **absolut konvergent**, falls die Reihe $\sum_{n=0}^{\infty} |a_n|$ konvergiert. Die alternierende harmonische Reihe ist zwar konvergent, aber nicht absolut konvergent, weil die harmonische Reihe nicht konvergiert.

BEMERKUNG 65. Die Summe einer absolut konvergenten Reihe ändert sich nicht, wenn man die Summanden umordnet, z.B.

$$\begin{aligned} \frac{2}{3} &= \frac{1}{1 - (-\frac{1}{2})} = \sum_{n=0}^{\infty} \frac{(-1)^n}{2^n} = \sum_{m=0}^{\infty} \frac{1}{2^{2m}} + \sum_{m=0}^{\infty} \frac{-1}{2^{2m+1}} = \sum_{m=0}^{\infty} \frac{1}{4^m} - \frac{1}{2} \sum_{m=0}^{\infty} \frac{1}{4^m} \\ &= \frac{1}{2} \sum_{m=0}^{\infty} \frac{1}{4^m} = \frac{1}{2} \cdot \frac{1}{1 - \frac{1}{4}} = \frac{2}{3}. \quad \square \end{aligned}$$

SATZ 83. Wenn $|a_n| \leq b_n$ für alle $n \geq n_0$ und $\sum_{n=n_0}^{\infty} b_n$ konvergiert, so konvergiert die Reihe $\sum_{n=0}^{\infty} a_n$ absolut und es gilt

$$\left| \sum_{n=0}^{\infty} a_n \right| \leq \left| \sum_{n=0}^{n_0-1} a_n \right| + \sum_{n=n_0}^{\infty} b_n.$$

BEISPIEL 142. Für $a_n = \frac{n+2}{3^n}$ und $b_n = \frac{2}{3}$ gilt $a_n = |a_n| \leq b_n$ für alle $n \geq 2$, denn für alle $n \geq 2$ gilt $n+2 \leq 2^n$. Daraus folgt

$$\sum_{n=0}^{\infty} \frac{n+2}{3^n} = \frac{2}{1} + \frac{3}{3} + \sum_{n=2}^{\infty} \frac{n+2}{3^n} \leq 2 + 1 + \sum_{n=2}^{\infty} \frac{2^n}{3^n} = 3 + \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^n = 3 + \sum_{n=2}^{\infty} b_n.$$

Die Reihe $\sum_{n=2}^{\infty} b_n$ konvergiert, weil die geometrische Reihe $\sum_{n=0}^{\infty} q^n$ für $|q| < 1$, also auch für $q = 2/3$, konvergiert. Es folgt (siehe Beispiel 140)

$$0 < \sum_{n=0}^{\infty} \frac{n+2}{3^n} \leq 3 + \sum_{n=2}^{\infty} \left(\frac{2}{3}\right)^n = 3 + \frac{4}{3} = \frac{13}{3}.$$

Mit Hilfe der Resultate aus Abschnitt 4 kann man die Summe der Reihe $\sum_{n=0}^{\infty} a_n$ sogar exakt berechnen.

FOLGERUNG 34. Wenn

$$\lim_{n \rightarrow \infty} \sqrt[n]{|a_n|} < 1 \text{ oder } \lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| < 1,$$

dann ist die Reihe $\sum_{n=0}^{\infty} a_n$ absolut konvergent.

BEWEIS. Es existieren $n_0 \in \mathbb{N}$ und q mit $|q| < 1$ und $|a_n| \leq q^n$ für alle $n \geq n_0$. $\square$

BEMERKUNG 66. Falls die Folgen $\{\sqrt[n]{|a_n|}\}_{n \in \mathbb{N}}$ bzw. $\{|\frac{a_{n+1}}{a_n}|\}_{n \in \mathbb{N}}$ zwar beschränkt sind, aber keinen Grenzwert besitzen, so betrachten man in Folgerung 34 stattdessen den größten Grenzwert einer konvergenten Teilfolge. Dieser heißt **Limes Superior** und wird mit $\limsup$ bezeichnet. $\square$

BEISPIEL 143. Für $a_n = \frac{n+2}{3^n}$ wie in Beispiel 142 gilt

$$\lim_{n \rightarrow \infty} \sqrt[n]{|a_n|} = \lim_{n \rightarrow \infty} \sqrt[n]{\left| \frac{n+2}{3^n} \right|} = \lim_{n \rightarrow \infty} \frac{\sqrt[n]{n+2}}{\sqrt[n]{3^n}} = \lim_{n \rightarrow \infty} \frac{\sqrt[n]{n+2}}{3} = \frac{1}{3} < 1$$

und

$$\lim_{n \rightarrow \infty} \left| \frac{a_{n+1}}{a_n} \right| = \lim_{n \rightarrow \infty} \left| \frac{\frac{n+1+2}{3^{n+1}}}{\frac{n+2}{3^n}} \right| = \lim_{n \rightarrow \infty} \left| \frac{(n+3)3^n}{3^{n+1}(n+2)} \right| = \frac{1}{3} \lim_{n \rightarrow \infty} \frac{n+3}{n+2} = \frac{1}{3} < 1.$$

SATZ 84 (Vergleich mit uneigentlichen Integralen). *Wenn eine Funktion f mit den folgenden drei Eigenschaften existiert*

- *f ist stetig und monoton fallend.*
- *Es existiert ein $n_0 \in \mathbb{N}$ mit $f(n) \geq |a_n|$ für alle $n \geq n_0$.*
- *Das uneigentliche Integral $\int_a^\infty f(x)dx$ für $a = n_0 - 1$ existiert.,*

dann konvergiert die Reihe $\sum_{n=0}^\infty a_n$ absolut.

BEWEIS. Für alle $n \geq n_0$ gilt $b_n := \int_{n-1}^n f(x)dx \geq |a_n|$. Wegen $\sum_{n=n_0}^\infty b_n = \int_a^\infty f(x)dx$ folgt die Behauptung aus Satz 83. $\square$

FOLGERUNG 35. *Die Reihe*

$$\sum_{n=1}^\infty \frac{1}{n^\alpha}$$

konvergiert genau dann, wenn $\alpha > 1$.

BEWEIS. Für $\alpha \leq 0$ bilden die Summanden $a_n := n^{-\alpha}$ keine Nullfolge.

Für $x > 0$ und $\alpha > 0$ ist die Funktion $f(x) = x^{-\alpha}$ streng monoton fallend. Weiterhin gilt $f(n) = n^{-\alpha} = a_n > 0$. Falls $\alpha > 1$, so konvergiert das uneigentliche Integral $\int_1^\infty x^{-\alpha} dx$ (siehe Lemma 24). Falls $0 < \alpha \leq 1$, so gilt $a_n \geq \frac{1}{n}$. Das harmonische Reihe gegen ∞ divergiert, konvergieren auch die Reihen $\sum_{n=1}^\infty n^{-\alpha}$ für $\alpha \leq 1$ nicht. $\square$

BEMERKUNG 67. Für die Reihen $\sum_{n=1}^\infty \frac{1}{n^\alpha}$ versagt das Konvergenzkriterium aus Folgerung 34, denn

$$\lim_{n \rightarrow \infty} \sqrt[n]{n^{-\alpha}} = \left(\lim_{n \rightarrow \infty} \sqrt[n]{n} \right)^{-\alpha} = 1 \text{ und } \lim_{n \rightarrow \infty} \frac{(n+1)^{-\alpha}}{n^{-\alpha}} = \left(\lim_{n \rightarrow \infty} \frac{n+1}{n} \right)^{-\alpha} = 1.$$

4. Potenzreihen

Eine reelle **Potenzreihe** um $x_0 \in \mathbb{R}$ ist eine Reihe der Form

$$(104) \quad \sum_{n=0}^\infty a_n (x - x_0)^n$$

mit Koeffizienten $a_n \in \mathbb{R}$ und einer Variablen $x \in \mathbb{R}$.

Läßt man den Grad eines Taylorpolynoms einer Funktion f um x_0 beliebig groß werden, so erhält man **Taylorreihen**

$$\sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x_0)(x-x_0)^n, \text{ also } a_n = \frac{1}{n!} f^{(n)}(x_0).$$

4.1. Konvergenzradius einer Potenzreihe.

SATZ 85. Es seien $a_n, x_0 \in \mathbb{R}$ und

$$(105) \quad R := \frac{1}{\limsup_{n \rightarrow \infty} \sqrt[n]{|a_n|}} = \limsup_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|.$$

Die Potenzreihe $\sum_{n=0}^{\infty} a_n(x-x_0)^n$ konvergiert für alle $x \in \mathbb{R}$ mit $|x-x_0| < R$ absolut.

Die Potenzreihe $\sum_{n=0}^{\infty} a_n(x-x_0)^n$ divergiert für alle $x \in \mathbb{R}$ mit $|x-x_0| > R$.

BEISPIEL 144. Die Taylorreihe der Exponentialfunktion e^x um $x_0 = 0$ ist

$$\sum_{n=0}^{\infty} \frac{1}{n!} x^n, \text{ also } a_n = \frac{1}{n!} \text{ und } \left| \frac{a_n}{a_{n+1}} \right| = \frac{(n+1)!}{n!} = n+1 \xrightarrow{n \rightarrow \infty} \infty = R.$$

Die Taylorreihe der Exponentialfunktion konvergiert für alle $x \in \mathbb{R}$.

BEISPIEL 145. Die Taylorreihe der Funktion $\sin x$ um $x_0 = 0$ ist

$$\sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}, \text{ also } a_n = \frac{(-1)^n}{(2n+1)!}.$$

Da für jedes $S \in \mathbb{R}$, $S > 1$, die Ungleichung $S^n < S^{2n+1} < (2n+1)!$ für fast alle $n \in \mathbb{N}$ gilt, ist die Folge $\sqrt[n]{(2n+1)!}$ unbeschränkt, also

$$R = \frac{1}{\limsup_{n \rightarrow \infty} \sqrt[n]{|a_n|}} = \frac{1}{\limsup_{n \rightarrow \infty} \sqrt[n]{\frac{1}{(2n+1)!}}} = \limsup_{n \rightarrow \infty} \sqrt[n]{(2n+1)!} = \infty.$$

Die Taylorreihe der Funktion $\sin x$ um $x_0 = 0$ konvergiert für alle $x \in \mathbb{R}$.

BEISPIEL 146. Die Taylorreihe der Funktion $\ln x$ um $x_0 > 0$ ist

$$\ln x_0 + \sum_{n=1}^{\infty} \frac{(-1)^n}{n x_0^n} (x-x_0)^n, \text{ also } a_n = \frac{(-1)^{n-1}}{n x_0^n}$$

und

$$\left| \frac{a_n}{a_{n+1}} \right| = \left| \frac{(-1)^{n-1}(n+1)x_0^{n+1}}{n x_0^n (-1)^n} \right| = x_0 \frac{n+1}{n} \xrightarrow{n \rightarrow \infty} x_0 = R.$$

Die Taylorreihe der Funktion $\ln x$ um $x_0 > 0$ konvergiert für alle $x \in \mathbb{R}$ mit $|x-x_0| < x_0$, also $0 < x < 2x_0$. Für $x = 0$ erhält man die divergente harmonische Reihe

$$\ln x_0 + \sum_{n=1}^{\infty} \frac{(-1)^n}{n x_0^n} (-x_0)^n = \ln x_0 - \sum_{n=1}^{\infty} \frac{1}{n}.$$

Für $x = 2x_0$ erhält man die konvergente alternierende harmonische Reihe

$$\ln x_0 + \sum_{n=1}^{\infty} \frac{(-1)^n}{nx_0^n} (2x_0 - x_0)^n = \ln x_0 + \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{n}.$$

BEISPIEL 147. Die geometrische Reihe $\sum_{n=0}^{\infty} x^n$ hat den Konvergenzradius

$$R = \lim_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right| = 1 \text{ und es gilt } \sum_{n=0}^{\infty} x^n = \frac{1}{1-x} \text{ für } |x| < 1.$$

4.2. Potenzreihenentwicklungen von Funktionen.

SATZ 86. Für alle $x \in \mathbb{R}$ gilt

$$(106) \quad e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n.$$

BEWEIS. Die Taylorreihe konvergiert gegen e^x , denn für ein festes $x \in \mathbb{R}$ gilt

$$\left| e^x - \sum_{n=0}^N \frac{1}{n!} x^n \right| \leq \frac{|x|^{N+1}}{(n+1)!} e^{|x|} \xrightarrow{N \rightarrow \infty} 0.$$

□

SATZ 87. Für alle $x \in \mathbb{R}$ gilt

$$(107) \quad \sin x = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}.$$

BEWEIS. Die Taylorreihe konvergiert gegen $\sin x$, denn für ein festes $x \in \mathbb{R}$ gilt

$$\left| \sin x - \sum_{n=0}^N \frac{(-1)^n}{(2n+1)!} x^{2n+1} \right| \leq \frac{|x|^{2N+2}}{(2N+2)!} \xrightarrow{N \rightarrow \infty} 0.$$

□

SATZ 88. Für jedes reelle Polynom p mit $\deg p = N$ und jedes $x_0 \in \mathbb{R}$ gilt

$$p(x) = \sum_{n=0}^N \frac{1}{n!} p^{(n)}(x_0) (x - x_0)^n \text{ für alle } x \in \mathbb{R}.$$

BEWEIS. Da $p^{(n+1)} \equiv 0$, gilt $|p(x) - T_{p, x_0, N}(x)| \leq \frac{|x - x_0|^{N+1}}{(N+1)!} \cdot 0 = 0$. □

Eine Funktion f heißt um x_0 **in eine Potenzreihe entwickelbar**, wenn

$$(108) \quad f(x) = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x_0) (x - x_0)^n \text{ für } |x - x_0| < R$$

für ein $R > 0$.

Polynome, die Exponentialfunktion und $\sin x$ sind für alle $x_0 \in \mathbb{R}$ in eine Potenzreihe entwickelbar. Die Summenformel für die geometrische Reihe $\sum_{n=0}^{\infty} x^n = \frac{1}{1-x}$ liefert eine Potenzreihenentwicklung der Funktion $\frac{1}{1-x}$ ist um $x_0 = 0$ mit Konvergenzradius 1.

SATZ 89. Für $x_0 \neq 0$ gilt

$$(109) \quad \frac{1}{x} = \sum_{n=0}^{\infty} \frac{(-1)^n}{x_0^{n+1}} (x - x_0)^n \text{ für } |x - x_0| < |x_0|.$$

BEWEIS.

$$\begin{aligned} \frac{1}{x} &= \frac{1}{x - x_0 + x_0} = \frac{1}{x_0} \cdot \frac{1}{1 + \frac{x-x_0}{x_0}} = \frac{1}{x_0} \cdot \frac{1}{1 - (-1)\frac{x-x_0}{x_0}} = \frac{1}{x_0} \sum_{n=0}^{\infty} \left(-\frac{x-x_0}{x_0} \right)^n \\ &= \frac{1}{x_0} \sum_{n=0}^{\infty} \frac{(-1)^n}{x_0^n} (x - x_0)^n = \sum_{n=0}^{\infty} \frac{(-1)^n}{x_0^{n+1}} (x - x_0)^n \end{aligned}$$

für $|\frac{x-x_0}{x_0}| < 1$, also $|x - x_0| < |x_0|$. $\square$

SATZ 90 (l'Hospitalsche Regel). Wenn f und g um x_0 in eine Potenzreihe entwickelbar sind, $f(x) = \sum_{n=k}^{\infty} a_n (x - x_0)^n$, $g(x) = \sum_{n=k}^{\infty} b_n (x - x_0)^n$ und $b_k \neq 0$, dann gilt

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} = \frac{a_k}{b_k}.$$

BEISPIEL 148. Für $f(x) = \sin x - x$, $g(x) = x^3$ und $x_0 = 0$ gilt $b_3 = 1$, $b_n = 0$ für $n \neq 3$ und

$$f(x) = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots - x = -\frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!},$$

also

$$\lim_{x \rightarrow 0} \frac{\sin x - x}{x^3} = \frac{-\frac{1}{3!}}{1} = -\frac{1}{3!} = -\frac{1}{6}.$$

BEMERKUNG 68. Die Funktion $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ mit $\varphi(x) = e^{-1/x^2}$ für $x \neq 0$ und $\varphi(0) = 0$ ist in $x_0 = 0$ beliebig oft differenzierbar. Es gilt $\varphi^{(n)}(0) = 0$ für alle $n \in \mathbb{N}$ (siehe Abschnitt 2.2). Die Taylorreihe der Funktion φ um $x_0 = 0$ ist also $\equiv 0$. Sie stimmt nur in $x_0 = 0$ mit der Funktion φ überein, obwohl ihr Konvergenzradius ∞ ist. Die Funktion φ ist um $x_0 = 0$ nicht in eine Potenzreihe entwickelbar. $\square$

4.3. Verknüpfungen von Potenzreihenentwicklungen.

SATZ 91. Wenn $f'(x) = \sum_{n=0}^{\infty} a_n (x - x_0)^n$ für $|x - x_0| < R$, dann

$$(110) \quad f(x) = f(x_0) + \sum_{n=0}^{\infty} \frac{a_n}{n+1} (x - x_0)^{n+1} = f(x_0) + \sum_{n=1}^{\infty} \frac{a_{n-1}}{n} (x - x_0)^n$$

für $|x - x_0| < R$.

BEISPIEL 149 (Potenzreihenentwicklung des Logarithmus). Für alle $x_0 > 0$ gilt

$$(111) \quad \ln x = \ln x_0 = \sum_{n=1}^{\infty} \frac{(-1)^{n-1}}{nx_0^n} (x - x_0)^n.$$

SATZ 92. Wenn $f(x) = \sum_{n=0}^{\infty} a_n(x - x_0)^n$ für $|x - x_0| < R$, dann

$$(112) \quad f'(x) = \sum_{n=0}^{\infty} na_n(x - x_0)^{n-1} = \sum_{n=1}^{\infty} na_n(x - x_0)^{n-1} = \sum_{n=0}^{\infty} (n+1)a_{n+1}(x - x_0)^n$$

für $|x - x_0| < R$.

BEISPIEL 150 (Potenzreihenentwicklung der Funktion $\cos x$). Aus der Potenzreihenentwicklung der Funktion $\sin x$ folgt durch Ableitung jedes Summanden

$$(113) \quad \cos = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} (2n+1)x^{2n} = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} x^{2n} \text{ für alle } x \in \mathbb{R}.$$

BEISPIEL 151. Für $|x| < 1$ gilt $(1-x)^{-1} = \sum_{n=0}^{\infty} x^n$. Daraus folgt

$$\left(\frac{1}{1-x} \right)' = \frac{1}{(1-x)^2} = \sum_{n=0}^{\infty} (n+1)x^n.$$

Für $x = 1/3$ erhält man

$$\begin{aligned} \sum_{n=0}^{\infty} \frac{n+2}{3^n} &= \sum_{n=0}^{\infty} \frac{n+1}{3^n} + \sum_{n=0}^{\infty} \frac{1}{3^n} = \sum_{n=0}^{\infty} (n+1) \left(\frac{1}{3} \right)^n + \sum_{n=0}^{\infty} \left(\frac{1}{3} \right)^n \\ &= \frac{1}{1 - \frac{1}{3}} + \frac{1}{(1 - \frac{1}{3})^2} = \frac{3}{2} + \frac{9}{4} = \frac{15}{4}. \end{aligned}$$

SATZ 93. Wenn g um x_0 und f um $g(x_0)$ in eine Potenzreihe entwickelbar sind, dann ist auch die Verkettung $f \circ g$ um x_0 in eine Potenzreihe entwickelbar.

BEWEIS. Die Funktion $f \circ g$ ist beliebig oft differenzierbar. Es gilt die Abschätzung

$$|f(g(x)) - T_{f,g(x_0),N}(x)| \leq \frac{|g(x) - g(x_0)|^{N+1}}{(N+1)!} \max_{|\eta - g(x_0)| \leq |g(x) - g(x_0)|} |f^{N+1}(\eta)|.$$

Es seien R_f der Konvergenzradius der Funktion f und R_g der Konvergenzradius der Funktion g . Die Konvergenz der Taylorreihe gegen den Funktionswert $f(g(x))$ für alle x mit $|x - x_0| < R_g$ und $|g(x) - g(x_0)| < R_f$ folgt aus der gleichmäßigen Konvergenz der Potenzreihen auf allen abgeschlossenen Intervallen innerhalb des offenen Konvergenzintervalls. $\square$

BEMERKUNG 69. Die Potenzreihenentwicklung von $f \circ g$ kann man bestimmen indem man die Ableitungen $(f \circ g)^{(n)}(x_0)$ berechnet oder die Potenzreihe von g in die von f einsetzt. $\square$

BEISPIEL 152. Wenn $a \in \mathbb{R}$ fest und $f(x) = \sum_{n=0}^{\infty} a_n x^n$ für $|x| < R$, dann gilt $f(ax) = \sum_{n=0}^{\infty} a_n (ax)^n = \sum_{n=0}^{\infty} a_n a^n x^n$ für alle $|x| < R/|a|$, zum Beispiel

$$e^{ax} = \sum_{n=0}^{\infty} \frac{a^n}{n!} x^n \text{ für alle } a, x \in \mathbb{R}.$$

BEISPIEL 153. Für $g(x) = -x^2$ erhält man

$$e^{-x^2} = \sum_{n=0}^{\infty} \frac{1}{n!} (-x^2)^n = \sum_{n=0}^{\infty} \frac{(-1)^n}{n!} x^{2n}.$$

SATZ 94. Wenn f und g um x_0 in eine Potenzreihe entwickelbar sind, dann sind auch $f + g$ und fg um x_0 in eine Potenzreihe entwickelbar. Es gilt

$$(114) \quad \sum_{n=0}^{\infty} a_n (x - x_0)^n + \sum_{n=0}^{\infty} b_n (x - x_0)^n = \sum_{n=0}^{\infty} (a_n + b_n) (x - x_0)^n$$

und

$$(115) \quad \left(\sum_{n=0}^{\infty} a_n (x - x_0)^n \right) \left(\sum_{n=0}^{\infty} b_n (x - x_0)^n \right) = \sum_{n=0}^{\infty} \left(\sum_{k=0}^n a_k b_{n-k} \right) (x - x_0)^n.$$

Wenn man die Exponentialfunktion nicht als Lösung der Differentialgleichung $y' = y$ (siehe Abschnitt 3), sondern als Potenzreihe $e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n$ definiert, so kann man die Potenzgesetze mit Hilfe des Produktsatzes beweisen.

FOLGERUNG 36. Für alle $a, b \in \mathbb{R}$ gilt $e^a e^b = e^{a+b}$.

BEWEIS. Für $f(x) = e^{ax}$ und $g(x) = e^{bx}$ mit $a, b \in \mathbb{R}$ gilt

$$f(x) = \sum_{n=0}^{\infty} \frac{a^n}{n!} x^n \text{ und } g(x) = \sum_{n=0}^{\infty} \frac{b^n}{n!} x^n$$

und

$$\begin{aligned} f(x)g(x) &= \left(f(x) = \sum_{n=0}^{\infty} \frac{a^n}{n!} x^n \right) \left(f(x) = \sum_{n=0}^{\infty} \frac{b^n}{n!} x^n \right) \\ &= \sum_{n=0}^{\infty} \left(\sum_{k=0}^n \frac{a^k}{k!} \frac{b^{n-k}}{(n-k)!} \right) x^n = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{k=0}^n \frac{n!}{k!(n-k)!} a^k b^{n-k} \right) x^n \\ &= \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{k=0}^n \binom{n}{k} a^k b^{n-k} \right) x^n = \sum_{n=0}^{\infty} \frac{1}{n!} (a+b)^n x^n = \sum_{n=0}^{\infty} \frac{(a+b)^n}{n!} x^n. \end{aligned}$$

Für $x = 1$ folgt die Behauptung aus Beispiel 152. $\square$

5. Die Hyperbelfunktionen

Wie jede Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ lässt sich auch die Exponentialfunktion als Summe einer geraden und einer ungeraden Funktion schreiben. Der gerade Anteil ist $\frac{1}{2}(f(x) + f(-x))$, der ungerade Anteil ist $\frac{1}{2}(f(x) - f(-x))$. Der ungerade Anteil der Funktion e^x heißt **Sinus hyperbolicus** und wird mit $\sinh$ bezeichnet. Der gerade Anteil der Funktion e^x heißt **Kosinus hyperbolicus** und wird mit $\cosh$ bezeichnet. Es gilt

$$(116) \quad \sinh x := \frac{1}{2}(e^x - e^{-x}) \text{ und } \cosh x := \frac{1}{2}(e^x + e^{-x}).$$

Abbildung 1 zeigt die Graphen der Exponentialfunktion und der Hyperbelfunktionen.

5.1. Potenzreihenentwicklung der Hyperbelfunktionen. Da die Potenzreihenentwicklung von e^x für alle $x \in \mathbb{R}$ die Funktion e^x darstellt und die Reihe für alle $x \in \mathbb{R}$ absolut konvergiert, kann man die Summanden der Reihe nach geraden und ungeraden Potenzen von x umordnen und erhält die Potenzreihenentwicklungen von $\sinh x$ und $\cosh x$. Es gilt

$$e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n = \sum_{m=0}^{\infty} \frac{1}{(2m)!} x^{2m} + \sum_{m=0}^{\infty} \frac{1}{(2m+1)!} x^{2m+1}$$

und damit

$$(117) \quad \sinh x = \sum_{m=0}^{\infty} \frac{1}{(2m+1)!} x^{2m+1} \text{ und } \cosh x = \sum_{m=0}^{\infty} \frac{1}{(2m)!} x^{2m} \text{ für alle } x \in \mathbb{R}.$$

5.2. Differentialgleichungen der Hyperbelfunktionen. Es gilt

$$\begin{aligned} (\sinh x)' &= \frac{1}{2}(e^x - e^{-x})' = \frac{1}{2}(e^x + e^{-x}) = \cosh x \\ (\cosh x)' &= \frac{1}{2}(e^x + e^{-x})' = \frac{1}{2}(e^x - e^{-x}) = \sinh x. \end{aligned}$$

Dieses Resultat erhält man auch, wenn man die Potenzreihenentwicklungen summandenweise ableitet.

$$\begin{aligned} (\sinh x)' &= \sum_{m=0}^{\infty} \frac{1}{(2m+1)!} (x^{2m+1})' = \sum_{m=0}^{\infty} \frac{2m+1}{(2m+1)!} x^{2m} = \sum_{m=0}^{\infty} \frac{1}{(2m)!} x^{2m} = \cosh x \\ (\cosh x)' &= \sum_{m=0}^{\infty} \frac{1}{(2m)!} (x^{2m})' = \sum_{m=0}^{\infty} \frac{2m}{(2m)!} x^{2m-1} = \sum_{m=1}^{\infty} \frac{1}{(2m-1)!} x^{2m-1} \\ &= \sum_{m=0}^{\infty} \frac{1}{(2m+1)!} x^{2m+1} = \sinh x \end{aligned}$$

Die Hyperbelfunktionen sind Lösungen der Differentialgleichung $y'' - y = 0$ zu den Anfangsbedingungen $y(0) = 0$ und $y'(0) = 1$ bzw. $y(0) = 1$ und $y'(0) = 0$, da $\sinh 0 = 0$ und $\cosh 0 = e^0 = 1$.

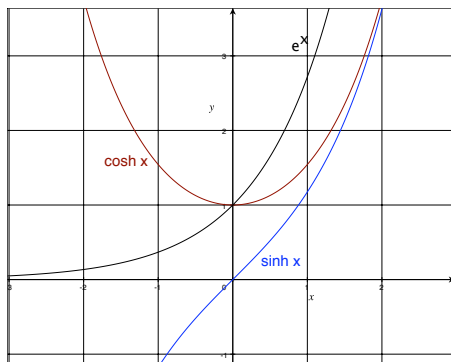


ABBILDUNG 1. Hyperbelfunktionen

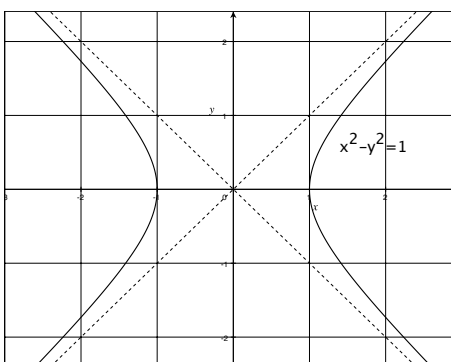


ABBILDUNG 2. Hyperbelgleichung

5.3. Identitäten für Hyperbelfunktionen.

Für alle $x \in \mathbb{R}$ gilt

$$(118) \quad \sinh^2 x + 1 = \cosh^2 x, \quad \sinh(2x) = 2 \sinh(x) \cosh(x),$$

$$\cosh(2x) = \cosh^2(x) + \sinh^2(x)$$

denn

$$\begin{aligned} \cosh^2 x &= \frac{1}{4}(e^x + e^{-x})^2 = \frac{1}{4}(e^{2x} + 2 + e^{-2x}) = 1 + \frac{1}{4}(e^{2x} - 2 + e^{-2x}) \\ &= 1 + \frac{1}{4}(e^x - e^{-x})^2 = 1 + \sinh^2 x \\ \sinh(2x) &= \frac{1}{2}(e^{2x} - e^{-2x}) = 2 \frac{1}{2}(e^x - e^{-x}) \frac{1}{2}e^x + e^{-x} \\ \cosh(2x) &= \frac{1}{2}(e^{2x} + e^{-2x}) = \frac{1}{4}(e^{2x} + 2 + e^{-2x}) + \frac{1}{4}(e^{2x} - 2 + e^{-2x}) \end{aligned}$$

Die Kurven $\gamma : \mathbb{R} \rightarrow \mathbb{R}^2$ mit

$$\gamma(t) = \begin{pmatrix} \cosh t \\ \sinh t \end{pmatrix} \text{ bzw. } \gamma(t) = \begin{pmatrix} -\cosh t \\ \sinh t \end{pmatrix}$$

parametrisieren die beiden Hyperbelarme, die Menge aller Punkte $(x, y) \in \mathbb{R}^2$ mit $x^2 - y^2 = 1$ (siehe Abbildung 2).

6. Komplexe Potenzreihen

Eine komplexe **Potenzreihe** um $z_0 \in \mathbb{C}$ ist eine Reihe der Form

$$(119) \quad \sum_{n=0}^{\infty} a_n (z - z_0)^n$$

mit Koeffizienten $a_n \in \mathbb{C}$ und einer Variablen $z \in \mathbb{C}$.

SATZ 95. Es seien $a_n, x_0 \in \mathbb{C}$ und

$$(120) \quad R := \frac{1}{\limsup_{n \rightarrow \infty} \sqrt[n]{|a_n|}} = \limsup_{n \rightarrow \infty} \left| \frac{a_n}{a_{n+1}} \right|.$$

Die Potenzreihe $\sum_{n=0}^{\infty} a_n(z-z_0)^n$ konvergiert für alle $z \in \mathbb{C}$ mit $|z-z_0| < R$ absolut.

Die Potenzreihe $\sum_{n=0}^{\infty} a_n(z-z_0)^n$ divergiert für alle $z \in \mathbb{C}$ mit $|z-z_0| > R$.

BEMERKUNG 70. Komplexe Potenzreihen konvergieren also für alle komplexen Zahlen z innerhalb des Kreises mit Radius R um z_0 . $\square$

Erlaubt man in einer Potenzreihenentwicklung $f(x) = \sum_{n=0}^{\infty} a_n(x-x_0)^n$ für $|x-x_0| < R$ auch komplexe Zahlen als Argumente, so erhält man eine natürliche Fortsetzung der Funktion f für komplexe Argumente, $f(z) := \sum_{n=0}^{\infty} a_n(z-z_0)^n$ für $|z-z_0| < R$. Zum Beispiel definiert man für alle $z \in \mathbb{C}$

$$(121) \quad e^z := \sum_{n=0}^{\infty} \frac{z^n}{n!},$$

$$(122) \quad \sin z := \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} z^{2n+1} \quad \text{und} \quad \cos z := \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n)!} z^{2n},$$

$$(123) \quad \sinh z := \sum_{n=0}^{\infty} \frac{1}{(2n+1)!} z^{2n+1} \quad \text{und} \quad \cosh z := \sum_{n=0}^{\infty} \frac{1}{(2n)!} z^{2n}.$$

Diese Definitionen sind mit der Definition $e^{it} = \cos t + i \sin t$ für $t \in \mathbb{R}$ verträglich, denn

$$\begin{aligned} e^{it} &= \sum_{n=0}^{\infty} \frac{(it)^n}{n!} = \sum_{n=0}^{\infty} \frac{i^n}{n!} t^n = \sum_{m=0}^{\infty} \frac{i^{2m}}{(2m)!} t^{2m} + \sum_{m=0}^{\infty} \frac{i^{2m+1}}{(2m+1)!} t^{2m+1} \\ &= \sum_{m=0}^{\infty} \frac{(-1)^m}{(2m)!} t^{2m} + i \sum_{m=0}^{\infty} \frac{(-1)^m}{(2m+1)!} t^{2m+1} = \cos t + i \sin t. \end{aligned}$$

Auch für alle $z \in \mathbb{C}$ gilt

$$(124) \quad \cos z = \frac{1}{2} (e^{iz} + e^{-iz}) \quad \text{und} \quad \sin z = \frac{1}{2i} (e^{iz} - e^{-iz}).$$

KAPITEL XI

Lineare Gleichungssysteme

Eine Gleichung der Form

$$(125) \quad a_1x_1 + a_2x_2 + \dots + a_nx_n = b$$

heißt **lineare Gleichung** in den n Variablen $x_1, x_2, \dots, x_n$. Die Faktoren $a_1, \dots, a_n$ vor den Variablen $x_1, \dots, x_n$ in der linearen Gleichung und auch die Konstante b auf der rechten Seite der Gleichung heißen **Koeffizienten**. Die lineare Gleichung 125 kann man auch in der Form $\sum_{j=1}^n a_jx_j = b$ schreiben.

Ein **lineares Gleichungssystem (LGS)** ist ein System linearer Gleichungen, es besteht also aus mehreren linearen Gleichungen. Das LGS

$$(126) \quad \sum_{j=1}^n a_{ij}x_j = b_i \quad \text{für alle } i = 1, \dots, m$$

besteht aus m linearen Gleichungen in den Variablen $x_1, \dots, x_n$.

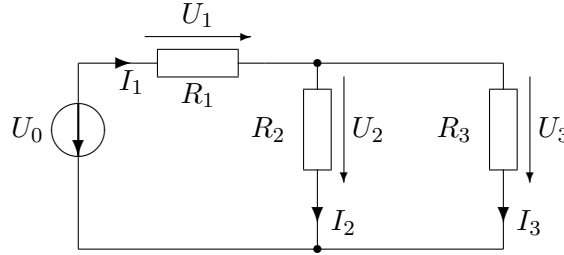
Eine **Lösung** eines LGSs ist ein n -Tupel $(\xi_1, \xi_2, \dots, \xi_n)$ mit $\sum_{j=1}^n a_{ij}\xi_j = b_i$ für alle $i = 1, \dots, m$.

Zu einem gegebenen LGS, also gegebenen Koeffizienten a_{ij} und b_i möchte man folgende Fragen beantworten:

- Besitzt das LGS eine Lösung? (Existenz einer Lösung)
- Besitzt das LGS mehrere Lösungen? (Eindeutigkeit der Lösung)
- Wie kann man die Menge aller Lösungen, die Lösungsmenge, beschreiben. Wie viele Parameter benötigt man, um die Lösungsmenge zu beschreiben? (Parametrisierung der Lösungsmenge)

Lineare Gleichungssysteme sind zum Beispiel zu lösen, wenn man die Ströme und Spannungen einer Schaltung mit Hilfe der Kirchhoffschen Gesetze (siehe Beispiel 154) oder die Konstanten einer Partialbruchzerlegung durch Koeffizientenvergleich zweier Polynome (siehe Beispiel 155) bestimmen möchte.

BEISPIEL 154 (LGS für die Stromstärken in einer Schaltung). In der folgenden Schaltung sind die Ohmschen Widerstände R_1, R_2, R_3 und die Spannung U_0 gegeben. Die Stromstärken I_1, I_2 und I_3 und die Spannungen U_1, U_2 und U_3 sind gesucht.



Aus dem Ohmschen Gesetz folgt $U_1 = R_1 I_1$, $U_2 = R_2 I_2$ und $U_3 = R_3 I_3$. Die Knotenregel liefert in beiden Knoten die Gleichung $I_1 - I_2 - I_3 = 0$. Die Schaltung besitzt zwei unabhängige Maschen. Sie liefern die Gleichungen $U_1 + U_2 - U_0 = 0$ und $U_2 - U_3 = 0$ für die gesuchten Spannungen. Beschreibt man die Spannungen U_j mit Hilfe der Stromstärken I_j , so erhält man das LGS

$$\begin{array}{rrcr} I_1 & -I_2 & -I_3 & = & 0 \\ & R_2 I_2 & -R_3 I_3 & = & 0 \\ R_1 I_1 & +R_2 I_2 & & = & U_0 \end{array}$$

für die Variablen I_1 , I_2 und I_3 . Setzt man für die gegebenen Größen R_1 , R_2 , R_3 und U_0 spezielle Werte ein, z.B. $R_1 = 1\Omega$, $R_2 = 2\Omega$, $R_3 = 3\Omega$ und $U_0 = 10V$, und kürzt die Einheiten, so erhält man das LGS

$$\begin{array}{rrcr} I_1 & -I_2 & -I_3 & = & 0 \\ & 2I_2 & -3I_3 & = & 0 \\ I_1 & +2I_2 & & = & 10 \end{array}$$

BEISPIEL 155 (LGS für die Konstanten einer Partialbruchzerlegung). Der Ansatz für die Partialbruchzerlegung der rationalen Funktion

$$\frac{2x^4 + x^3 + 5x^2 - 2x + 6}{(x-1)(x^2+1)^2}$$

ist

$$\frac{2x^4 + x^3 + 5x^2 - 2x + 6}{(x-1)(x^2+1)^2} = \frac{\alpha_1}{x-1} + \frac{\alpha_2 x + \alpha_3}{x^2+1} + \frac{\alpha_4 x + \alpha_5}{(x^2+1)^2}.$$

Um die Konstanten $\alpha_1, \dots, \alpha_5$ zu bestimmen, kann man den Ansatz mit dem Hauptnenner $(x-1)(x^2+1)^2$ multiplizieren und die Koeffizienten der Polynome auf beiden Seiten der Gleichung vergleichen. Man erhält

$$\begin{aligned} 2x^4 + x^3 + 5x^2 - 2x + 6 &= \alpha_1(x^2+1)^2 + (\alpha_2 x + \alpha_3)(x-1)(x^2+1) \\ &\quad + (\alpha_4 x + \alpha_5)(x-1) \\ &= \alpha_1(x^4 + 2x^2 + 1) + (\alpha_2 x + \alpha_3)(x^3 - x^2 + x - 1) \\ &\quad + (\alpha_4 x + \alpha_5)(x-1) \\ &= (\alpha_1 + \alpha_2)x^4 + (\alpha_3 - \alpha_2)x^3 + (2\alpha_1 + \alpha_2 - \alpha_3 + \alpha_4)x^2 \\ &\quad + (\alpha_3 - \alpha_2 + \alpha_5 - \alpha_4)x + (\alpha_1 - \alpha_3 - \alpha_5) \end{aligned}$$

und durch Vergleich der Koeffizienten vor x^4 , x^3 , x^2 , x und x^0 das aus fünf linearen Gleichungen bestehende LGS

$$\begin{array}{rrrrr} \alpha_1 & +\alpha_2 & & & = & 2 \\ & -\alpha_2 & +\alpha_3 & & = & 1 \\ 2\alpha_1 & +\alpha_2 & -\alpha_3 & +\alpha_4 & = & 5 \\ & -\alpha_2 & +\alpha_3 & -\alpha_4 & +\alpha_5 & = -2 \\ \alpha_1 & & -\alpha_3 & & -\alpha_5 & = 6 \end{array}$$

für die fünf Variablen $\alpha_1, \dots, \alpha_5$.

1. Die Einsetzmethode

Stellt man in einem LGS

$$\sum_{j=1}^n a_{ij}x_j = b_i \text{ für alle } i = 1, \dots, m$$

eine der m Gleichungen nach einer der Variablen $x_1, \dots, x_n$ um, zum Beispiel

$$x_n = \frac{1}{a_{mn}} (b_m - a_{m1}x_1 - a_{m2}x_2 - \dots - a_{mn-1}x_{n-1}) = \frac{b_m}{a_{mn}} - \sum_{j=1}^{n-1} \frac{a_{mj}}{a_{mn}} x_j,$$

falls $a_{mn} \neq 0$, und drückt x_n in den restlichen $m-1$ Gleichungen durch die Variablen $x_1, \dots, x_{n-1}$ aus, so erhält man ein aus $m-1$ linearen Gleichungen bestehendes LGS in $n-1$ Variablen $x_1, \dots, x_{n-1}$:

$$\sum_{j=1}^n \left(a_{ij} - \frac{a_{mj}}{a_{mn}} \right) x_j = b_i - \frac{b_m}{a_{mn}} \text{ für alle } i = 1, \dots, m-1$$

Die Auswertung einer Gleichung kann also die Anzahl der Variablen um 1 verringern. Im Idealfall kann man durch mehrfache Anwendung dieser Idee eine Lösung des Gleichungssystems.

BEISPIEL 156. Das LGS aus Beispiel 154 lautet:

$$\begin{array}{rrrr} x_1 & -x_2 & -x_3 & = & 0 \\ & 2x_2 & -3x_3 & = & 0 \\ x_1 & +2x_2 & & = & 10 \end{array}$$

Aus der dritten Gleichung erhält man $x_1 = 10 - 2x_2$. Setzt man dies in die erste Gleichung ein, so erhält man $10 - 2x_2 - x_2 - x_3 = 0$, also $x_3 = 10 - 3x_2$. Setzt man dies in die zweite Gleichung ein, so erhält man $2x_2 - 3(10 - 3x_2) = 0$, also $x_2 = 30/11$, $x_3 = 10 - 90/11 = 20/11$ und $x_1 = 10 - 60/11 = 50/11$. Das LGS hat genau eine Lösung:

$$(x_1, x_2, x_3) = \left(\frac{50}{11}, \frac{30}{11}, \frac{20}{11} \right)$$

Diese Lösung entsprechen in Beispiel 154 die Stromstärken $I_1 = \frac{50}{11} A$, $I_2 = \frac{30}{11} A$ und $I_3 = \frac{20}{11} A$.

BEISPIEL 157. Das LGS aus Beispiel 155 lautet:

$$\begin{array}{rrrrrr}
 \alpha_1 & +\alpha_2 & & & & = & 2 \\
 & -\alpha_2 & +\alpha_3 & & & = & 1 \\
 2\alpha_1 & +\alpha_2 & -\alpha_3 & +\alpha_4 & & = & 5 \\
 & -\alpha_2 & +\alpha_3 & -\alpha_4 & +\alpha_5 & = & -2 \\
 \alpha_1 & & -\alpha_3 & & -\alpha_5 & = & 6
 \end{array}$$

Aus der ersten Gleichung erhält man $\alpha_1 = 2 - \alpha_2$. Aus der zweiten Gleichung erhält man $\alpha_3 = 1 + \alpha_2$. Setzt man dies in die dritte Gleichung ein, so erhält man $2(2 - \alpha_2) + \alpha_2 - (1 + \alpha_2) + \alpha_4 = 5$, also $\alpha_4 = 2 + 2\alpha_2$. Setzt man dies in die vierte Gleichung ein, so erhält man $-\alpha_2 + (1 + \alpha_2) - (2 + 2\alpha_2) + \alpha_5 = -2$, also $\alpha_5 = 2\alpha_2 - 1$. Setzt man dies in die fünfte Gleichung ein, so erhält man $(2 - \alpha_2) - (1 + \alpha_2) - (2\alpha_2 - 1) = 6$, also $\alpha_2 = -1$. Aus α_2 kann man jetzt alle α_j berechnen. Es folgt $\alpha_1 = 2 - (-1) = 3$, $\alpha_2 = -1$, $\alpha_3 = 1 + (-1) = 0$, $\alpha_4 = 2 + 2(-1) = 0$ und $\alpha_5 = 2(-1) - 1 = -3$. Dies ist die eindeutige Lösung des LGSs, da jede Gleichung genau ein Mal verarbeitet wurde. Dieser Lösung entspricht in Beispiel 155 die Partialbruchzerlegung

$$\frac{2x^4 + x^3 + 5x^2 - 2x + 6}{(x-1)(x^2+1)^2} = \frac{3}{x-1} - \frac{x}{x^2+1} - \frac{3}{(x^2+1)^2}.$$

BEMERKUNG 71. Die Einsetzmethode bietet sehr viele Freiheiten. Man die Reihenfolge, in der die einzelnen Gleichungen ausgewertet werden, frei wählen. In jedem Schritt kann man die gewählte Gleichung nach einer beliebigen Variable umstellen, falls der Koeffizient vor dieser Variable verschieden von 0 ist. Jedoch muss man darauf achten, dass jede Gleichung genau ein Mal ausgewertet wird. $\square$

Die Einsetzmethode ist nur für „kleine“ lineare Gleichungssysteme, d.h. wenige Variablen und wenige Gleichungen, geeignet, da das Umstellen der Gleichungen viel Schreibarbeit erfordert. Man kann mit Hilfe der Einsetzmethode erkennen, ob das LGS genau eine Lösung besitzt (siehe Beispiele 157 und 156), nicht lösbar ist (siehe Beispiel 158) oder mehrere Lösungen besitzt (siehe Beispiel 159). Falls das LGS mehrere Lösungen besitzt, so ist die Beschreibung der Lösungsmenge mit Hilfe der Einsetzmethode umständlicher als nach Durchführung des Gaußverfahrens, das in Abschnitt 2 beschrieben wird.

BEISPIEL 158 (Nicht lösbares LGS). Das LGS

$$\begin{array}{rrrr}
 x_1 & & +2x_3 & = & 1 \\
 x_1 & +2x_2 & +4x_3 & = & 4 \\
 & x_2 & +x_3 & = & 1
 \end{array}$$

besitzt keine Lösung, denn aus der ersten Gleichung folgt $x_1 = 1 - 2x_3$ und aus der dritten Gleichung folgt $x_2 = 1 - x_3$. Setzt man dies in die zweite Gleichung ein, so erhält man $(1 - 2x_3) + 2(1 - x_3) + 4x_3 = 4$, also $3 = 4$.

BEISPIEL 159 (LGS mit mehreren Lösungen). Das LGS

$$\begin{array}{rcl} x_1 & +2x_3 & = 1 \\ x_1 & +2x_2 & +4x_3 = 5 \\ & x_2 & +x_3 = 2 \end{array}$$

besitzt mehrere Lösungen, denn aus der ersten Gleichung folgt $x_1 = 1 - 2x_3$ und aus der dritten Gleichung folgt $x_2 = 2 - x_3$. Setzt man dies in die zweite Gleichung ein, so erhält man $(1 - 2x_3) + 2(2 - x_3) + 4x_3 = 4$, also $4 = 4$. Sind die erste und die dritte Gleichung erfüllt, so ist hier die zweite Gleichung auch erfüllt. Man kann $x_2 = t \in \mathbb{R}$ beliebige wählen und erhält für jede Wahl eine Lösung des linearen Gleichungssystems. Die Lösungsmenge ist $\{(1 - 2t, 2 - t, t) : t \in \mathbb{R}\}$.

2. Das Gaußverfahren

Ein lineares Gleichungssystem wird beim Gaußverfahren mit Hilfe elementarer Umformungen systematisch in eine Normalform, der Zeilenstufenform, transformiert, da die Lösungsmenge eines linearen Gleichungssystems in Zeilenstufenform einfach zu beschreiben ist.

Bei der Transformation des linearen Gleichungssystems in eine Zeilenstufenform mit Hilfe des Gaußverfahrens genügt es, die Veränderung der Koeffizienten zu beobachten. Diese können effizient in einer Matrix gespeichert werden.

2.1. Elementare Transformationen eines linearen Gleichungssystems. Die Lösungsmenge des linearen Gleichungssystems

$$\sum_{j=1}^n a_{ij}x_j = b_i \quad \forall i = 1, \dots, m$$

bleibt unverändert, wenn man eine der folgenden Umformungen durchführt:

- Vertauschung zweier Gleichungen
- Multiplikation einer Gleichung mit einer von Null verschiedenen Zahl
- Addition des Vielfachen einer Gleichung zu einer anderen Gleichung

Bei diesen drei elementaren Umformungen werden die Summanden innerhalb einer Gleichung nicht vertauscht. Der j -te Summand ist immer ein Vielfaches der Variablen x_j .

BEMERKUNG 72. Die Lösungsmenge eines linearen Gleichungssystems ändert sich auch nicht, wenn man die Variablen umbenennt oder die Summanden innerhalb einer Gleichung vertauscht. Jedoch müßte man diese Veränderungen buchführen, wenn man nur die erweiterte Koeffizientenmatrix in Stufenform transformiert. $\square$

2.2. Die erweiterte Koeffizientenmatrix. Die **Koeffizientenmatrix** A des linearen Gleichungssystems $\sum_{j=1}^n a_{ij}x_j = b_i$ für alle $i = 1, \dots, m$ ist

$$(127) \quad A := \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}.$$

Die Koeffizientenmatrix A besteht aus m Zeilen und n Spalten. Die **erweiterte Koeffizientenmatrix** $(A|b)$ des linearen Gleichungssystems $\sum_{j=1}^n a_{ij}x_j = b_i$ für alle $i = 1, \dots, m$ ist

$$(128) \quad (A|b) := \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ a_{21} & a_{22} & \cdots & a_{2n} & b_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{pmatrix}.$$

Die erweiterte Koeffizientenmatrix $(A|b)$ besteht aus m Zeilen und $n + 1$ Spalten. Um die Koeffizienten, die in den Gleichungen des linearen Gleichungssystems auf verschiedenen Seiten des Gleichheitszeichens stehen auch in der erweiterten Koeffizientenmatrix voneinander zu trennen, kann man auch die Schreibweise

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

benutzen.

Die Vertauschung zweier Gleichungen des linearen Gleichungssystems entspricht der Vertauschung zweier Zeilen in der erweiterten Koeffizientenmatrix:

$$\left(\begin{array}{cccc|c} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{k1} & a_{k2} & \cdots & a_{kn} & b_k \\ \vdots & \vdots & & \vdots & \vdots \\ a_{l1} & a_{l2} & \cdots & a_{ln} & b_l \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{array} \right) \xrightarrow{(l) \leftrightarrow (k)} \left(\begin{array}{cccc|c} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{l1} & a_{l2} & \cdots & a_{ln} & b_l \\ \vdots & \vdots & & \vdots & \vdots \\ a_{k1} & a_{k2} & \cdots & a_{kn} & b_k \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{array} \right)$$

Der Multiplikation der k -ten Gleichung des linearen Gleichungssystems mit $\lambda \neq 0$ entspricht die Multiplikation der k -ten Zeile der erweiterten Koeffizientenmatrix mit λ :

$$\begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{k1} & a_{k2} & \cdots & a_{kn} & b_k \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{pmatrix} \xrightarrow{(k) \rightarrow \lambda(k)} \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ \vdots & \vdots & & \vdots & \vdots \\ \lambda a_{k1} & \lambda a_{k2} & \cdots & \lambda a_{kn} & \lambda b_k \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{pmatrix}$$

Der Addition des λ -fachen der l -ten Gleichung zur k -ten Gleichung des linearen Gleichungssystems entspricht die Addition des λ -fachen der l -ten Zeile zur k -ten Zeile der erweiterten Koeffizientenmatrix:

$$\begin{pmatrix} a_{11} & \cdots & a_{1n} & b_1 \\ \vdots & & \vdots & \vdots \\ a_{k1} & \cdots & a_{kn} & b_k \\ \vdots & & \vdots & \vdots \\ a_{l1} & \cdots & a_{ln} & b_l \\ \vdots & & \vdots & \vdots \\ a_{m1} & \cdots & a_{mn} & b_m \end{pmatrix} \xrightarrow{(k) \rightarrow (k) + \lambda(l)} \begin{pmatrix} a_{11} & \cdots & a_{1n} & b_1 \\ \vdots & & \vdots & \vdots \\ a_{k1} + \lambda a_{l1} & \cdots & a_{kn} + \lambda a_{ln} & b_k + \lambda b_l \\ \vdots & & \vdots & \vdots \\ a_{l1} & \cdots & a_{ln} & b_l \\ \vdots & & \vdots & \vdots \\ a_{m1} & \cdots & a_{mn} & b_m \end{pmatrix}$$

2.3. Die Stufenform. Eine $m \times n$ -Matrix

$$A = (a_{ij})_{i=1, \dots, m, j=1, \dots, n} = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{m1} & \cdots & a_{mn} \end{pmatrix}$$

hat eine **Zeilenstufenform**, falls sie die folgende Eigenschaften hat:

$$a_{i^*j^*} \neq 0 \text{ und } a_{ij} = 0 \quad \forall j < j^* \Rightarrow a_{ij} = 0 \quad \forall i > i^*, j \leq j^*$$

Dies bedeutet, dass die Matrix A unterhalb einer stufenförmigen Grenzlinie nur noch die Einträge 0 hat. Die Höhe der Stufen ist immer eine Zeile. Jedoch kann eine Stufe auch mehr als eine Spalte breit sein:

$$\begin{pmatrix} * & * & * & \cdots & * & * & * \\ 0 & * & * & \cdots & * & * & * \\ 0 & 0 & * & & \vdots & \vdots & * \\ \vdots & \vdots & & & * & * & \vdots \\ 0 & 0 & 0 & \cdots & 0 & 0 & * \end{pmatrix}$$

SATZ 96 (Gaußverfahren). Jede Matrix kann durch wiederholte Anwendung der drei elementaren Transformationen

- Vertauschung zweier Zeilen
- Multiplikation einer Zeile mit einer von 0 verschiedenen Zahl
- Addition eines Vielfachen einer Zeile zu einer anderen Zeile

in eine Stufenform gebracht werden.

BEWEIS. Wenn $a_{ij} = 0$ für alle (i, j) mit $j < l$ und $i \geq k$ und $a_{kl} \neq 0$, dann erzeugt Addition des $-\frac{a_{il}}{a_{kl}}$ -fache der k -ten Zeile zur i -ten Zeile Nullen in der l -ten Spalte für alle Zeilen $i > k$. Dabei werden die Nullen in den Zeilen $j < k$ nicht zerstört. $\square$

BEISPIEL 160. Die erweiterte Koeffizientenmatrix des Beispiels 157 ist

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 2 & 1 & -1 & 1 & 0 & 5 \\ 0 & -1 & 1 & -1 & 1 & -2 \\ 1 & 0 & -1 & 0 & -1 & 6 \end{pmatrix}.$$

Wir addieren Vielfache der ersten Zeile zur dritten und zur fünften Zeile, $(3)-2(1)$, $(5)-(1)$, um Nullen in der ersten Spalte zu erzeugen:

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 0 & -1 & -1 & 1 & 0 & 1 \\ 0 & -1 & 1 & -1 & 1 & -2 \\ 0 & -1 & -1 & 0 & -1 & 4 \end{pmatrix}$$

Wir subtrahieren die zweite Zeile von der dritten, vierten und fünften Zeile, um Nullen in der zweiten Spalte zu erzeugen:

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 0 & 0 & -2 & 1 & 0 & 0 \\ 0 & 0 & 0 & -1 & 1 & -3 \\ 0 & 0 & -2 & 0 & -1 & 3 \end{pmatrix}$$

Wir subtrahieren die dritte Zeile von der fünften Zeile, um Nullen in der dritten Spalte zu erzeugen:

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 0 & 0 & -2 & 1 & 0 & 0 \\ 0 & 0 & 0 & -1 & 1 & -3 \\ 0 & 0 & 0 & -1 & -1 & 3 \end{pmatrix}$$

Wir subtrahieren die vierte Zeile von der fünften Zeile, um Nullen in der vierten Spalte zu erzeugen:

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 0 & 0 & -2 & 1 & 0 & 0 \\ 0 & 0 & 0 & -1 & 1 & -3 \\ 0 & 0 & 0 & 0 & -2 & 6 \end{pmatrix}$$

Dies ist eine Stufenform der Koeffizientenmatrix.

BEISPIEL 161. Die erweiterte Koeffizientenmatrix des Beispiels 156 ist

$$\begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 2 & -3 & 0 \\ 1 & 2 & 0 & 10 \end{pmatrix}.$$

Durch elementare Umformungen erhält man eine Stufenform:

$$\begin{aligned} \begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 2 & -3 & 0 \\ 1 & 2 & 0 & 10 \end{pmatrix} &\sim \begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 2 & -3 & 0 \\ 0 & 3 & 1 & 10 \end{pmatrix} \sim \begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 2 & -3 & 0 \\ 0 & 1 & 4 & 10 \end{pmatrix} \\ &\sim \begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 1 & 4 & 10 \\ 0 & 2 & -3 & 0 \end{pmatrix} \sim \begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 1 & 4 & 10 \\ 0 & 0 & -11 & -20 \end{pmatrix} \end{aligned}$$

BEISPIEL 162. Die erweiterte Koeffizientenmatrix des Beispiels 158 ist

$$\begin{pmatrix} 1 & 0 & 2 & 1 \\ 1 & 2 & 4 & 4 \\ 0 & 1 & 1 & 1 \end{pmatrix}.$$

Durch elementare Umformungen erhält man eine Stufenform:

$$\begin{aligned} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 1 & 2 & 4 & 4 \\ 0 & 1 & 1 & 1 \end{pmatrix} &\stackrel{(2) \rightarrow (2) - (1)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 2 & 2 & 3 \\ 0 & 1 & 1 & 1 \end{pmatrix} \stackrel{(2) \leftrightarrow (3)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 2 & 2 & 3 \end{pmatrix} \\ &\stackrel{(3) \rightarrow (3) - 2(2)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

BEISPIEL 163. Die erweiterte Koeffizientenmatrix des Beispiels 159 ist

$$\begin{pmatrix} 1 & 0 & 2 & 1 \\ 1 & 2 & 4 & 5 \\ 0 & 1 & 1 & 2 \end{pmatrix}.$$

Durch elementare Umformungen erhält man eine Stufenform:

$$\begin{aligned} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 1 & 2 & 4 & 5 \\ 0 & 1 & 1 & 2 \end{pmatrix} &\stackrel{(2) \rightarrow (2) - (1)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 2 & 2 & 4 \\ 0 & 1 & 1 & 2 \end{pmatrix} \stackrel{(2) \leftrightarrow (3)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 2 \\ 0 & 2 & 2 & 4 \end{pmatrix} \\ &\stackrel{(3) \rightarrow (3) - 2(2)}{\sim} \begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 2 \\ 0 & 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

2.4. Lösungsmenge eines linearen Gleichungssystems. Wird die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems mit Hilfe des Gaußverfahrens in eine Stufenform transformiert, so ändert sich die Lösungsmenge nicht. Wir nehmen nun an, dass die erweiterte Koeffizientenmatrix eines linearen Gleichungssystems Stufenform besitzt.

Wenn $a_{kj} = 0$ für alle $j = 1, \dots, n$ und $b_k \neq 0$ für einen Index k , dann besitzt das LGS keine Lösung, denn der k -ten Zeile entspricht die Gleichung $0 = b_k \neq 0$.

Wenn kein Index k mit $b_k \neq 0$ und $a_{kj} = 0$ für alle $j = 1, \dots, n$ existiert, dann besitzt das LGS eine Lösung. Für jede Zeile i sei j_i der kleinste Spaltenindex mit $a_{ij_i} \neq 0$. Die entsprechende Gleichung kann man nach x_{j_i} auflösen:

$$x_{j_i} = \frac{b_{j_i}}{a_{ij_i}} - \sum_{j=j_i+1}^n \frac{a_{ij}}{a_{ij_i}} x_j$$

Die Indexmenge

$$J := \{j_1\} \cup \{j_2\} \cup \dots \cup \{j_m\}$$

beinhaltet die Indizes aller abhängigen Variablen. Die Menge $\bar{J} := \{1, \dots, n\} \setminus J$ beinhaltet die Indizes aller freien Variablen. Zur Beschreibung der Lösungsmenge des linearen Gleichungssystems benötigt man $|\bar{J}|$ (reelle) Parameter, also für jedes Element in $\bar{J}$ einen Parameter.

BEISPIEL 164. Das LGS aus den Beispielen 155, 157 und 160 besitzt die Stufenform

$$\begin{pmatrix} 1 & 1 & 0 & 0 & 0 & 2 \\ 0 & -1 & 1 & 0 & 0 & 1 \\ 0 & 0 & -2 & 1 & 0 & 0 \\ 0 & 0 & 0 & -1 & 1 & -3 \\ 0 & 0 & 0 & 0 & -2 & 6 \end{pmatrix}.$$

Es gilt $J = \{1, 2, 3, 4, 5, 6\}$. Das LGS ist eindeutig lösbar.

Aus der letzten Zeile folgt $-2\alpha_5 = 6$, also $\alpha_5 = -3$.

Aus der vierten Zeile folgt $-\alpha_4 + \alpha_5 = -3$, also $\alpha_4 = 0$.

Aus der dritten Zeile folgt $-2\alpha_3 + \alpha_4 = 0$, also $\alpha_3 = 0$.

Aus der zweiten Zeile folgt $-\alpha_2 + \alpha_3 = 1$, also $\alpha_2 = -1$.

Aus der ersten Zeile folgt $\alpha_1 + \alpha_2 = 2$, also $\alpha_1 = 3$.

Diese Lösung stimmt mit der in Abschnitt 1 gefundenen überein.

BEISPIEL 165. Das LGS aus den Beispielen 154, 156 und 161 besitzt die Stufenform

$$\begin{pmatrix} 1 & -1 & -1 & 0 \\ 0 & 1 & 4 & 10 \\ 0 & 0 & -11 & -20 \end{pmatrix}.$$

Es gilt $J = \{1, 2, 3\}$. Das LGS ist eindeutig lösbar.

Aus der dritten Zeile folgt $-11I_3 = -20$, also $I_3 = 20/11$.

Aus der zweiten Zeile folgt $I_2 + 4I_3 = 10$, also $I_2 = 30/11$.

Aus der ersten Zeile folgt $I_1 - I_2 - I_3 = 0$, also $I_1 = 50/11$.

Diese Lösung stimmt mit der in Abschnitt 1 gefundenen überein.

BEISPIEL 166. Das LGS aus den Beispielen 159 und 163 besitzt die Stufenform

$$\begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Es gilt $J = \{1, 2\}$ und $\bar{J} = \{3\}$. Das LGS ist lösbar. Die Lösungsmenge kann man mit Hilfe der freien Variable x_3 beschreiben.

Aus der zweiten Zeile folgt $x_2 + x_3 = 1$, also $x_2 = 1 - x_3$.

Aus der ersten Zeile folgt $x_1 + 2x_3 = 1$, also $x_1 = 1 - 2x_3$.

Die Lösungsmenge des linearen Gleichungssystems ist $\{(1 - 2t, 1 - t, t) : t \in \mathbb{R}\}$.

BEISPIEL 167. Das LGS aus den Beispielen 158 und 162 besitzt die Stufenform

$$\begin{pmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

Da $a_{31} = a_{32} = a_{33} = 0$ und $b_3 = 1 \neq 0$, besitzt das LGS keine Lösung.

BEMERKUNG 73 (Aspekte der Implementierung). Das Gaußverfahren liefert einen Algorithmus zur Lösung linearer Gleichungssysteme. Für ein LGS aus m Gleichungen in n Variablen benötigt maximal $\frac{m}{2}(m+1)$ elementare Umformungen der erweiterten Koeffizientenmatrix, um eine Stufenform zu erreichen.

Sind die Anzahl der Variablen und die Anzahl der Gleichungen, also n und m , sehr groß und die Anzahl der von Null verschiedenen Koeffizienten sehr klein, so existieren für diese *schwach besetzte* erweiterte Koeffizientenmatrix eventuell effektivere Speicherformate als Matrizen und ein schnellerer Algorithmus als das Standardgaußverfahren.

Sind die Eingangsdaten, die Einträge der erweiterten Koeffizientenmatrix, ganze Zahlen, so erhält man bei der Transformation in Stufenform Matrizen, deren Einträge Brüche sind. Da nur eine begrenzte Bitzahl zur Speicherung ganzer Zahlen zur Verfügung steht, versucht man, Brüche so weit wie möglich zu kürzen. Hier könnte der Euklidische Algorithmus zur Berechnung des größten gemeinsamen Teilers Anwendung finden.

Ist man gezwungen, Dezimalbrüche mit einer festen Anzahl von Nachkommastellen zu benutzen, so können Rundungsfehler auftreten, die unter Umständen nicht nur die Genauigkeit der berechneten Lösungen beeinflussen, sondern sogar qualitative Aussagen wie die Anzahl der Lösungen verfälschen. Zum Beispiel gilt für die erweiterte Koeffizientenmatrix

$$\left(\begin{array}{cc|c} 1 & 0,00001 & 0 \\ 0,00001 & 0 & 0 \end{array} \right) \sim \left(\begin{array}{cc|c} 1 & 0,00001 & 0 \\ 0 & 0,0000000001 & 0 \end{array} \right) = \left(\begin{array}{cc|c} 1 & 0,00001 & 0 \\ 0 & 0 & 0 \end{array} \right)$$

bei der Beachtung von nur 8 Nachkommastellen. Obwohl das LGS eindeutig lösbar ist, liefert das (numerische) Gaußverfahren einen freien Parameter in der Lösungsmenge. Man muss Methoden zur Vermeidung bzw. Abschätzung

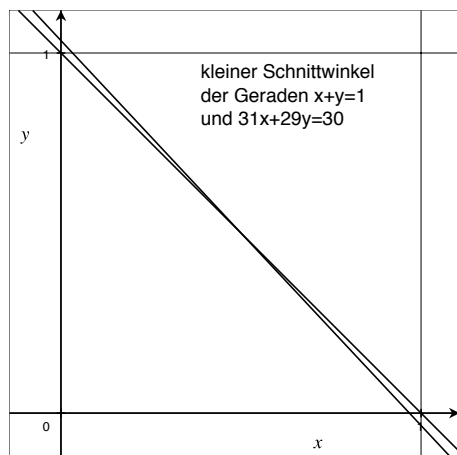


ABBILDUNG 1

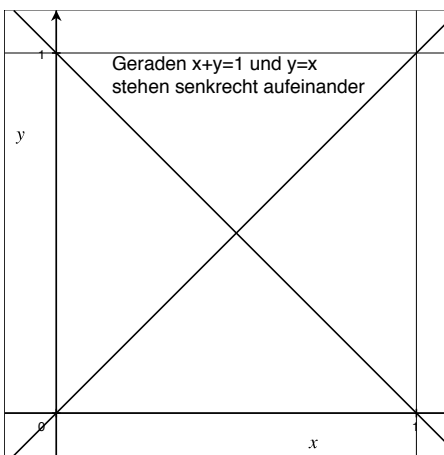


ABBILDUNG 2

von Rundungsfehlern entwickeln. Als Faustregel gilt: Stehen die Zeilen der Koeffizientenmatrix fast senkrecht aufeinander, so sind die Rundungsfehler klein (siehe Abbildung 2). Bilden Zeilen der Koeffizientenmatrix einen kleinen Winkel, so ist der Rundungsfehler groß und der Schnittpunkt der Geraden bzw. Hyperebenen in Skizzen unscharf (siehe Abbildung 1). $\square$

3. Rang einer Matrix

Der **Rang** einer Matrix M in Zeilenstufenform ist die Anzahl der Zeilen, die mindestens einen von Null verschiedenen Eintrag besitzen. Der Rang einer Matrix M ist der Rang einer Matrix in Stufenform, die man aus M durch elementare Transformationen (Gaußverfahren) bilden kann. Er wird mit $\text{rang}(M)$ bezeichnet. Wir werden in Abschnitt XII sehen, dass der Rang einer Matrix gerade die Anzahl der linear unabhängigen Zeilen der Matrix ist.

BEISPIEL 168. Die erweiterte Koeffizientenmatrix aus Beispiel 160 hat den Rang 5. Auch die dazugehörige Koeffizientenmatrix hat Rang 5.

BEISPIEL 169. Die erweiterte Koeffizientenmatrix aus Beispiel 161 hat den Rang 3. Auch die dazugehörige Koeffizientenmatrix hat Rang 3.

BEISPIEL 170. Die erweiterte Koeffizientenmatrix aus Beispiel 163 hat den Rang 2. Auch die dazugehörige Koeffizientenmatrix hat Rang 2.

BEISPIEL 171. Die erweiterte Koeffizientenmatrix aus Beispiel 162 hat den Rang 3. Die dazugehörige Koeffizientenmatrix hat nur Rang 2.

SATZ 97. *Das lineare Gleichungssystem*

$$\sum_{j=1}^n a_{ij}x_j = b_i \quad \forall i = 1, \dots, m$$

besitzt genau dann eine Lösung, falls der Rang der Koeffizientenmatrix A gleich dem Rang der erweiterten Koeffizientenmatrix $(A|b)$ ist, also $\text{rang}(A) = \text{rang}(A|b)$. In dem Fall benötigt man $n - \text{rang } A$ Parameter, um die Lösungsmenge zu beschreiben.

4. Interpolationspolynome

Sucht man zu einer gegebenen Funktion f ein Polynom, das die Funktion f möglichst gut approximiert, so muss man konkretisieren, welche Kriterien für die Güte der Näherung ausschlaggebend sind. Die optimale Lösung wird von diesen Kriterien abhängen.

4.1. Taylorpolynome. Betrachtet man nur einen Bezugspunkt $x_0 \in [a, b]$ und fordert, dass p im Punkt x_0 und in der Nähe des Punktes x_0 so gut wie möglich mit f übereinstimmt, so sucht man ein Polynom p vom Grad d mit der Eigenschaft $p(x_0) = f(x_0)$ und $p^{(n)}(x_0) = f^{(n)}(x_0)$ für $n = 1, \dots, d$. Diese Aufgabenstellung ist nur sinnvoll, wenn die Funktion f mindestens d mal differenzierbar ist. Die eindeutige Lösung dieses Problems ist das Taylorpolynom (siehe Abschnitt 2) von f um x_0 vom Grad d

$$p(x) = \sum_{n=0}^d \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n.$$

4.2. Newtonpolynome. Betrachtet man $m + 1$ Bezugspunkte x_0, x_1 bis x_m und fordert $p(x_i) = f(x_i) =: y_i$ für $i = 0, \dots, m$, so erhält man aus diesen $m + 1$ Bedingungen mit dem Ansatz

$$p(x) = \sum_{j=0}^m a_j x^j = a_0 + a_1 x + a_2 x^2 + \dots + a_m x^m$$

das lineare Gleichungssystem

$$\sum_{j=0}^m x_i^j a_j = y_i \quad \text{für } i = 0, \dots, m$$

für die Variablen $a_0, \dots, a_m$. Es besitzt die erweiterte Koeffizientenmatrix

$$(129) \quad (A|b) = \left(\begin{array}{cccc|c} 1 & x_0 & x_0^2 & \cdots & x_0^m & y_0 \\ 1 & x_1 & x_1^2 & \cdots & x_1^m & y_1 \\ \vdots & \vdots & \vdots & & \vdots & \vdots \\ 1 & x_m & x_m^2 & \cdots & x_m^m & y_m \end{array} \right).$$

BEMERKUNG 74. Die Koeffizientenmatrix A und auch ihre Determinante $\det A$ in Gleichung 129 heißen Vandermondsche in $x_0, \dots, x_m$. Es gilt $\text{rang } A = m + 1 \Leftrightarrow \det A \neq 0 \Leftrightarrow x_j \neq x_k \forall j \neq k$. $\square$

BEISPIEL 172. Für $f(x) = \sin x$, $x_0 = \pi/2$, $x_1 = \pi$ und $x_2 = 3\pi/2$ ist ein Polynom zweiten Grades $p(x) = a_0 + a_1x + a_2x^2$ gesucht und das lineare Gleichungssystem in Gleichung 129 wird zu

$$\begin{pmatrix} 1 & \frac{1}{2}\pi & \frac{1}{4}\pi^2 & | & 1 \\ 1 & \pi & \pi^2 & | & 0 \\ 1 & \frac{3}{2}\pi & \frac{9}{4}\pi^2 & | & -1 \end{pmatrix} \sim \begin{pmatrix} 1 & \frac{1}{2}\pi & \frac{1}{4}\pi^2 & | & 1 \\ 0 & \frac{1}{2}\pi & \frac{3}{4}\pi^2 & | & -1 \\ 0 & \pi & 2\pi^2 & | & -2 \end{pmatrix} \sim \begin{pmatrix} 1 & \frac{1}{2}\pi & \frac{1}{4}\pi^2 & | & 1 \\ 0 & \frac{1}{2}\pi & \frac{3}{4}\pi^2 & | & -1 \\ 0 & 0 & \frac{1}{2}\pi^2 & | & 0 \end{pmatrix}$$

$$\sim \begin{pmatrix} 1 & \frac{1}{2}\pi & 0 & | & 1 \\ 0 & \frac{1}{2}\pi & 0 & | & -1 \\ 0 & 0 & 1 & | & 0 \end{pmatrix} \sim \begin{pmatrix} 1 & 0 & 0 & | & 2 \\ 0 & \frac{1}{2}\pi & 0 & | & -1 \\ 0 & 0 & 1 & | & 0 \end{pmatrix}.$$

Die eindeutige Lösung ist $a_0 = 2$, $a_1 = -2/\pi$ und $a_2 = 0$, also

$$p(x) = 2 - \frac{2}{\pi}x = -\frac{2}{\pi}(x - \pi).$$

In Abbildung 3 ist der Graph des Polynoms p die blaue Gerade. Das einzige Polynom p vom Grad 2 mit der Eigenschaft $p(\pi/2) = 1$, $p(\pi) = 0$ und $p(3\pi/2) = -1$ ist eine lineare Funktion.

BEISPIEL 173. Für $f(x) = \sin x$, $x_0 = 0$, $x_1 = \pi/2$ und $x_2 = \pi$ ist ein Polynom zweiten Grades $p(x) = a_0 + a_1x + a_2x^2$ mit $p(0) = p(\pi) = 0$ und $p(\pi/2) = 1$ gesucht. Da dieses Polynom die Nullstellen $x_0 = 0$ und $x_2 = \pi$ hat, gilt $p(x) = x(x - \pi)c$ für $c \in \mathbb{R}$. Dieser Ansatz führt schneller zu einer Lösung als die Behandlung des linearen Gleichungssystems mit dem Gaußverfahren. Die Bedingung $p(\pi/2) = -c\pi^2/4 = 1$ liefert $c = -\frac{4}{\pi^2}$, also

$$p(x) = -\frac{4}{\pi^2}x(x - \pi).$$

In Abbildung 3 ist der Graph dieses Polynoms p die grüne Parabel.

BEISPIEL 174. Für $f(x) = \sin x$, $x_0 = 0$, $x_1 = \pi/2$, $x_2 = \pi$, $x_3 = \pi/2$ und $x_4 = 2\pi$ ist ein Polynom vierten Grades $p(x) = a_0 + a_1x + a_2x^2 + a_3x^3 + a_4x^4$ mit $p(0) = p(\pi) = p(2\pi) = 0$, $p(\pi/2) = 1$ und $p(3\pi/2) = -1$ gesucht. Da dieses Polynom die Nullstellen $x_0 = 0$, $x_2 = \pi$ und $x_4 = 2\pi$ hat, gilt $p(x) = x(x - \pi)(x - 2\pi)(ax + b)$ für $a, b \in \mathbb{R}$. Dieser Ansatz führt schneller zu einer Lösung als die Behandlung des linearen Gleichungssystems in Gleichung 129 mit dem Gaußverfahren. Die Bedingungen $p(\pi/2) = 1$ und $p(3\pi/2) = -1$ liefern die Gleichungen

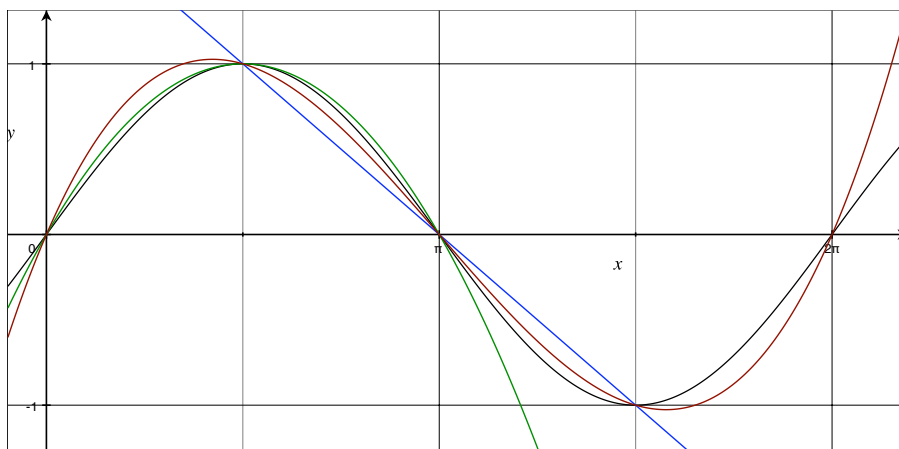
$$1 = \frac{3}{8}\pi^3 \left(a\frac{1}{2}\pi + b \right) \quad \text{und} \quad -1 = -\frac{3}{8}\pi^3 \left(a\frac{3}{2}\pi + b \right)$$

für die Unbekannten a und b . Daraus folgt $a = 0$ und $b = \frac{8}{3\pi^3}$, also

$$p(x) = \frac{8}{3\pi^3}x(x - \pi)(x - 2\pi).$$

In Abbildung 3 ist der Graph dieses Polynoms dritten Grades rot gefärbt.

Es seien $x_i, y_i \in \mathbb{R}$ für $i = 0, \dots, m$ gegeben. Ein Polynom p mit $\deg p \leq m$ und $p(x_i) = y_i$ für alle $i = 0, \dots, m$ heißt **Newtonpolynom**.

ABBILDUNG 3. Verschiedene Newtonpolynome für $\sin x$

SATZ 98 (Explizite Formel des Newtonpolynoms). Sind $x_0, x_1, \dots, x_m \in \mathbb{R}$ paarweise verschieden und $y_0, y_1, \dots, y_m \in \mathbb{R}$, so existiert genau ein Polynom p vom Grad $\leq m$ mit $p(x_i) = y_i$ für alle $i = 0, 1, \dots, m$. Es gilt

$$(130) \quad p(x) = \sum_{i=0}^m y_i \frac{\prod_{j \neq i} (x - x_j)}{\prod_{j \neq i} (x_i - x_j)}.$$

BEMERKUNG 75. Gleichung 130 liefert eine explizite Formel für die eindeutige Lösung des linearen Gleichungssystems 129. Allerdings ist jeder der $m + 1$ Summanden des Polynoms p ein Quotient, dessen Zähler und Nenner aus m Faktoren bestehen. Es ist relativ zeitaufwendig in der Gleichung 130 alle Terme auszumultiplizieren und nach Potenzen von x zu ordnen. Außerdem muss man sie einzelnen Summanden bei der Hinzunahme weiterer Stützstellen x_j immer wieder neu berechnen. Es existieren aber Methoden zur effektiven Bestimmung der Koeffizienten des Newtonpolynoms. $\square$

Möchte man eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ auf dem Intervall $[a, b]$ durch ein Polynom interpolieren und wählt man sehr viele Stützstellen $x_j \in [a, b]$, so erhält man i.A. ein Polynom hohen Grades, das an den Intervallenden zu starken Oszillationen neigt. Der schwarze Graph in Abbildung 4 zeigt das Newtonpolynom p mit $p(0) = p(3) = p(4) = p(8) = 1$, $p(1) = p(2) = p(6) = p(7) = 2$ und $p(5) = 3$.

Eine Idee zur Vermeidung dieser starken Schwankungen ist die Benutzung kubischer Splines auf Intervallen, die genau drei Stützstellen enthalten. Der rote Graph in Abbildung 4 ist aus vier kubischen Splines zusammengesetzt. Er nimmt an den Stützstellen $x_j = j$ die gleichen Funktionswerte an wie das Newtonpolynom p .

4.3. Kubische Splines. Es seien $x_0, x_1, x_2 \in \mathbb{R}$ und $y_0, y_1, y_2, y'_0 \in \mathbb{R}$ gegeben. Ein **kubischer Spline** ist ein Polynom p mit $\deg p \leq 3$, $p(x_i) = y_i$

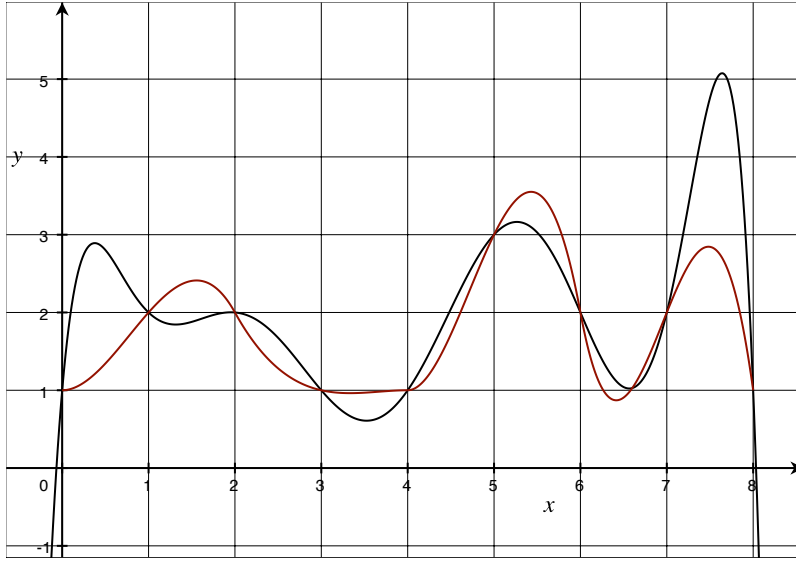


ABBILDUNG 4. Newtonpolynom und kubische Splines

für $i = 0, 1, 2$ und $p'(x_0) = y'_0$. Zusätzlich zu den Funktionswerten in drei Argumenten ist auch noch die erste Ableitung in einem Punkt bei einem Spline vorgegeben.

Mit dem Ansatz $p(x) = a_0 + a_1x + a_2x^2 + a_3x^3$ ergibt sich wegen $p'(x) = a_1 + 2a_2x + 3a_3x^2$ aus den vier Bedingungen das lineare Gleichungssystem mit der Koeffizientenmatrix

$$\begin{aligned}
 A &= \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 1 & x_1 & x_1^2 & x_1^3 \\ 1 & x_2 & x_2^2 & x_2^3 \\ 0 & 1 & 2x_0 & 3x_0^2 \end{pmatrix} \sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & x_1 - x_0 & x_1^2 - x_0^2 & x_1^3 - x_0^3 \\ 0 & x_2 - x_0 & x_2^2 - x_0^2 & x_2^3 - x_0^3 \\ 0 & 1 & 2x_0 & 3x_0^2 \end{pmatrix} \\
 &\sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & 1 & x_1 + x_0 & x_1^2 + x_1x_0 + x_0^2 \\ 0 & 1 & x_2 + x_0 & x_2^2 + x_2x_0 + x_0^2 \\ 0 & 1 & 2x_0 & 3x_0^2 \end{pmatrix} \\
 &\sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & 1 & x_1 + x_0 & x_1^2 + x_1x_0 + x_0^2 \\ 0 & 0 & x_2 - x_1 & (x_0 + x_1 + x_2)(x_2 - x_1) \\ 0 & 0 & x_0 - x_1 & (x_0 - x_1)(2x_0 + x_1) \end{pmatrix} \\
 &\sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & 1 & x_1 + x_0 & x_1^2 + x_1x_0 + x_0^2 \\ 0 & 0 & 1 & x_0 + x_1 + x_2 \\ 0 & 0 & 1 & 2x_0 + x_1 \end{pmatrix}
 \end{aligned}$$

$$\sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & 1 & x_1 + x_0 & x_1^2 + x_1x_0 + x_0^2 \\ 0 & 0 & 1 & x_0 + x_1 + x_2 \\ 0 & 0 & 0 & x_0 - x_2 \end{pmatrix} \sim \begin{pmatrix} 1 & x_0 & x_0^2 & x_0^3 \\ 0 & 1 & x_1 + x_0 & x_1^2 + x_1x_0 + x_0^2 \\ 0 & 0 & 1 & x_0 + x_1 + x_2 \\ 0 & 0 & 0 & 1 \end{pmatrix},$$

also $\text{rang } A = 4$, falls $x_0 \neq x_1, x_2$ und $x_1 \neq x_2$. Falls x_0, x_1, x_2 drei paarweise verschiedene Stützstellen sind, so existiert immer ein eindeutiger kubischer Spline.

Möchte man die kubischen Splines p_1 mit den Stützstellen $x_0 < x_1 < x_2$ und p_2 mit den Stützstellen $x_2 < x_3 < x_4$ so zusammensetzen, dass der Übergang von p_1 zu p_2 in $x = x_2$ im Graph nicht sichtbar ist, so benutzt man für p_2 die Bedingungen $p_2(x_2) = p_1(x_2)$ und $p_2'(x_2) = p_1'(x_2)$. Die Funktionswerte und die ersten Ableitungen von p_1 und p_2 in $x = x_2$ stimmen also überein.

KAPITEL XII

Lineare Vektorräume

Ein linearer $\mathbb{R}$ -**Vektorraum** ist eine Menge V zusammen mit zwei Abbildungen $V \times V \rightarrow V$, $(v, w) \mapsto v + w$, und $\mathbb{R} \times V \rightarrow V$, $(\lambda, v) \mapsto \lambda v$, die folgende Eigenschaften haben:

- Für beliebige $u, v, w \in V$ gilt $(u + v) + w = u + (v + w)$ und $v + w = w + v$. (Die Vektoraddition ist assoziativ und kommutativ.)
- Für beliebige $\lambda, \mu \in \mathbb{R}$ und $v \in V$ gilt $\lambda(\mu v) = (\lambda\mu)v$. (Die Multiplikation mit Skalaren ist assoziativ und kommutativ.)
- Distributivgesetze: Für beliebige $\lambda, \mu \in \mathbb{R}$ und $v, w \in V$ gilt $\lambda(v + w) = \lambda v + \lambda w$ und $(\lambda + \mu)v = \lambda v + \mu v$.

Die Elemente eines Vektorraumes heißen **Vektoren**.

BEMERKUNG 76. Der Begriff des Vektors wurde für das Verständnis gerichteter Größen, zum Beispiel der Kraft, entwickelt. Der Vektoraddition entspricht die Summe zweier Kräfte, die am gleichen Punkt angreifen. Der Skalarmultiplikation entspricht die Vervielfachung der Kraft unter Beibehaltung der Richtung. Aber auch weniger anschauliche Objekte, zum Beispiel Funktionen, können als Elemente eines Vektorraumes betrachtet werden, weil man sie addieren und mit reellen Zahlen multiplizieren kann und wieder ein Objekt derselben Art erhält. $\square$

BEISPIEL 175. Die Menge $\mathbb{R}^2 := \{(x, y) : x, y \in \mathbb{R}\}$ ist die Menge aller Paare reeller Zahlen. Die komponentenweise definierte Addition und Skalarmultiplikation

$$(x_1, y_1) + (x_2, y_2) := (x_1 + x_2, y_1 + y_2) \text{ und } \lambda(x, y) := (\lambda x, \lambda y)$$

für beliebige $x, y, \lambda \in \mathbb{R}$ macht $\mathbb{R}^2$ zu einem $\mathbb{R}$ -Vektorraum.

BEISPIEL 176. Die Menge $\mathbb{R}^3 := \{(x, y, z) : x, y, z \in \mathbb{R}\}$ ist die Menge aller Tripel reeller Zahlen. Die komponentenweise definierte Addition und Skalarmultiplikation

$$(x_1, y_1, z_1) + (x_2, y_2, z_2) := (x_1 + x_2, y_1 + y_2, z_1 + z_2) \text{ und } \lambda(x, y, z) := (\lambda x, \lambda y, \lambda z)$$

für beliebige $x, y, z, \lambda \in \mathbb{R}$ macht $\mathbb{R}^3$ zu einem $\mathbb{R}$ -Vektorraum.

BEISPIEL 177. Die Menge $\mathbb{R}^n := \{\vec{x} = (x_1, x_2, \dots, x_n) : x_j \in \mathbb{R} \forall j\}$ ist die Menge aller n -Tupel reeller Zahlen. Die komponentenweise definierte Addition und Skalarmultiplikation

$$\vec{x} + \vec{y} := (x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \text{ und } \lambda \vec{x} := (\lambda x_1, \lambda x_2, \dots, \lambda x_n)$$

für beliebige $\lambda \in \mathbb{R}$ und $\vec{x}, \vec{y} \in \mathbb{R}^n$ mit $\vec{x} = (x_1, \dots, x_n), \vec{y} = (y_1, \dots, y_n)$ macht $\mathbb{R}^n$ zu einem $\mathbb{R}$ -Vektorraum.

BEMERKUNG 77. Man kann ein n -Tupel auch als $x = (x_1, \dots, x_n)$ (ohne Pfeil) notieren. Möchte man aber unterstreichen, dass das n -Tupel ein Vektor ist, der aus n reellen Zahlen besteht, und keine reelle Zahl, so ist die Notation $\vec{x} = (x_1, \dots, x_n)$ aussagekräftiger. $\square$

BEISPIEL 178. Die Menge reeller Zahlenfolgen mit der Addition $\{x_n\}_{n \in \mathbb{R}} + \{y_n\}_{n \in \mathbb{R}} = \{x_n + y_n\}_{n \in \mathbb{R}}$ und der Skalarmultiplikation $\lambda \{x_n\}_{n \in \mathbb{R}} = \{\lambda x_n\}_{n \in \mathbb{R}}$ ist ein $\mathbb{R}$ -Vektorraum.

BEISPIEL 179. $\mathcal{C}([a, b])$ bezeichnet die Menge aller auf dem Intervall $[a, b]$ stetigen Funktionen. Wenn $f, g \in \mathcal{C}([a, b])$, dann ist $f + g$ mit $(f + g)(x) = f(x) + g(x)$ wieder auf $[a, b]$ stetig. Wenn $\lambda \in \mathbb{R}$ und $f \in \mathcal{C}([a, b])$, dann ist λf mit $(\lambda f)(x) = \lambda f(x)$ wieder auf $[a, b]$ stetig. Die Menge $\mathcal{C}([a, b])$ ist ein $\mathbb{R}$ -Vektorraum.

BEISPIEL 180. $\mathbb{R}[x]$ bezeichnet die Menge aller Polynome mit reellen Koeffizienten in der Variablen x . Es gilt $\mathbb{R}[x] \subset \mathcal{C}([a, b])$. Wenn $p, q \in \mathbb{R}[x]$ und $\lambda \in \mathbb{R}$, dann $p + q \in \mathbb{R}[x]$ und $\lambda p \in \mathbb{R}[x]$. Die reellen Polynome bilden einen $\mathbb{R}$ -Vektorraum. $\mathbb{R}[x]$ ist ein Untervektorraum von $\mathcal{C}([a, b])$.

1. Basis, Koordinaten, Dimension

Es seien V ein $\mathbb{R}$ -Vektorraum und $\vec{v}_1, \dots, \vec{v}_n \in V$ Vektoren in V . Eine **Linearkombination** der Vektoren $\vec{v}_1, \dots, \vec{v}_n$ ist ein Vektor der Form

$$\alpha_1 \vec{v}_1 + \dots + \alpha_n \vec{v}_n = \sum_{j=1}^n \alpha_j \vec{v}_j \text{ mit } \alpha_1, \dots, \alpha_n \in \mathbb{R}.$$

Eine Menge $\{\vec{v}_1, \dots, \vec{v}_n\} \subset V$ heißt **linear unabhängig**, falls $\alpha_1 = \dots = \alpha_n = 0$ für jede Linearkombination $\alpha_1 \vec{v}_1 + \dots + \alpha_n \vec{v}_n = \vec{0}$ des Nullvektors $\vec{0}$ gilt. Eine Menge $\mathcal{B} = \{\vec{v}_1, \dots, \vec{v}_n\} \subset V$ heißt **Basis** des Vektorraumes V , falls $\mathcal{B}$ linear unabhängig ist und sich jeder Vektor aus V als Linearkombination der Vektoren in $\mathcal{B}$ schreiben läßt. Die Koeffizienten α_j der Linearkombination $\vec{v} = \alpha_1 \vec{v}_1 + \dots + \alpha_n \vec{v}_n$ heißen **Koordinaten** des Vektors $\vec{v}$ bezüglich der Basis $\mathcal{B}$.

BEISPIEL 181. Die Menge $\{(1, 0), (0, 1)\} \subset \mathbb{R}^2$ ist eine Basis des Vektorraumes $\mathbb{R}^2$, denn $(x, y) = x(1, 0) + y(0, 1)$ für alle $x, y \in \mathbb{R}$ und

$$\alpha_1(1, 0) + \alpha_2(0, 1) = (\alpha_1, \alpha_2) = (0, 0) \Leftrightarrow \alpha_1 = \alpha_2 = 0.$$

Die Koordinaten des Paares (x, y) bezüglich dieser Basis sind gerade die beiden Komponenten des Paares.

BEISPIEL 182. Die Menge $\{(1, 0, 0), (0, 1, 0), (0, 0, 1)\} \subset \mathbb{R}^3$ ist eine Basis des Vektorraumes $\mathbb{R}^3$, denn $(x, y, z) = x(1, 0, 0) + y(0, 1, 0) + z(0, 0, 1)$ für alle $x, y, z \in \mathbb{R}$ und

$$\alpha_1(1, 0, 0) + \alpha_2(0, 1, 0) + \alpha_3(0, 0, 1) = (\alpha_1, \alpha_2, \alpha_3) = (0, 0, 0) \Leftrightarrow \alpha_1 = \alpha_2 = \alpha_3 = 0.$$

Die Koordinaten des Tripels (x, y, z) bezüglich dieser Basis sind gerade die drei Komponenten des Tripels.

Die Menge $\mathcal{B} = \{\vec{e}_1, \vec{e}_2, \dots, \vec{e}_n\} \subset \mathbb{R}^n$ mit $\vec{e}_1 = (1, 0, \dots, 0)$, $\vec{e}_2 = (0, 1, 0, \dots, 0)$ usw., also

$$\vec{e}_j = (x_1, x_2, \dots, x_n) \text{ mit } x_i = \begin{cases} 0 & , \text{ falls } i \neq j \\ 1 & , \text{ falls } i = j \end{cases},$$

ist die **Standardbasis** des $\mathbb{R}^n$. Für jeden Vektor $\vec{x} \in \mathbb{R}^n$ mit $\vec{x} = (x_1, x_2, \dots, x_n)$ gilt $\vec{x} = x_1\vec{e}_1 + \dots + x_n\vec{e}_n$. Die Komponenten des n -Tupels sind die Koordinaten des n -Tupels bezüglich der Standardbasis.

Jeder Vektorraum besitzt eine Basis. Diese Basis ist nicht eindeutig. Sind $\mathcal{B}$ und $\mathcal{B}'$ zwei Basen des Vektorraumes V , so gilt $|\mathcal{B}| = |\mathcal{B}'|$, d.h. alle Basen eines Vektorraumes besitzen gleich viel Elemente. Die Anzahl der Elemente einer Basis $\mathcal{B}$ eines Vektorraumes V ist die **Dimension** $\dim V := |\mathcal{B}|$ des Vektorraumes. Zum Beispiel gilt $\dim \mathbb{R}^n = n$.

BEISPIEL 183. Die Menge der Monome

$$\{1, x, x^2, x^3, x^4, \dots\} = \{x^n : n \in \mathbb{N}\}$$

ist eine Basis des Vektorraumes $\mathbb{R}[x]$. Die Koeffizienten eines Polynoms sind die Koordinaten des Polynoms bezüglich dieser Basis. Es gilt $\dim \mathbb{R}[x] = \infty$.

LEMMA 27. Es seien $\vec{v}_1, \dots, \vec{v}_m \in \mathbb{R}^n$ und A die $n \times m$ -Matrix der Koordinaten der Vektoren $\vec{v}_j$ bezüglich der Standardbasis, also

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1m} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nm} \end{pmatrix} \text{ mit } \vec{v}_j = \sum_{i=1}^n a_{ij}\vec{e}_i \text{ für alle } j = 1, \dots, m.$$

Die Menge $\{\vec{v}_1, \dots, \vec{v}_m\}$ ist genau dann linear unabhängig, wenn $\text{rang } A = m$.

BEWEIS. Das lineare Gleichungssystem $\vec{0} = \alpha_1\vec{v}_1 + \dots + \alpha_m\vec{v}_m$ für die Variablen α_j besitzt genau dann eine eindeutige Lösung, wenn $\text{rang } A = m$. $\square$

SATZ 99. Es seien $\vec{v}_1, \dots, \vec{v}_n \in \mathbb{R}^n$ und A die $n \times n$ -Matrix der Koordinaten der Vektoren $\vec{v}_j$ bezüglich der Standardbasis, also

$$A = \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{pmatrix} \text{ mit } \vec{v}_j = \sum_{i=1}^n a_{ij}\vec{e}_i \text{ für alle } j = 1, \dots, n.$$

Die Menge $\{\vec{v}_1, \dots, \vec{v}_n\}$ ist genau dann eine Basis, wenn $\text{rang } A = n$.

BEWEIS. Das lineare Gleichungssystem $\vec{v} = \alpha_1\vec{v}_1 + \dots + \alpha_n\vec{v}_n$ für die Variablen α_j besitzt genau dann eine eindeutige Lösung, wenn $\text{rang } A = n$. $\square$

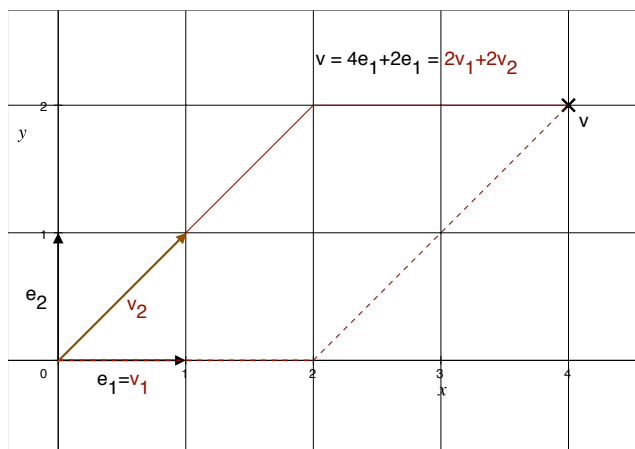


ABBILDUNG 1

BEISPIEL 184. Die Vektoren $\vec{v}_1 := (1, 0)$ und $\vec{v}_2 := (1, 1)$ bilden eine Basis des $\mathbb{R}^2$, denn $\vec{v}_1 = 1\vec{e}_1 + 0\vec{e}_2$, $\vec{v}_2 = 1\vec{e}_1 + 1\vec{e}_2$ und die Matrix

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

hat Rang 2. Abbildung 1 zeigt die Standardbasis des $\mathbb{R}^2$ (schwarz), die Basisvektoren $\vec{v}_1$ und $\vec{v}_2$ (rot) und den Vektor $\vec{v} = 4\vec{e}_1 + 2\vec{e}_2$. Dieser Vektor hat bezüglich der Basis $\{\vec{v}_1, \vec{v}_2\}$ die Koordinaten (2, 2).

2. Norm und Längenmessung

Eine **Norm** eines Vektorraumes V ist eine Abbildung $V \rightarrow \mathbb{R}$, $v \mapsto \|v\|$, mit folgenden Eigenschaften:

- Für alle $v \in V$ gilt $\|v\| \geq 0$. Es gilt $\|v\| = 0$ genau dann, wenn $v = 0$.
- Für alle $\lambda \in \mathbb{R}$ und alle $v \in V$ gilt $\|\lambda v\| = |\lambda| \|v\|$.
- Für alle $v, w \in V$ gilt $\|v + w\| \leq \|v\| + \|w\|$. (Dreiecksungleichung)

Die Norm misst die Länge eines Vektors.

BEISPIEL 185 (Euklidische Norm auf dem $\mathbb{R}^n$). Auf dem Vektorraum $\mathbb{R}^n$ wird durch

$$\|\vec{x}\| := \sqrt{x_1^2 + \dots + x_n^2} = \sqrt{\sum_{j=1}^n x_j^2} \text{ für } \vec{x} = (x_1, \dots, x_n)$$

eine Norm definiert. Sie heißt **Euklidische Norm** und verallgemeinert den Satz von Pythagoras. Die Standardbasisvektoren $\vec{e}_j$ haben bezüglich der Euklidischen Norm die Länge 1, d.h. $\|\vec{e}_j\| = 1$ für alle j .

Ist ein Vektorraum V mit einer Norm versehen, so kann man auch den **Abstand**, die Entfernung, zweier Vektoren messen. Für $v, w \in V$ definiert

man den Abstand mit Hilfe der Norm durch $d(v, w) := \|v - w\|$. Diese Definition ist sinnvoll, denn es gilt $v = (v - w) + w$ (siehe Abbildung 2). Diese Abstandsdefinition hat die folgenden Eigenschaften:

- Für alle $v, w \in V$ gilt $d(v, w) \geq 0$. Es gilt $d(v, w) = 0 \Leftrightarrow v = w$.
- Für alle $v, w \in V$ gilt $d(v, w) = d(w, v)$.
- Für alle $u, v, w \in V$ gilt $d(u, v) + d(v, w) \geq d(u, w)$. (Dreiecksungleichung)

BEISPIEL 186 (Maximumnorm). Auf dem Vektorraum $\mathcal{C}([a, b])$ wird durch

$$\|f\|_\infty := \max\{|f(x)| : x \in [a, b]\}$$

eine Norm definiert. Der daraus abgeleitete Abstandsbegriff $d(f, g)$ misst die maximale Differenz der Funktionswerte von f und g auf dem Intervall $[a, b]$.

BEMERKUNG 78. Eine Funktionenfolge $\{f_n\}_{n \in \mathbb{N}}$ konvergiert genau dann bezüglich der Maximumnorm gegen eine Funktion f , wenn für jedes $\varepsilon > 0$ ein $n_0 \in \mathbb{N}$ mit $|f(x) - f_n(x)| < \varepsilon$ für alle $x \in [a, b]$ und $n \geq n_0$ existiert. Diese Konvergenz nennt man gleichmäßige Konvergenz der Funktionenfolge $\{f_n\}$ gegen die Funktion f auf dem Intervall $[a, b]$.

Zum Beispiel benutzt man konvergente Funktionenfolgen zur Konstruktion bzw. Approximation von Lösungen von Differentialgleichungen. $\square$

BEISPIEL 187 (Nichteuklidische Normen auf $\mathbb{R}^n$). Auch durch

$$\|\vec{x}\|_1 := \sum_{j=1}^n |x_j| \text{ und } \|\vec{x}\|_\infty := \max\{|x_1|, |x_2|, \dots, |x_n|\} \text{ für } \vec{x} = (x_1, \dots, x_n)$$

werden sinnvolle Normen auf $\mathbb{R}^n$ definiert. Die Norm $\|\cdot\|_1$ misst Länge des Vektors $\vec{x}$, falls man sich nur in Richtung der Basisvektoren bewegen darf. Die Norm $\|\cdot\|_\infty$ misst, ähnlich wie in Beispiel 186 das Maximum der Komponenten. Der Vektor $\vec{v} := (3, 4) \in \mathbb{R}^2$ hat die euklidische Norm $\|\vec{v}\| = \sqrt{3^2 + 4^2} = 5$, die Maximumsnorm $\|\vec{v}\|_\infty = \max(|3|, |4|) = 4$ und $\|\vec{v}\|_1 = |3| + |4| = 7$ (siehe Abbildung 3).

BEMERKUNG 79. Wenn $\vec{v} \in V$ und $\vec{v} \neq \vec{0}$, dann gilt

$$\left\| \frac{1}{\|\vec{v}\|} \vec{v} \right\| = \frac{1}{\|\vec{v}\|} \|\vec{v}\| = 1.$$

Die Skalierung eines von $\vec{0}$ verschiedenen Vektors mit dem Faktor $\|\vec{v}\|^{-1}$ heißt **Normierung** des Vektors. Normiert man die Vektoren einer Basis, so erhält man eine Basis aus Vektoren der Länge 1. $\square$

3. Skalarprodukt und Winkelmessung

Ein **Skalarprodukt** auf einem $\mathbb{R}$ -Vektorraum V ist eine Abbildung

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}, \quad \langle v, w \rangle,$$

mit folgenden Eigenschaften:

- Es gilt $\langle v, w \rangle = \langle w, v \rangle$ für beliebige $v, w \in V$. (Symmetrie)

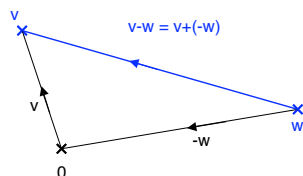


ABBILDUNG 2

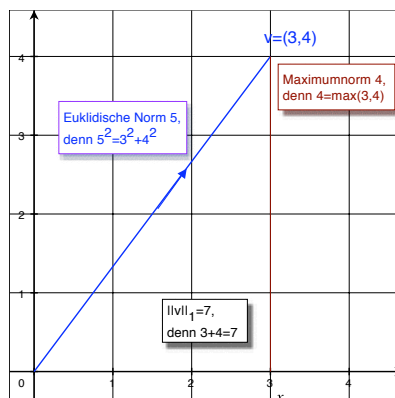


ABBILDUNG 3

- Es gilt $\langle \lambda v, w \rangle = \lambda \langle v, w \rangle$ und $\langle u + v, w \rangle = \langle u, w \rangle + \langle v, w \rangle$ für beliebige $\lambda \in \mathbb{R}$ und $u, v, w \in V$. (Bilinearität)
- Für alle $v \in V$ gilt $\langle v, v \rangle \geq 0$. Es gilt $\langle v, v \rangle = 0 \Leftrightarrow v = 0$.

Zwei Vektoren $v, w \in V$ sind **orthogonal** zueinander (stehen senkrecht aufeinander), falls $\langle v, w \rangle = 0$. Man schreibt auch $v \perp w$.

BEISPIEL 188 (Skalarprodukt auf $\mathbb{R}^n$). Auf $\mathbb{R}^n$ wird durch

$$\langle \vec{x}, \vec{y} \rangle := x_1 y_1 + x_2 y_2 + \dots + x_n y_n = \sum_{j=1}^n x_j y_j \text{ für } \vec{x} = (x_1, \dots, x_n), \vec{y} = (y_1, \dots, y_n)$$

das **Standardskalarprodukt** definiert. Es gilt

$$\langle \vec{e}_j, \vec{e}_k \rangle = \begin{cases} 1 & , \text{ falls } j = k \\ 0 & , \text{ falls } j \neq k \end{cases}.$$

Die Basisvektoren der Standardbasis stehen bezüglich dieses Skalarprodukts senkrecht aufeinander. Für alle $\vec{x} = (x_1, \dots, x_n) \in \mathbb{R}^n$ gilt $x_j = \langle \vec{x}, \vec{e}_j \rangle$ für alle j .

Ist der Vektorraum V mit einem Skalarprodukt versehen, so ist durch

$$(131) \quad \|v\| := \sqrt{\langle v, v \rangle}$$

eine Norm auf V definiert. Aus dem Standardskalarprodukt auf $\mathbb{R}^n$ erhält man auf diesem Wege die Euklidische Norm auf $\mathbb{R}^n$, denn $\langle \vec{x}, \vec{x} \rangle = x_1^2 + \dots + x_n^2$ für alle $\vec{x} = (x_1, \dots, x_n) \in \mathbb{R}^n$.

Eine Basis eines Vektorraumes mit Skalarprodukt heißt **Orthonormalbasis**, falls alle Basisvektoren Länge 1 haben und paarweise orthogonal zueinander sind. Eine Basis $\{\vec{v}_1, \dots, \vec{v}_n\}$ ist genau dann eine Orthonormalbasis, wenn $\langle \vec{v}_j, \vec{v}_k \rangle = 0$ für alle $j \neq k$ und $\langle \vec{v}_j, \vec{v}_j \rangle = 1$ für alle j . Die Standardbasis des $\mathbb{R}^n$ ist eine Orthonormalbasis.

BEMERKUNG 80. Mit Hilfe des Schmidtschen Orthonormalisierungsverfahrens kann für jeden Vektorraum mit Skalarprodukt aus einer Basis eine Orthonormalbasis konstruieren. $\square$

SATZ 100. Es seien $\vec{v}_1, \dots, \vec{v}_n \in \mathbb{R}^n$. Wenn $\|\vec{v}_j\| = 1$ für alle j und $\langle \vec{v}_j, \vec{v}_k \rangle = 0$ für alle $j \neq k$, dann ist $\{\vec{v}_1, \dots, \vec{v}_n\}$ eine Orthonormalbasis des $\mathbb{R}^n$.

BEISPIEL 189. Für beliebige $\varphi \in \mathbb{R}$ ist $\{(\cos \varphi, \sin \varphi), (-\sin \varphi, \cos \varphi)\}$ eine Orthonormalbasis des $\mathbb{R}^2$, denn

$$\begin{aligned}\|(\cos \varphi, \sin \varphi)\| &= \sqrt{\cos^2 \varphi + \sin^2 \varphi} = 1 \\ \|(-\sin \varphi, \cos \varphi)\| &= \sqrt{(-\sin \varphi)^2 + \cos^2 \varphi} = \sqrt{\sin^2 \varphi + \cos^2 \varphi} = 1 \\ \langle (\cos \varphi, \sin \varphi), (-\sin \varphi, \cos \varphi) \rangle &= -\cos \varphi \sin \varphi + \sin \varphi \cos \varphi = 0.\end{aligned}$$

Wenn $\{\vec{v}_1, \dots, \vec{v}_n\}$ eine Orthonormalbasis des Vektorraumes V ist, dann lassen sich die Koordinaten bezüglich dieser Basis leicht mit Hilfe des Skalarprodukts bestimmen, denn es gilt

$$(132) \quad \vec{v} = \sum_{j=1}^n \langle \vec{v}, \vec{v}_j \rangle \vec{v}_j \text{ für alle } \vec{v} \in V.$$

Die Abbildung $\pi_j : V \rightarrow \mathbb{R}\vec{v}_j$, $\vec{v} \mapsto \langle \vec{v}, \vec{v}_j \rangle \vec{v}_j$, heißt Orthogonalprojektion auf den Untervektorraum $\mathbb{R}\vec{v}_j$. Für alle $\vec{v} \in V$ gilt $\pi_j(\pi_j(\vec{v})) = \pi_j(\vec{v})$. Eine **Projektion** ist eine lineare Abbildung $\pi : V \rightarrow V$ mit der Eigenschaft $\pi \circ \pi = \pi$.

Für den **Winkel** $\varphi := \angle(\vec{v}, \vec{w})$ zwischen zwei Vektoren $\vec{v}, \vec{w} \in V$ gilt

$$(133) \quad \cos \varphi = \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \|\vec{w}\|}.$$

Die Schwarzsche Ungleichung (Satz 13) sichert

$$-1 \leq \frac{\langle \vec{v}, \vec{w} \rangle}{\|\vec{v}\| \|\vec{w}\|} \leq 1.$$

BEISPIEL 190. Der Winkel φ zwischen den Vektoren $(1, 0), (1, 1) \in \mathbb{R}^2$ beträgt 45° , denn die Punkte $(0, 0)$, $(1, 0)$ und $(1, 1)$ bilden ein gleichschenkliges, rechtwinkliges Dreieck. Wir bestätigen

$$\frac{\langle (1, 0), (1, 1) \rangle}{\|(1, 0)\| \|(1, 1)\|} = \frac{1 \cdot 1 + 0 \cdot 1}{\sqrt{1^2 + 0^2} \sqrt{1^2 + 1^2}} = \frac{1}{\sqrt{2}} = \frac{\sqrt{2}}{2} = \cos 45^\circ.$$

4. Untervektorräume

Eine Teilmenge $U \subset V$ eines Vektorraumes V heißt **Untervektorraum**, falls $\lambda u \in U$ und $u + w \in U$ für alle $\lambda \in \mathbb{R}$ und $u, w \in U$.

BEISPIEL 191 (Triviale Untervektorräume). Für jeden Vektorraum V sind $\{\vec{0}\}$ und V Untervektorräume von V .

4.1. Aufgespannte Unterräume. Der kleinste Untervektorraum, der die Vektoren $\vec{v}_1, \dots, \vec{v}_m \in V$ beinhaltet, also in V von den Vektoren $\vec{v}_1, \dots, \vec{v}_m$ aufgespannt wird, ist

$$(134) \quad \text{span}\{\vec{v}_1, \dots, \vec{v}_m\} := \{\alpha_1 \vec{v}_1 + \dots + \alpha_m \vec{v}_m : \alpha_j \in \mathbb{R} \forall j\}.$$

BEISPIEL 192 (Eindimensionale Untervektorräume). Wenn V ein Vektorraum ist, dann spannt jeder Vektor $\vec{v} \in V$ mit $\vec{v} \neq \vec{0}$ den eindimensionalen Untervektorraum $\text{span } \vec{v} := \{\lambda \vec{v} : \lambda \in \mathbb{R}\} \subset V$ auf. Dem Vektorraum $\text{span } \vec{v}$ entspricht die Gerade durch den Ursprung in Richtung $\vec{v}$.

BEISPIEL 193. Zwei linear unabhängige Vektoren $\vec{v}_1, \vec{v}_2 \in V$ spannen die Ebene $E = \text{span}\{\vec{v}_1, \vec{v}_2\}$ durch den Ursprung auf.

4.2. Komplementäre Unterräume. Wenn V ein n -dimensionaler Vektorraum mit Skalarprodukt ist, dann definiert jeder Vektor $\vec{v} \in V$ mit $\vec{v} \neq \vec{0}$ den $(n-1)$ -dimensionalen Untervektorraum

$$(135) \quad \vec{v}^\perp := \{\vec{w} \in V : \vec{w} \perp \vec{v}\} = \{\vec{w} \in V : \langle \vec{w}, \vec{v} \rangle = 0\} \subset V.$$

Der Vektorraum $\vec{v}^\perp$ heißt **Orthogonalkomplement** von $\vec{v}$ bzw. $\text{span } \vec{v}$. Er besteht aus allen Vektoren in V , die senkrecht auf $\vec{v}$ stehen.

BEISPIEL 194. Das Orthogonalkomplement von $\vec{v} = (2, 1) \in \mathbb{R}^2$ ist ein eindimensionaler Untervektorraum. Er wird von einem Vektor, der senkrecht auf $\vec{v}$ steht, aufgespannt, z.B. $\vec{v}^\perp = \text{span}(-1, 2)$, denn $\langle (2, 1), (-2, 1) \rangle = 0$. Abbildung 4 zeigt die eindimensionalen Unterräume $\text{span } \vec{v}$ (blau) und $\vec{v}^\perp$ (rot).

Die lineare Gleichung $a_1 x_1 + a_2 x_2 + \dots + a_n x_n = 0$ kann man mit Hilfe des Skalarproduktes auf $\mathbb{R}^n$ und den Bezeichnungen $\vec{a} = (a_1, \dots, a_n)$ und $\vec{x} = (x_1, \dots, x_n)$ kurz als $\langle \vec{a}, \vec{x} \rangle = 0$ schreiben. Die Lösungsmenge der linearen Gleichung $\langle \vec{a}, \vec{x} \rangle = 0$ ist der $(n-1)$ -dimensionale Unterraum $\vec{a}^\perp \subset \mathbb{R}^n$, falls $\vec{a} \neq \vec{0}$.

4.3. Durchschnitt von Unterräumen. Sind $V_1, V_2 \subset V$ zwei Untervektorräume, dann ist auch der Durchschnitt $V_1 \cap V_2$ ein Untervektorraum von V .

5. Lösungsmenge linearer Gleichungssysteme

Die Lösungsmenge des linearen Gleichungssystems $\sum_{j=1}^n a_{ij} x_j = 0$ für $i = 1, \dots, m$ ist der Durchschnitt der Orthogonalkomplemente der Zeilenvektoren $\vec{a}_i = (a_{i1}, \dots, a_{in})$ für $i = 1, \dots, m$, also der Untervektorraum

$$(136) \quad \vec{a}_1^\perp \cap \dots \cap \vec{a}_m^\perp.$$

Es sei $\vec{x}_0$ eine Lösung des LGSs $\sum_{j=1}^n a_{ij} x_j = b_i$, $i = 1, \dots, m$. Dann ist $\vec{x} \in \mathbb{R}^n$ genau dann eine Lösung des LGSs $\sum_{j=1}^n a_{ij} x_j = b_i$ für $i =$

$1, \dots, m$, wenn $\vec{x} - \vec{x}_0$ eine Lösung des dazugehörigen *homogenen* linearen Gleichungssystems $\sum_{j=1}^n a_{ij}x_j = 0$ für $i = 1, \dots, m$ ist, denn

$$\langle \vec{a}_i, \vec{x} - \vec{x}_0 \rangle = \langle \vec{a}_i, \vec{x} \rangle - \langle \vec{a}_i, \vec{x}_0 \rangle = \langle \vec{a}_i, \vec{x} \rangle - b_i = 0 \Leftrightarrow \langle \vec{a}_i, \vec{x} \rangle = b_i.$$

Die Lösungsmenge des LGSs $\sum_{j=1}^n a_{ij}x_j = b_i$ für $i = 1, \dots, m$ ist

$$(137) \quad \vec{x}_0 + \vec{a}_1^\perp \cap \dots \cap \vec{a}_m^\perp.$$

6. Geraden in $\mathbb{R}^2$

Eine Gerade in $\mathbb{R}^n$ ist durch zwei verschiedene Punkte $\vec{v}_0, \vec{v}_1 \in \mathbb{R}^n$ eindeutig bestimmt.

6.1. Parameterform. Die Gerade durch die Punkte $\vec{v}_0, \vec{v}_1 \in \mathbb{R}^n$ ist die Gerade durch $\vec{v}_0$ in Richtung $\vec{v} := \vec{v}_1 - \vec{v}_0$, also die Menge

$$\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\} \subset \mathbb{R}^n.$$

Diese Beschreibung einer Gerade heißt **Parameterform**, weil die Punkte der Geraden mit Hilfe der Variablen t parametrisiert werden. Der Vektor $\vec{v}$ heißt **Richtungsvektor** der Geraden.

6.2. Der Normalenvektor. Falls $\vec{v} \in \mathbb{R}^2$ und $\vec{v} \neq \vec{0}$, so wird das Orthogonalkomplement zu $\vec{v}$ von einem von $\vec{0}$ verschiedenen Vektor aufgespannt, der senkrecht auf $\vec{v}$ steht. Für $\vec{v} = (v_1, v_2)$ gilt

$$\vec{v}^\perp = \text{span}(-v_2, v_1) = \{\lambda(-v_2, v_1) : \lambda \in \mathbb{R}\},$$

denn $\langle \vec{v}, (-v_2, v_1) \rangle = 0$. Jeder von $\vec{0}$ verschiedene Vektor in $\vec{v}^\perp$, z.B. $(-v_2, v_1)$, heißt **Normalenvektor** der Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\} \subset \mathbb{R}^2$ und wird mit $\vec{n}$ bezeichnet. Es gilt $\vec{n}^\perp = \text{span } \vec{v}$.

6.3. Lineare Form. Ein Punkt $\vec{x} \in \mathbb{R}^2$ liegt genau dann auf der Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$, wenn $\vec{x} = \vec{v}_0 + t\vec{v}$ für ein $t \in \mathbb{R}$. Da $\langle \vec{v}_0 + t\vec{v}, \vec{n} \rangle = \langle \vec{v}_0, \vec{n} \rangle + t\langle \vec{v}, \vec{n} \rangle = \langle \vec{v}_0, \vec{n} \rangle$, gilt

$$\begin{aligned} \vec{x} = \vec{v}_0 + t\vec{v} \text{ für ein } t \in \mathbb{R} &\Leftrightarrow \vec{x} - \vec{v}_0 \in \text{span } \vec{v} \Leftrightarrow \langle \vec{x} - \vec{v}_0, \vec{n} \rangle = 0 \\ &\Leftrightarrow \langle \vec{x}, \vec{n} \rangle = \langle \vec{v}_0, \vec{n} \rangle. \end{aligned}$$

Die Bedingung $\langle \vec{x}, \vec{n} \rangle = \langle \vec{v}_0, \vec{n} \rangle$ ist eine lineare Gleichung für die Koordinaten x_1, x_2 des Vektors $\vec{x} = (x_1, x_2)$. Diese lineare Gleichung heißt **lineare Form** der Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$, weil sie die Gerade als Lösungsmenge einer linearen Gleichung beschreibt.

BEISPIEL 195. Mit $\vec{v} = (2, 1)$ und $\vec{v}_0 = (-1, 1)$ erhält man die Gerade

$$\left\{ \begin{pmatrix} -1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 2 \\ 1 \end{pmatrix} : t \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} -1 + 2t \\ 1 + t \end{pmatrix} : t \in \mathbb{R} \right\}.$$

Ein Normalenvektor ist $\vec{n} = (-1, 2)$. Er liefert die lineare Gleichung $\langle \vec{x}, \vec{n} \rangle = \langle (-1, 1), (-1, 2) \rangle$, also $-x_1 + 2x_2 = 3$. Man sieht leicht (Abbildung 5), dass die Lösungsmenge dieser linearen Gleichung die Gerade durch die Punkte $(-3, 0)$ und $(0, 3/2)$ ist.

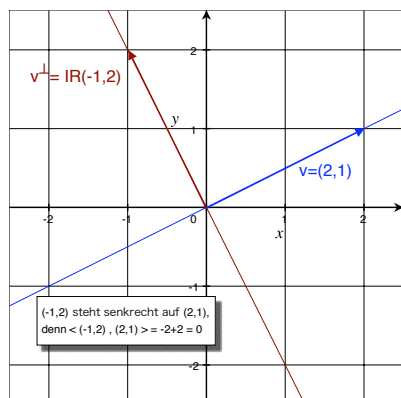


ABBILDUNG 4

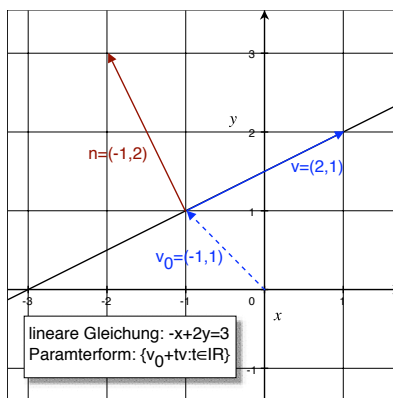


ABBILDUNG 5

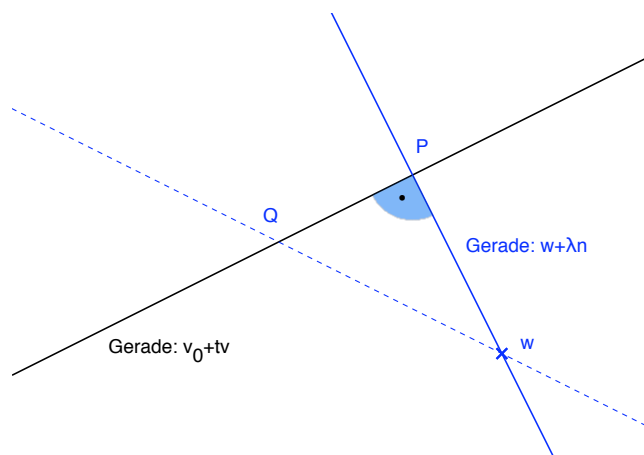


ABBILDUNG 6

6.4. Abstand eines Punktes zu einer Geraden. Der Abstand des Punktes $\vec{w}$ zur Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$ ist das Minimum der Abstände zwischen einem Punkt auf der Geraden und $\vec{v}$, also $\min\{d(\vec{w}, \vec{v}_0 + t\vec{v}) : t \in \mathbb{R}\}$.

Es seien $\vec{n}$ ein Normalenvektor der Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$, $\vec{P}$ der Schnittpunkt der Geraden $\{\vec{w} + \lambda\vec{n} : \lambda \in \mathbb{R}\}$ mit $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$ und $\vec{Q}$ ein beliebiger Punkt auf der Geraden $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$ (siehe Abbildung 6). Dann existieren $t, \lambda \in \mathbb{R}$ mit

$$\begin{aligned} d(\vec{Q}, \vec{w})^2 &= \|\vec{Q} - \vec{w}\|^2 = \|\vec{P} + t\vec{v} - \vec{w}\|^2 = \|\vec{P} - \vec{w} + t\vec{v}\|^2 \\ &= \langle \vec{P} - \vec{w} + t\vec{v}, \vec{P} - \vec{w} + t\vec{v} \rangle \\ &= \langle \vec{P} - \vec{w}, \vec{P} - \vec{w} \rangle + 2t\langle \vec{P} - \vec{w}, \vec{v} \rangle + t^2\langle \vec{v}, \vec{v} \rangle \\ &= \|\vec{P} - \vec{w}\|^2 + 2t\langle \lambda\vec{n}, \vec{v} \rangle + t^2\|\vec{v}\|^2 \geq d(\vec{P}, \vec{w})^2, \end{aligned}$$

da $\vec{n} \perp \vec{v}$ und $\|\vec{v}\|^2 \geq 0$. Der Punkt $\vec{P}$ heißt **Fußpunkt des Lotes** von $\vec{w}$ auf die Gerade $\{\vec{v}_0 + t\vec{v} : t \in \mathbb{R}\}$.

SATZ 101. Es seien $\vec{w}, \vec{n} \in \mathbb{R}^2$, $\vec{n} \neq \vec{0}$, $c \in \mathbb{R}$ und

$$(138) \quad \lambda := \frac{b - \langle \vec{w}, \vec{n} \rangle}{\|\vec{n}\|^2}.$$

Der Abstand des Punktes $\vec{w}$ von der Geraden $\langle \vec{n}, \vec{x} \rangle = b$ ist $|\lambda| \|\vec{n}\|$. Der Fußpunkt des Lotes von $\vec{w}$ auf die Gerade $\langle \vec{n}, \vec{x} \rangle = b$ ist $\vec{w} + \lambda \vec{n}$.

BEWEIS. Wenn $\vec{w} + \lambda \vec{n}$ auf der Geraden $\langle \vec{n}, \vec{x} \rangle = b$ liegt, dann gilt $b = \langle \vec{n}, \vec{w} + \lambda \vec{n} \rangle = \lambda \|\vec{n}\|^2 + \langle \vec{n}, \vec{w} \rangle$. $\square$

6.5. Spiegelung an einer Geraden. Die Spiegelung an der Geraden, die durch die lineare Gleichung $\langle \vec{n}, \vec{x} \rangle = b$ definiert ist, ist die Abbildung

$$(139) \quad \sigma : \mathbb{R}^2 \rightarrow \mathbb{R}^2, \quad \vec{x} \mapsto \vec{x} + 2 \frac{b - \langle \vec{x}, \vec{n} \rangle}{\|\vec{n}\|^2} \vec{n}.$$

Es gilt genau dann $\sigma(\vec{x}) = \vec{x}$, wenn $\langle \vec{n}, \vec{x} \rangle = b$. Die Spiegelung σ an der Geraden $\langle \vec{n}, \vec{x} \rangle = b$ bewegt die Punkte auf dieser Geraden nicht. Für alle $\vec{x}, \vec{y} \in \mathbb{R}^2$ gilt $d(\vec{x}, \vec{y}) = d(\sigma(\vec{x}), \sigma(\vec{y}))$, d.h. bei einer Spiegelung bleiben die Abstände zwischen Punkte unverändert.

7. Das Kreuzprodukt auf $\mathbb{R}^3$

7.1. Normalenvektor einer Ebene. Es seien $\vec{a}_1$ und $\vec{a}_2$ zwei Vektoren in $\mathbb{R}^3$ mit $\vec{a}_1 = (a_{11}, a_{12}, a_{13})$ und $\vec{a}_2 = (a_{21}, a_{22}, a_{23})$. Ein Vektor $\vec{x} = (x_1, x_2, x_3)$ steht genau dann senkrecht auf $\vec{a}_1$ und $\vec{a}_2$, wenn $\langle \vec{a}_1, \vec{x} \rangle = 0$ und $\langle \vec{a}_2, \vec{x} \rangle = 0$. Diese Bedingung ist ein lineares Gleichungssystem für die Unbekannten x_1, x_2, x_3 mit der erweiterten Koeffizientenmatrix

$$A = \left(\begin{array}{ccc|c} a_{11} & a_{12} & a_{13} & 0 \\ a_{21} & a_{22} & a_{23} & 0 \end{array} \right).$$

Die Lösungsmenge $L := \vec{a}_1^\perp \cap \vec{a}_2^\perp$ ist genau dann ein eindimensionaler Untervektorraum, wenn die Vektoren $\vec{a}_1$ und $\vec{a}_2$ eine Ebene aufspannen, also $\text{rang } A = 2$ gilt. Für jeden Vektor $\vec{x} \in L$ mit $\vec{x} \neq \vec{0}$ ist dann das Orthogonalkomplement zu $\vec{x}$ gerade die von $\vec{a}_1$ und $\vec{a}_2$ aufgespannte Ebene, also $\vec{x}^\perp = \text{span}\{\vec{a}_1, \vec{a}_2\}$.

BEISPIEL 196. Die Vektoren $\vec{x} \in \mathbb{R}^3$, die senkrecht auf $\vec{a}_1 = (1, 2, 1)$ und $\vec{a}_2 = (2, 6, 0)$ stehen sind die Lösungen des linearen Gleichungssystems

$$\left(\begin{array}{ccc|c} 1 & 2 & 1 & 0 \\ 2 & 6 & 0 & 0 \end{array} \right) \sim \left(\begin{array}{ccc|c} 1 & 2 & 1 & 0 \\ 0 & 2 & -2 & 0 \end{array} \right) \sim \left(\begin{array}{ccc|c} 1 & 0 & 3 & 0 \\ 0 & 1 & -1 & 0 \end{array} \right),$$

also gerade die Elemente der Menge $\{(-3t, t, t) : t \in \mathbb{R}\} = \mathbb{R}(-3, 1, 1)$.

7.2. Eigenschaften des Kreuzproduktes. Das Kreuzprodukt $\vec{a} \times \vec{b}$ zweier Vektoren $\vec{a}, \vec{b} \in \mathbb{R}^3$ ist ein Vektor in $\mathbb{R}^3$ mit folgenden Eigenschaften:

- Der Vektor $\vec{a} \times \vec{b}$ steht *senkrecht* auf $\vec{a}$ und $\vec{b}$.

$$(140) \quad \langle \vec{a} \times \vec{b}, \vec{a} \rangle = 0 \text{ für beliebige } \vec{a}, \vec{b} \in \mathbb{R}^3$$

- Die *Länge* des Vektors $\vec{a} \times \vec{b}$ ist der Flächeninhalt des von den Vektoren $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms

$$P(\vec{a}, \vec{b}) = \{t\vec{a} + s\vec{b} : t, s \in [0, 1]\}.$$

- Der Vektor $\vec{a} \times \vec{b}$ ist so *orientiert*, dass $\vec{a}, \vec{b}$ und $\vec{a} \times \vec{b}$ ein Rechtssystem bilden.

Da je zwei Basisvektoren $\vec{e}_j \in \mathbb{R}^3$ der Standardorthonormalbasis ein Quadrat mit Seitenlänge 1 erzeugen, gilt

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

BEISPIEL 197. Liegen die Vektoren $\vec{a}$ und $\vec{b}$ in der x_1, x_2 -Ebene, also $\vec{a} = (a_1, a_2, 0)$ und $\vec{b} = (b_1, b_2, 0)$, so gilt $\text{span}(0, 0, 1) = \vec{a}^\perp \cap \vec{b}^\perp$, das Kreuzprodukt muss also ein Vielfaches des Vektors $(0, 0, 1)$ sein. Wir berechnen den Flächeninhalt $|P|$ des Parallelogramms $P = \{(ta_1 + sb_1, ta_2 + sb_2) : t, s \in [0, 1]\} \subset \mathbb{R}^2$, das in Abbildung 7 blau gefärbt ist. Vom Rechteck mit den Seitenlängen $a_1 + b_1$ und $a_2 + b_2$ muss jeweils das Doppelte der Flächeninhalte des gelben und roten Dreiecks und des grünen Rechtecks abgezogen werden. Es gilt

$$\begin{aligned} |P| &= (a_1 + b_1)(a_2 + b_2) - 2b_1a_2 - 2\frac{1}{2}a_1a_2 - 2\frac{1}{2}b_1b_2 \\ &= a_1a_2 + a_1b_2 + b_1a_2 + b_1b_2 - 2b_1a_2 - a_1a_2 - b_1b_2 = a_1b_2 - a_2b_1. \end{aligned}$$

Die verträgliche Orientierung des Kreuzproduktes ist $\vec{a} \times \vec{b} = (0, 0, a_1b_2 - b_1a_2)$, da man für $a_1 = b_2 = 1$ und $a_2 = b_1 = 1$ das Kreuzprodukt $\vec{a} \times \vec{b} = \vec{e}_1 \times \vec{e}_2 = \vec{e}_3$ erhalten sollte.

Das Kreuzprodukt ist *antisymmetrisch*, d.h. Vertauschung der Faktoren kehrt das Vorzeichen (die Orientierung) um, und *bilinear*, d.h. die Bildung des Kreuzproduktes ist mit Vektoraddition und Skalarmultiplikation in einem Faktor vertauschbar. Für beliebige $\vec{a}, \vec{b}, \vec{c} \in \mathbb{R}^3$ und $\lambda \in \mathbb{R}$ gilt:

$$(141) \quad \vec{a} \times \vec{b} = -\vec{b} \times \vec{a}, \vec{a} \times (\lambda\vec{b}) = \lambda(\vec{a} \times \vec{b}), \vec{a} \times (\vec{b} + \vec{c}) = \vec{a} \times \vec{b} + \vec{a} \times \vec{c}$$

7.3. Explizite Formel des Kreuzproduktes. Das Kreuzprodukt ist eine Abbildung $\mathbb{R}^3 \times \mathbb{R}^3 \times \mathbb{R}^3, (\vec{a}, \vec{b}) \mapsto \vec{a} \times \vec{b}$, mit

$$(142) \quad \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} \times \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} = \begin{pmatrix} a_2b_3 - b_2a_3 \\ a_3b_1 - b_3a_1 \\ a_1b_2 - b_1a_2 \end{pmatrix} \text{ für } \vec{a} = \begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} \text{ und } \vec{b} = \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$$

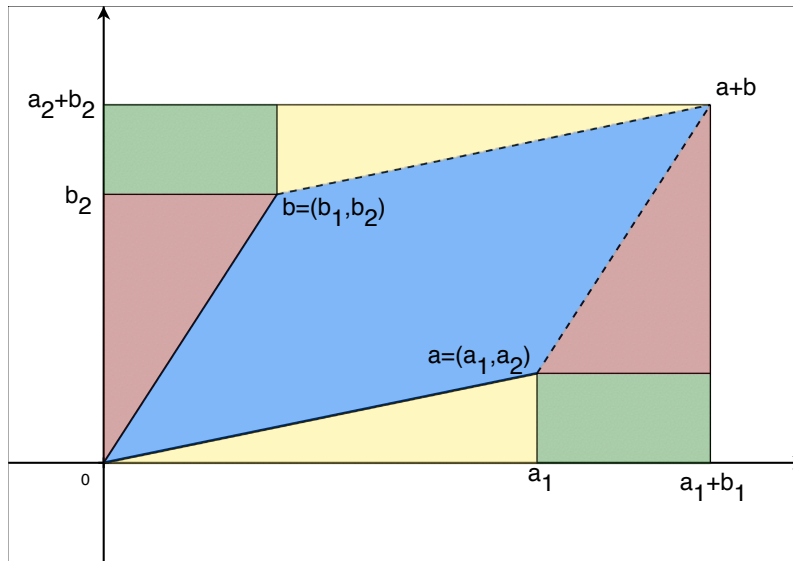


ABBILDUNG 7

BEISPIEL 198. Für $\vec{a} = (1, 2, 1)$ und $\vec{b} = (2, 6, 0)$ erhält man mit der expliziten Formel 142

$$\begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix} \times \begin{pmatrix} 2 \\ 6 \\ 0 \end{pmatrix} = \begin{pmatrix} 2 \cdot 0 - 6 \cdot 1 \\ 1 \cdot 2 - 0 \cdot 1 \\ 1 \cdot 6 - 2 \cdot 2 \end{pmatrix} = \begin{pmatrix} -6 \\ 2 \\ 2 \end{pmatrix}.$$

Der Vektor $(-6, 2, 2)$ ist ein Vielfaches des Vektors $(-3, 1, 1)$ (siehe Beispiel 196).

8. Lösungsmenge einer linearen Differentialgleichung

Die Exponentialfunktion $x \mapsto e^x$ ist die Funktion $y : \mathbb{R} \rightarrow \mathbb{R}$ mit den Eigenschaften $y'(x) = y(x)$ für alle $x \in \mathbb{R}$ und $y(0) = 1$. Die Exponentialfunktion ist also eine Lösung der Differentialgleichung $y' = y$ mit der Anfangsbedingung $y(0) = 1$.

Wie viele differenzierbare Funktionen $y : \mathbb{R} \rightarrow \mathbb{R}$ mit der Eigenschaft $y'(x) = y(x)$ für alle $x \in \mathbb{R}$ gibt es? Wenn $y'(x) = y(x)$ für alle $x \in \mathbb{R}$, dann ist die Funktion $h(x) := y(x)e^{-x}$ konstant, denn für alle $x \in \mathbb{R}$ gilt

$$h'(x) = y'(x)e^{-x} - y(x)e^{-x} = (y'(x) - y(x))e^{-x} = 0.$$

Die Lösungsmenge der Differentialgleichung $y' = y$ ist $\{ce^x : c \in \mathbb{R}\}$. Dies ist der eindimensionale Untervektorraum im Vektorraum aller differenzierbaren Funktionen, der vom Vektor e^x aufgespannt wird.

8.1. Homogene, lineare Differentialgleichungen. Eine homogene, lineare Differentialgleichung n -ter Ordnung ist eine Gleichung der Form

$$(143) \quad a_n(x)y^{(n)} + a_{n-1}(x)y^{(n-1)} + \dots + a_1(x)y' + a_0(x)y = 0,$$

wobei $a_0, \dots, a_n : (a, b) \rightarrow \mathbb{R}$ auf dem Intervall (a, b) stetige Funktionen mit $a_n \neq 0$ sind. Eine **Lösung** der Differentialgleichung 143 ist eine Funktion $y : (a, b) \rightarrow \mathbb{R}$ mit der Eigenschaft $\sum_{j=0}^n a_j(x)y^{(j)}(x) = 0$ für alle $x \in \mathbb{R}$.

SATZ 102. *Die Lösungen einer homogenen, linearen Differentialgleichung bilden einen Vektorraum.*

BEWEIS. Wenn y_1 und y_2 Lösungen der homogenen, linearen Differentialgleichung $\sum_{j=0}^n a_j(x)y^{(j)} = 0$ sind und $c \in \mathbb{R}$, dann gilt $(cy_1)^{(j)}(x) = cy_1^{(j)}(x)$ und $(y_1 + y_2)^{(j)}(x) = y_1^{(j)}(x) + y_2^{(j)}(x)$, also

$$\begin{aligned} \sum_{j=0}^n a_j(x)(y_1 + y_2)^{(j)}(x) &= \sum_{j=0}^n a_j(x)(y_1^{(j)}(x) + y_2^{(j)}(x)) \\ &= \sum_{j=0}^n a_j(x)y_1^{(j)}(x) + \sum_{j=0}^n a_j(x)y_2^{(j)}(x) = 0 \\ \sum_{j=0}^n a_j(x)(cy_1)^{(j)}(x) &= \sum_{j=0}^n a_j(x)cy_1^{(j)}(x) = c \sum_{j=0}^n a_j(x)y_1^{(j)}(x) = 0. \end{aligned}$$

Auch die Funktionen $y_1 + y_2$ und cy_1 sind Lösungen der homogenen linearen Differentialgleichung. $\square$

Gibt man zusätzlich zu der Differentialgleichung 143 den Funktionswert einer Lösungsfunktion $y : (a, b) \rightarrow \mathbb{R}$ und ihrer Ableitungen $y^{(j)}$, $j = 1, \dots, n-1$, die zum Zeitpunkt x_0 vor, so spricht man von einer Differentialgleichung mit Anfangsbedingungen oder einer **Anfangswertaufgabe**. Zu gegebenen $\xi \in (a, b)$ und $\eta_0, \dots, \eta_{n-1} \in \mathbb{R}$ ist eine Lösung y der homogenen, linearen Differentialgleichung 143 eine Lösung der Anfangswertaufgabe, falls $y^{(j)}(\xi) = \eta_j$ für alle $j = 0, \dots, n-1$.

BEMERKUNG 81. Falls die Funktion eine Bewegung eines Körpers beschreibt, so sind $y(\xi)$ die Anfangsposition, $y'(\xi)$ die Anfangsgeschwindigkeit, $y''(\xi)$ die Anfangsbeschleunigung usw. $\square$

SATZ 103. *Jede Anfangswertaufgabe zu einer linearen Differentialgleichung n -ter Ordnung ist eindeutig lösbar.*

SATZ 104. *Die Lösungen einer homogenen, linearen Differentialgleichung n -ter Ordnung bilden einen n -dimensionalen Vektorraum.*

Für jedes $x \in (a, b)$ bilden eindeutigen Lösungen y_k , $k = 0, \dots, n-1$, der Anfangswertaufgaben $y_k^{(j)}(\xi) = 0$ für $k \neq j$ und $y_k^{(k)}(\xi) = 1$ eine Basis dieses Vektorraumes.

8.2. Inhomogene, lineare Differentialgleichungen. Eine inhomogene, lineare Differentialgleichung n -ter Ordnung ist eine Gleichung der Form

$$(144) \quad a_n(x)y^{(n)} + a_{n-1}(x)y^{(n-1)} + \dots + a_1(x)y' + a_0(x)y = f(x),$$

wobei $f, a_0, \dots, a_n : (a, b) \rightarrow \mathbb{R}$ auf dem Intervall (a, b) stetige Funktionen mit $a_n \neq 0$ sind. Eine **Lösung** der Differentialgleichung 144 ist eine Funktion $y : (a, b) \rightarrow \mathbb{R}$ mit der Eigenschaft $\sum_{j=0}^n a_j(x)y^{(j)}(x) = f(x)$ für alle $x \in \mathbb{R}$.

SATZ 105. *Wenn y_1 und y_2 Lösungen der inhomogenen, linearen Differentialgleichung $\sum_{j=0}^n a_j(x)y^{(j)} = f(x)$ sind, dann ist $y_1 - y_2$ Lösung der homogenen, linearen Differentialgleichung $\sum_{j=0}^n a_j(x)y^{(j)} = 0$.*

BEWEIS. Da $(y_1 - y_2)^{(j)} = y_1^{(j)} - y_2^{(j)}$, folgt

$$\sum_{j=0}^n a_j(x)(y_1 - y_2)^{(j)} = \sum_{j=0}^n a_j(x)y_1^{(j)} - \sum_{j=0}^n a_j(x)y_2^{(j)} = f(x) - f(x) = 0.$$

□

Ist y_1 eine Lösung der inhomogenen Differentialgleichung $\sum_{j=0}^n a_j(x)y^{(j)}(x) = f(x)$, so ist die Menge aller Lösungen dieser inhomogenen Differentialgleichung

$$y_1 + \left\{ y_0 : \sum_{j=0}^n a_j(x)y_0^{(j)} = 0 \right\}.$$

KAPITEL XIII

Lineare Abbildungen

Es seien V und W Vektorräume. Eine Abbildung $T : V \rightarrow W$ heißt **linear**, falls $T(v + v') = T(v) + T(v')$ und $T(\lambda v) = \lambda T(v)$ für alle $v, v' \in V$ und alle $\lambda \in \mathbb{R}$. Insbesondere gilt immer $T(\vec{0}) = \vec{0}$.

Wählt man eine Basis $\{\vec{v}_1, \dots, \vec{v}_n\}$ des Vektorraumes V so gilt für jeden Vektor $\vec{v} \in V$ mit $\vec{v} = \alpha_1 \vec{v}_1 + \dots + \alpha_n \vec{v}_n$

$$T(\vec{v}) = T(\alpha_1 \vec{v}_1 + \dots + \alpha_n \vec{v}_n) = \alpha_1 T(\vec{v}_1) + \dots + \alpha_n T(\vec{v}_n).$$

Es genügt also, die Bilder der Basisvektoren anzugeben, um eine lineare Abbildung zu beschreiben. Wählt man auch eine Basis $\{\vec{w}_1, \dots, \vec{w}_m\}$ des Vektorraumes W , so kann man jeden Vektor $T(\vec{v}_j)$ als Linearkombination der Vektoren $\vec{w}_1, \dots, \vec{w}_m$ schreiben, $T(\vec{v}_j) = a_{1j} \vec{w}_1 + \dots + a_{mj} \vec{w}_m$.

Die j -te Spalte der **Matrix** A der linearen Abbildung T bezüglich den Basen $\{\vec{v}_1, \dots, \vec{v}_n\} \subset V$ und $\{\vec{w}_1, \dots, \vec{w}_m\} \subset W$ sind die Koordinaten des Vektors $T(\vec{v}_j)$ bezüglich der Basis $\{\vec{w}_1, \dots, \vec{w}_m\}$, also

$$(145) \quad A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}$$

BEMERKUNG 82. Die Matrix einer linearen Abbildung von einem n -dimensionalen Vektorraum in einen m -dimensionalen Vektorraum besteht aus m Zeilen und n Spalten. $\square$

BEISPIEL 199 (Identität). Für jeden Vektorraum V und bezüglich jeder Basis ist die Matrix der linearen Abbildung $V \rightarrow V$, $v \mapsto v$, die Einheitsmatrix

$$E := \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{pmatrix} \quad \text{mit } e_{ij} = \begin{cases} 1 & (i = j) \\ 0 & (i \neq j) \end{cases}.$$

BEISPIEL 200 (Spiegelung an der x_1 -Achse). Die Abbildung $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $\vec{x} = (x_1, x_2) \mapsto (x_1, -x_2)$, beschreibt die Spiegelung an der x_1 -Achse im $\mathbb{R}^2$. Da $T(\vec{e}_1) = \vec{e}_1$ und $T(\vec{e}_2) = -\vec{e}_2$ ist die Matrix dieser Abbildung bezüglich der Standardbasis

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

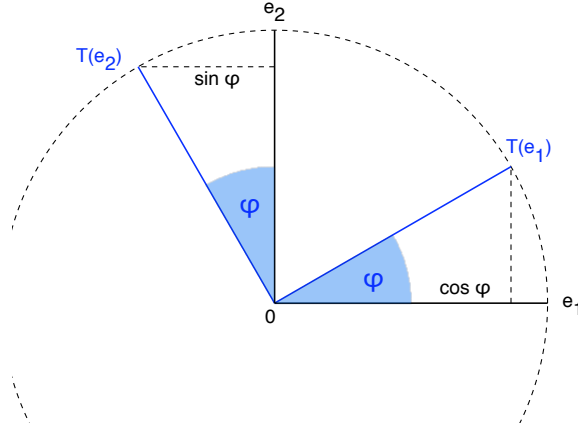


ABBILDUNG 1

Identifiziert man $\mathbb{R}^2$ mit $\mathbb{C}$, so entspricht die Spiegelung an der x_1 -Achse der Spiegelung an der reellen Achse, also der komplexen Konjugation.

BEISPIEL 201 (Spiegelung an der x_2 -Achse). Die Abbildung $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, $\vec{x} = (x_1, x_2) \mapsto (-x_1, x_2)$, beschreibt die Spiegelung an der x_2 -Achse im $\mathbb{R}^2$. Da $T(\vec{e}_1) = -\vec{e}_1$ und $T(\vec{e}_2) = \vec{e}_2$ ist die Matrix dieser Abbildung bezüglich der Standardbasis

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

BEISPIEL 202 (Drehung um den Winkel φ). Es sei $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Drehung man die Ebene im Koordinatenursprung um den Winkel φ . Die ist eine lineare Abbildung. Da $T(\vec{e}_1) = T(1, 0) = (\cos \varphi, \sin \varphi)$ und $T(\vec{e}_2) = T(0, 1) = (-\sin \varphi, \cos \varphi)$ (siehe Abbildung 1) gilt, ist die Matrix der Drehung bezüglich der Standardbasis

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

Interpretiert man $\mathbb{R}^2$ als $\mathbb{C}$, so entspricht die Drehung um den Winkel φ der Multiplikation mit $e^{i\varphi}$.

BEISPIEL 203. Für jeden Richtungsvektor $\vec{v} \in \mathbb{R}^2$ mit $\vec{v} = (v_1, v_2) \neq \vec{0}$ ist die Spiegelung S an der Geraden $\{t\vec{v} : t \in \mathbb{R}\}$ eine lineare Abbildung, denn ein Normalenvektor ist $\vec{n} = (-v_2, v_1)$ und aus Formel 139 folgt

$$\begin{aligned} S(\vec{x}) &= \vec{x} - 2 \frac{\langle \vec{x}, \vec{n} \rangle}{\|\vec{n}\|^2} \vec{n} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} - 2 \frac{-v_2 x_1 + v_1 x_2}{v_1^2 + v_2^2} \begin{pmatrix} -v_2 \\ v_1 \end{pmatrix} \\ &= \frac{1}{v_1^2 + v_2^2} \begin{pmatrix} (v_1^2 - v_2^2)x_1 - 2v_1 v_2 x_2 \\ 2v_1 v_2 x_1 + (-v_1^2 + v_2^2)x_2 \end{pmatrix}. \end{aligned}$$

Für die Gerade mit dem Richtungsvektor $\vec{v} = (2, 1)$ (siehe Abbildung 4) ist die Spiegelung definiert als

$$\begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \mapsto \frac{1}{5} \begin{pmatrix} 3x_1 - 4x_2 \\ 4x_1 - 3x_2 \end{pmatrix} = \begin{pmatrix} \frac{3}{5}x_1 - \frac{4}{5}x_2 \\ \frac{4}{5}x_1 - \frac{3}{5}x_2 \end{pmatrix}.$$

Die Matrix der Spiegelung bezüglich der Standardbasis ist

$$A = \begin{pmatrix} \frac{3}{5} & -\frac{4}{5} \\ \frac{4}{5} & -\frac{3}{5} \end{pmatrix} = \frac{1}{5} \begin{pmatrix} 3 & -4 \\ 4 & -3 \end{pmatrix}.$$

Auch der Richtungsvektor $\vec{v}$ und der Normalenvektor $\vec{n}$ bilden eine Basis des $\mathbb{R}^2$. Bezüglich der Basis $\{\vec{v}, \vec{n}\}$ besitzt die Spiegelung die (einfache) Matrix

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

BEISPIEL 204 (Projektion von $\mathbb{R}^3$ auf $\mathbb{R}^2$). Die Abbildung $\mathbb{R}^3 \rightarrow \mathbb{R}^2$ mit $(x_1, x_2, x_3) \mapsto (x_1, x_2)$ ist eine lineare Abbildung. Sie vergisst die dritte Koordinate. Ihre Matrix bezüglich der Standardbasen in $\mathbb{R}^3$ und $\mathbb{R}^2$ ist

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

denn $\vec{e}_1 = (1, 0, 0) \mapsto (1, 0) = \vec{e}_1$, $\vec{e}_2 = (0, 1, 0) \mapsto (0, 1) = \vec{e}_2$ und $\vec{e}_3 = (0, 0, 1) \mapsto (0, 0) = \vec{0}$.

BEISPIEL 205 (Spiegelung an der x_1, x_2 -Ebene in $\mathbb{R}^3$). Die Matrix der linearen Abbildung $\mathbb{R}^3 \rightarrow \mathbb{R}^3$, $(x_1, x_2, x_3) \mapsto (x_1, x_2, -x_3)$, bezüglich der Standardbasis ist

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

1. Matrizenmultiplikation

Es seien $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$ und $S : \mathbb{R}^m \rightarrow \mathbb{R}^k$ lineare Abbildungen mit den Matrizen $A = (a_{ij})$ bzw. $B = (b_{ki})$ (bezüglich der Standardbasen). Die Verkettung der $S \circ T$ der linearen Abbildungen T und S ist wieder eine lineare Abbildung, denn für beliebige $\vec{x}, \vec{y} \in \mathbb{R}^n$ und $\lambda \in \mathbb{R}$ gilt

$$(S \circ T)(\lambda \vec{x}) = S(T(\lambda \vec{x})) = S(\lambda(T(\vec{x}))) = \lambda S(T(\vec{x})) = \lambda(S \circ T)(\vec{x})$$

und

$$\begin{aligned} (S \circ T)(\vec{x} + \vec{y}) &= S(T(\vec{x} + \vec{y})) = S(T(\vec{x}) + T(\vec{y})) = S(T(\vec{x})) + S(T(\vec{y})) \\ &= (S \circ T)(\vec{x}) + (S \circ T)(\vec{y}). \end{aligned}$$

Die Spalten der Matrix $C = (c_{lj})$ der linearen Abbildung $S \circ T$ sind die Bilder der Standardbasisvektoren, also $(S \circ T)(\vec{e}_1), \dots, (S \circ T)(\vec{e}_n)$. Es gilt

$$\begin{aligned} S(T(\vec{e}_j)) &= S\left(\sum_{i=1}^m a_{ij}\vec{e}_i\right) = \sum_{i=1}^m a_{ij}S(\vec{e}_i) = \sum_{i=1}^m a_{ij}\left(\sum_{l=1}^k b_{li}\vec{e}_l\right) \\ &= \sum_{i=1}^m \sum_{l=1}^k (a_{ij}b_{li}\vec{e}_l) = \sum_{i=1}^m \left(\sum_{l=1}^k b_{li}a_{ij}\right)\vec{e}_l. \end{aligned}$$

Sind $A = (a_{ij})$ eine $(m \times n)$ -Matrix und $B = (b_{li})$ eine $(k \times m)$ -Matrix, so ist die **Produktmatrix** $C = (c_{lj})$ eine $(k \times n)$ -Matrix mit den Einträgen

$$c_{lj} := \sum_{i=1}^m b_{li}a_{ij} \text{ für } l = 1, \dots, k, j = 1, \dots, n.$$

Der Eintrag c_{lj} der Matrix C ist das Skalarprodukt der l -ten Zeile der Matrix B mit der j -ten Spalte der Matrix A .

BEISPIEL 206 (Hintereinanderausführung zweier Drehungen). Die Matrizen

$$A = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}, \quad B = \begin{pmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{pmatrix}$$

stellen bezüglich der Standardbasis des $\mathbb{R}^2$ Drehungen um die Winkel α bzw. β dar (siehe Beispiel 202). Es gilt

$$\begin{aligned} BA &= \begin{pmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{pmatrix} \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \\ &= \begin{pmatrix} \cos \beta \cos \alpha + (-\sin \beta) \sin \alpha & \cos \beta (-\sin \alpha) - \sin \beta \cos \alpha \\ \sin \beta \cos \alpha + \cos \beta \sin \alpha & \sin \beta (-\sin \alpha) + \cos \beta \cos \alpha \end{pmatrix} \\ &= \begin{pmatrix} \cos \beta \cos \alpha - \sin \beta \sin \alpha & -(\cos \beta \sin \alpha + \sin \beta \cos \alpha) \\ \sin \beta \cos \alpha + \cos \beta \sin \alpha & -\sin \beta \sin \alpha + \cos \beta \cos \alpha \end{pmatrix} \\ &= \begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix}. \end{aligned}$$

Führt man erst eine Drehung um den Winkel α und dann eine Drehung um den Winkel β durch, so erhält man eine Drehung um den Winkel $\alpha + \beta$.

BEMERKUNG 83. Im Beispiel 206 gilt $AB = BA$. Dies liegt daran, dass die Drehungen im $\mathbb{R}^2$ als Teilmenge der komplexen Zahlen $\{e^{i\varphi} : \varphi \in \mathbb{R}\}$ angesehen werden können und die komplexen Zahlen ein Körper sind. Im Allgemeinen ist die Matrixmultiplikation nicht kommutativ, d.h. $AB \neq BA$. Falls $k \neq n$, so ist das Produkt AB sogar nicht definiert. $\square$

2. Die Determinante

Es seien $\vec{a}_1, \dots, \vec{a}_n \in \mathbb{R}^n$. Die Determinante misst das orientierte Volumen des von den Vektoren $\vec{a}_1, \dots, \vec{a}_n$ aufgespannten *Parallelepipeds*

$$P(\vec{a}_1, \dots, \vec{a}_n) := \{t_1\vec{a}_1 + \dots + t_n\vec{a}_n : t_j \in [0, 1]\}.$$

Die Determinante $\det A$ einer (quadratischen) $n \times n$ -Matrix $A = (a_{ij})$ ist die Determinante der n -Spalten der Matrix A :

$$(146) \quad \det A := \det(\vec{a}_1, \dots, \vec{a}_n) := \text{vol}(P(\vec{a}_1, \dots, \vec{a}_n)) \text{ mit } \vec{a}_j = \begin{pmatrix} a_{1j} \\ \vdots \\ a_{nj} \end{pmatrix}$$

Für $n = 1$ und $A = (a)$ gilt $\det A = \det(a) = a$. Die Determinante misst die Länge des Vektors $a\vec{e}_1$ und ob er in Richtung des Basisvektors $\vec{e}_1$ oder die entgegengesetzte Richtung zeigt.

Für $n = 2$ ist das Parallelepiped $P(\vec{a}_1, \vec{a}_2)$ ein Parallelogramm und aus Beispiel 197 folgt

$$(147) \quad \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = a_{11}a_{22} - a_{21}a_{12}.$$

2.1. Determinantenberechnung durch Anwendung des Gaußverfahrens auf Zeilen und Spalten. Für die Standardbasisvektoren $\vec{e}_1, \dots, \vec{e}_n$ ist das Parallelepiped $P(\vec{e}_1, \dots, \vec{e}_n)$ ein n -dimensionaler Würfel der Seitenlänge 1 und Volumen ist $1^n = 1$. Daraus folgt

$$\det E = \det(\vec{e}_1, \dots, \vec{e}_n) = 1 \text{ für alle } n \in \mathbb{N}.$$

Für beliebige $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ ist das Parallelepiped $P(\lambda_1\vec{e}_1, \dots, \lambda_n\vec{e}_n)$ ein n -dimensionaler Quader mit den Seitenlängen $\lambda_1, \dots, \lambda_n$ und Volumen $\lambda_1 \cdots \lambda_n$, also

$$\det \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{pmatrix} = \det(\lambda_1\vec{e}_1, \lambda_2\vec{e}_2, \dots, \lambda_n\vec{e}_n) = \lambda_1 \cdots \lambda_n.$$

Mit Hilfe des Gaußverfahrens kann man eine Matrix A erst in Stufenform bringen und dann, falls $\text{rang } A = n$ sogar in eine Diagonalmatrix verwandeln. Bei den Zeilenoperationen des Gaußverfahrens (siehe Satz 96) ändert sich die Determinante folgendermaßen:

- Vertauschung zweier Zeilen ändert das Vorzeichen (die Orientierung) der Determinante, für $k \neq l$ gilt

$$(148) \quad \det(\vec{a}_1, \dots, \vec{a}_k, \dots, \vec{a}_l, \dots, \vec{a}_n) = (-1) \det(\vec{a}_1, \dots, \vec{a}_l, \dots, \vec{a}_k, \dots, \vec{a}_n).$$

- Multiplikation einer Zeile mit einer Zahl $\lambda \in \mathbb{R}$ vervielfacht auch die Determinante (das Volumen) um diesen Faktor, für beliebige $\lambda_1, \dots, \lambda_n \in \mathbb{R}$ gilt

$$(149) \quad \det(\lambda_1\vec{a}_1, \dots, \lambda_n\vec{a}_n) = \lambda_1 \cdots \lambda_n \det(\vec{a}_1, \dots, \vec{a}_n).$$

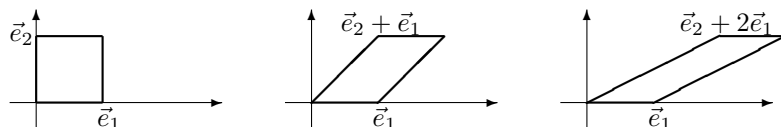
- Addition eines Vielfachen einer Zeile zu einer anderen Zeile ändert die Determinante nicht, für beliebige $\lambda \in \mathbb{R}$ und $k \neq l$ gilt

$$(150) \quad \det(\vec{a}_1, \dots, \vec{a}_k, \dots, \vec{a}_l, \dots, \vec{a}_n) = \det(\vec{a}_1, \dots, \vec{a}_k, \dots, \vec{a}_l + \lambda\vec{a}_k, \dots, \vec{a}_n).$$

Betrachtet man die Zeilenvektoren $\vec{a}_i$ mit $\vec{a}_i = (a_{i1}, \dots, a_{in})$ der Matrix A , so gilt $\det A = \det(\vec{a}_1, \dots, \vec{a}_n)$, d.h. das Volumen des von den Spaltenvektoren aufgespannten Parallelepipeds ist gleich dem Volumen des von den Zeilenvektoren aufgespannten Parallelepipeds.

Vertauscht man in der Matrix A zwei Zeilenvektoren, oder multipliziert einen Zeilenvektor mit einer Zahl λ oder addiert das Vielfache einer Zeile zu einer anderen, so ändert sich die Determinante auch bei diesen Zeilenoperationen wie in den Gleichungen 148, 149 und 150 beschrieben.

BEISPIEL 207. Es gilt $\det(\vec{e}_1, \vec{e}_2) = \det(\vec{e}_1, \vec{e}_2 + \vec{e}_1) = \det(\vec{e}_1, \vec{e}_2 + 2\vec{e}_1) = 1$, denn die Parallelogramme $P(\vec{e}_1, \vec{e}_2)$, $P(\vec{e}_1, \vec{e}_2 + \vec{e}_1)$ und $P(\vec{e}_1, \vec{e}_2 + 2\vec{e}_1)$ haben alle eine Grundseite ($\vec{e}_1$) der Länge 1 und Höhe 1, denn der Abstand des Punktes $\vec{e}_2 + \lambda\vec{e}_1$ zur Gerade $\mathbb{R}\vec{e}_1$ ist für alle $\lambda \in \mathbb{R}$ gleich.



BEISPIEL 208. Es gilt

$$\begin{aligned} \det \begin{pmatrix} 1 & 3 & 0 \\ 2 & -1 & 1 \\ 1 & 2 & 3 \end{pmatrix} &= \begin{vmatrix} 1 & 3 & 0 \\ 2 & -1 & 1 \\ 1 & 2 & 3 \end{vmatrix} = \begin{vmatrix} 1 & 3 & 0 \\ 0 & -7 & 1 \\ 0 & -1 & 3 \end{vmatrix} = (-1) \begin{vmatrix} 1 & 3 & 0 \\ 0 & -1 & 3 \\ 0 & -7 & 1 \end{vmatrix} \\ &= (-1)^2 \begin{vmatrix} 1 & 3 & 0 \\ 0 & 1 & -3 \\ 0 & -7 & 1 \end{vmatrix} = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & -3 \\ 0 & -7 & 1 \end{vmatrix} \\ &= \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & -3 \\ 0 & 0 & -20 \end{vmatrix} = \begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -20 \end{vmatrix} = -20. \end{aligned}$$

2.2. Explizite Formel der Determinante. Die Abbildung

$$\det : (\mathbb{R}^n)^n = \mathbb{R}^n \times \dots \times \mathbb{R}^n \rightarrow \mathbb{R}$$

ist eine alternierende Multilinearform. Durch die Normierung $\det(\vec{e}_1, \dots, \vec{e}_n) = 1$ ist diese Multilinearform eindeutig bestimmt.

Mit Hilfe der Permutationen der n Elemente $(1 \dots n)$ (siehe Abschnitt 6.1) kann man eine explizite Formel für die Determinante einer Matrix aufstellen. Die Menge aller Permutationen von n Elementen wird mit $\mathcal{S}_n$ bezeichnet.

Jede Permutation läßt sich durch Hintereinanderausführung von Vertauschungen genau zweier Elemente erzeugen. Das **Signum** $\text{sgn } \pi$ ist 1, wenn die Permutation π durch eine gerade Anzahl von Vertauschungen entsteht, und -1 , wenn man eine ungerade Anzahl von Vertauschungen benötigt, um π zu erzeugen. Es gilt die **Leibnizformel**

$$(151) \quad \det A = \sum_{\pi \in \mathcal{S}_n} \text{sgn } \pi a_{1\pi(1)} a_{2\pi(2)} \dots a_{n\pi(n)}$$

BEMERKUNG 84. Die Menge $\mathcal{S}_n$ besteht aus $n!$ Elementen. In der Formel 151 tauchen damit $n!$ Summanden auf. Dies bedeutet, dass sie sich für große n nicht zur Berechnung der Determinante eignet. $\square$

Für $n = 2$ hat die Menge $\mathcal{S}_2$ genau zwei Elemente $\mathcal{S}_2 = \{(12), (21)\}$. Es gilt $\text{sgn}(12) = 1$, denn man braucht keine Vertauschung, und $\text{sgn}(21) = -1$, denn man vertauscht die Elemente 1 und 2 einmal miteinander. Daraus folgt

$$(152) \quad \det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}.$$

Für $n = 3$ hat die Menge $\mathcal{S}_3$ genau sechs Elemente

$$\mathcal{S}_3 = \{(123), (213), (321), (132), (231), (312)\}.$$

Es gilt $\text{sgn}(123) = 1$, $\text{sgn}(213) = \text{sgn}(321) = \text{sgn}(132) = -1$ und $\text{sgn}(231) = \text{sgn}(312) = 1$. Daraus folgt

$$(153) \quad \det A = \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} \\ = a_{11}a_{22}a_{33} - a_{12}a_{21}a_{33} - a_{13}a_{22}a_{31} - a_{11}a_{23}a_{32} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32}.$$

2.3. Rekursionsformel, Laplacescher Entwicklungssatz. Zerlegt man in Formel 151 die Menge $\mathcal{S}_n$ für einen festen Index i in die Teilmengen $M_j := \{\pi : \pi(i) = j\}$ für $j = 1, \dots, n$, so kann die Mengen M_j mit $\mathcal{S}_{n-1}$ identifizieren und erhält die Rekursionsformel

$$(154) \quad \det A = \sum_{j=1}^n (-1)^{i+j} a_{ij} \det A_{ij} = \sum_{i=1}^n (-1)^{i+j} a_{ij} \det A_{ij},$$

wobei die $(n-1) \times (n-1)$ -Matrix A_{ij} aus der Matrix A durch Streichung der i -ten Zeile und der j -ten Spalte entsteht. Die Gleichung 154 heißt auch Entwicklung der Determinante nach der i -ten Zeile bzw. der j -ten Spalte.

BEISPIEL 209. Die Entwicklung einer 2×2 -Matrix nach der ersten Zeile liefert wegen $A_{11} = (a_{22})$ und $A_{12} = (a_{21})$

$$\det \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = a_{11}a_{22} - a_{12}a_{21}.$$

BEISPIEL 210. Die Entwicklung einer 3×3 -Matrix nach der dritten Spalte liefert wegen

$$A_{13} = \begin{pmatrix} a_{21} & a_{22} \\ a_{31} & a_{32} \end{pmatrix}, A_{23} = \begin{pmatrix} a_{11} & a_{12} \\ a_{31} & a_{32} \end{pmatrix} \text{ und } A_{33} = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix}$$

die Formel

$$\begin{aligned}
 \det \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{pmatrix} &= a_{13} \det A_{13} - a_{23} \det A_{23} + a_{33} \det A_{33} \\
 &= a_{13}(a_{21}a_{32} - a_{31}a_{22}) - a_{23}(a_{11}a_{32} - a_{31}a_{12}) \\
 &\quad + a_{33}(a_{11}a_{22} - a_{21}a_{12}) \\
 &= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} \\
 &\quad - a_{12}a_{21}a_{33} - a_{13}a_{22}a_{31} - a_{11}a_{23}a_{32}
 \end{aligned}$$

wie in Gleichung 153. Konkret folgt für die Matrix A aus Beispiel 208

$$\begin{aligned}
 \det \begin{pmatrix} 1 & 3 & 0 \\ 2 & -1 & 1 \\ 1 & 2 & 3 \end{pmatrix} &= -a_{23} \det A_{23} + a_{33} \det A_{33} \\
 &= -1(1 \cdot 2 - 1 \cdot 3) + 3(1 \cdot (-1) - 2 \cdot 3) = -20.
 \end{aligned}$$

BEMERKUNG 85. Bei konkreten Berechnungen läßt sich der Laplacesche Entwicklungssatz sinnvoll anwenden, wenn die Matrix A eine Spalte oder eine Zeile mit vielen Nullen besitzt.

2.4. Produktregel. Die Determinante ist eine multiplikative Abbildung.

SATZ 106. Für beliebige $n \times n$ -Matrizen A und B gilt

$$(155) \quad \det(AB) = (\det A)(\det B) = \det(BA).$$

BEMERKUNG 86. Es gilt also $\det(AB) = \det(BA)$, obwohl die Matrizenmultiplikation nicht kommutativ ist, also in den meisten Fällen $AB \neq BA$ gilt. $\square$

3. Die Inverse Matrix

Da die Determinante einer Matrix in Stufenform das Produkt der Diagonalelemente ist (siehe Abschnitt 2.1) und sich bei der Anwendung der Zeilenoperationen des Gaußverfahrens die Matrix nur um einen Faktor $\neq 0$ ändert, gilt

SATZ 107. Eine $n \times n$ -Matrix A hat genau dann den Rang n , wenn $\det A \neq 0$.

Es sei A eine $n \times n$ -Matrix. Das lineare Gleichungssystem $A\vec{x} = \vec{0}$ besitzt genau dann eine *eindeutige* Lösung, nämlich $\vec{x} = \vec{0}$, wenn $\text{rang } A = n$ gilt. Wenn $\text{rang } A = n$, dann besitzt das lineare Gleichungssystem $A\vec{x} = \vec{b}$ für jeden festen Vektor $\vec{b} \in \mathbb{R}^n$ genau eine Lösung. Daraus folgt

SATZ 108. Eine lineare Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^n$ ist genau dann bijektiv, wenn $\det A \neq 0$.

Die Umkehrabbildung A^{-1} , die den Vektor $A\vec{x}$ auf $\vec{x}$ abbildet, ist auch eine lineare Abbildung, denn

$$A^{-1}(A\vec{x} + A\vec{y}) = A^{-1}A(\vec{x} + \vec{y}) = \vec{x} + \vec{y} = A^{-1}(A\vec{x}) + A^{-1}(A\vec{y})$$

und $A^{-1}(\lambda A\vec{x}) = A^{-1}A(\lambda\vec{x}) = \lambda\vec{x} = \lambda A\vec{x}$. Die Matrix A^{-1} der Umkehrabbildung heißt **inverse Matrix**.

3.1. Bestimmung der inversen Matrix mit Hilfe des Gaußverfahrens. Die Spalte $\vec{a}_j$ der Matrix A^{-1} ist der Vektor $A^{-1}(\vec{e}_j)$. Es gilt $\vec{a}_j = A^{-1}(\vec{e}_j)$ genau dann, wenn $A\vec{a}_j = AA^{-1}\vec{e}_j = \vec{e}_j$. Die Spalte $\vec{a}_j$ ist Lösung des linearen Gleichungssystems $A\vec{x} = \vec{e}_j$. Die eindeutigen Lösungen n linearer Gleichungssysteme für $j = 1, \dots, n$ können mit Hilfe des Gaußverfahrens gleichzeitig bestimmt werden. Dazu wandelt man die erweiterte Koeffizientenmatrix $(A|E)$ durch Vertauschung von Zeilen, Multiplikation von Zeilen mit von 0 verschiedenen Zahlen und Addition von Vielfachen einer Zeile zu einer anderen Zeile in eine Matrix der Form $(E|B)$. Auf der rechten Seite erhält man die inverse Matrix, also $B = A^{-1}$.

BEISPIEL 211. Die Matrix

$$A = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}$$

ist invertierbar, denn $\det A = 1 \neq 0$. Mit Hilfe des Gaußverfahrens erhält man

$$(A|E) = \left(\begin{array}{cc|cc} 1 & 1 & 1 & 0 \\ 1 & 2 & 0 & 1 \end{array} \right) \sim \left(\begin{array}{cc|cc} 1 & 1 & 1 & 0 \\ 0 & 1 & -1 & 1 \end{array} \right) \sim \left(\begin{array}{cc|cc} 1 & 0 & 2 & -1 \\ 0 & 1 & -1 & 1 \end{array} \right),$$

also

$$A^{-1} = \begin{pmatrix} 2 & -1 \\ -1 & 1 \end{pmatrix}.$$

3.2. Explizite Formel der inversen Matrix. Wenn $A : \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung mit $\det A \neq 0$ ist, dann kann man für jeden Vektor $\vec{b} \in \mathbb{R}^n$ die eindeutige Lösung des linearen Gleichungssystems $A\vec{x} = \vec{b}$ mit Hilfe der Determinante bestimmen.

Wenn $A\vec{x} = \vec{b}$ für den Spaltenvektor $\vec{x} = (x_1, \dots, x_n)$ ist, dann gilt für alle Indizes $i = 1, \dots, n$

$$(156) \quad x_i = \frac{1}{\det A} \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{b}, \vec{a}_{i+1}, \dots, \vec{a}_n),$$

denn

$$\begin{aligned}
 \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{b}, \vec{a}_{i+1}, \dots, \vec{a}_n) &= \det(\vec{a}_1, \dots, \vec{a}_{i-1}, A\vec{x}, \vec{a}_{i+1}, \dots, \vec{a}_n) \\
 &= \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \sum_{j=1}^n \vec{a}_j x_j, \vec{a}_{i+1}, \dots, \vec{a}_n) \\
 &= \sum_{j=1}^n x_j \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{a}_j, \vec{a}_{i+1}, \dots, \vec{a}_n) \\
 &= x_i \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{a}_i, \vec{a}_{i+1}, \dots, \vec{a}_n) \\
 &= x_i \det A,
 \end{aligned}$$

wobei $\vec{a}_j$ die Spaltenvektoren der Matrix A sind.

Für den Vektor $\vec{b} = \vec{e}_j$ erhält man aus Formel 156

$$A \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \vec{e}_j \Leftrightarrow x_i = \frac{1}{\det A} \det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{e}_j, \vec{a}_{i+1}, \dots, \vec{a}_n) \quad \forall i = 1, \dots, n.$$

Entwicklung der Determinante im Zähler nach der i -ten Spalte liefert

$$\det(\vec{a}_1, \dots, \vec{a}_{i-1}, \vec{e}_j, \vec{a}_{i+1}, \dots, \vec{a}_n) = (-1)^{i+j} \det A_{ji}$$

und damit eine explizite Formel der inversen Matrix

$$(157) \quad A^{-1} = (\tilde{a}_{ij}) \text{ mit } \tilde{a}_{ij} = (-1)^{i+j} \frac{\det A_{ji}}{\det A}.$$

4. Orthogonale Abbildungen

Eine lineare Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ bzw. ihre Matrix A bezüglich der Standardbasis des $\mathbb{R}^n$ heißen **orthogonal**, falls $\langle \vec{x}, \vec{y} \rangle = \langle A\vec{x}, A\vec{y} \rangle$ für alle $\vec{x}, \vec{y} \in \mathbb{R}^n$. Insbesondere gilt für eine orthogonale Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^n$

- $\|A\vec{x}\| = \|\vec{x}\|$ für alle $\vec{x} \in \mathbb{R}^n$, d.h. die Länge der Vektoren bleibt gleich.
- $\vec{x} \perp \vec{y} \Leftrightarrow (A\vec{x}) \perp (A\vec{y})$, d.h. die Winkel zwischen Vektoren bleiben gleich.

4.1. Die transponierte Matrix. Es sei $A = (a_{ij})$ eine $m \times n$ -Matrix. Die $n \times m$ -Matrix $B = (b_{ij})$ mit $b_{ij} = a_{ji}$ heißt **transponierte** der Matrix A und wird mit A^T bezeichnet, also

$$(158) \quad A^T = \begin{pmatrix} a_{11} & a_{21} & \cdots & a_{m1} \\ a_{12} & a_{22} & \cdots & a_{m2} \\ \vdots & \vdots & & \vdots \\ a_{1n} & a_{2n} & \cdots & a_{mn} \end{pmatrix} \text{ für } A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix}.$$

Transponiert man einen Zeilenvektor, so erhält man einen Spaltenvektor. Transponiert man einen Spaltenvektor, so erhält man einen Zeilenvektor.

$$(x_1 \ x_2 \ \cdots \ x_n)^T = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \text{ und } \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}^T = (x_1 \ x_2 \ \cdots \ x_n)$$

SATZ 109. Für jede quadratische Matrix A gilt $\det A = \det(A^T)$. Wenn A eine $m \times n$ -Matrix und B eine $k \times m$ -Matrix sind, dann $(BA)^T = A^T B^T$.

BEWEIS. Die Gleichung $\det A = \det(A^T)$ folgt aus der Leibnizformel der Determinante, weil die Permutationen eine Gruppe bilden. Für die Matrizen $B = (b_{ij})$ und $A = (a_{jk})$ gilt

$$\begin{aligned} ((BA)^T)_{ij} &= (BA)_{ji} = \sum_{k=1}^m b_{jk} a_{ki} = \sum_{k=1}^m (B^T)_{kj} (A^T)_{ik} \\ &= \sum_{k=1}^m (A^T)_{ik} (B^T)_{kj} = (A^T B^T)_{ij} \end{aligned}$$

□

4.2. Skalarprodukt und Matrizenmultiplikation. Elemente des Vektorraumes $\mathbb{R}^n$ sind Spaltenvektoren, also Matrizen, die aus n Zeilen und einer Spalte bestehen. Auch wenn man $\vec{x} = (x_1, \dots, x_n) \in \mathbb{R}^n$ schreibt, meint man den Spaltenvektor

$$\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Nun gilt für $\vec{x} = (x_1, \dots, x_n), \vec{y} = (y_1, \dots, y_n) \in \mathbb{R}^n$

$$\langle \vec{x}, \vec{y} \rangle = x_1 y_1 + \dots + x_n y_n = (x_1 \ \cdots \ x_n) \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} = \vec{x}^T \vec{y}.$$

4.3. Charakterisierung orthogonaler Matrizen.

SATZ 110. Es sei A die Matrix einer linearen Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$ bezüglich der Standardbasis. Die folgenden Bedingungen sind äquivalent:

- (1) A ist orthogonal.
- (2) Die Spalten der Matrix A bilden eine Orthonormalbasis des $\mathbb{R}^n$.
- (3) $AA^T = E$
- (4) $A^{-1} = A^T$

BEWEIS. Es gilt $\langle A\vec{x}, A\vec{y} \rangle = (A\vec{x})^T (A\vec{y}) = \vec{x}^T A^T A \vec{y}$. Die Behauptung folgt aus $\langle \vec{e}_j, \vec{e}_j \rangle = (A^T A)_{ij}$. □

FOLGERUNG 37. Wenn A orthogonal ist, dann gilt $\det A = \pm 1$.

BEWEIS. $1 = \det E = \det(AA^T) = (\det A)(\det A^T) = (\det A)^2$. $\square$

4.4. Drehungen und Spiegelungen. Die Drehung A um den Winkel φ in $\mathbb{R}^2$ (siehe Beispiel 202) ist eine orthogonale Abbildung, denn

$$A^T = \begin{pmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{pmatrix}^T = \begin{pmatrix} \cos \varphi & \sin \varphi \\ -\sin \varphi & \cos \varphi \end{pmatrix} = \begin{pmatrix} \cos(-\varphi) & -\sin(-\varphi) \\ \sin(-\varphi) & \cos(-\varphi) \end{pmatrix}$$

ist eine Drehung um den Winkel $-\varphi$, also $A^T = A^{-1}$. Die Spaltenvektoren $\vec{a}_1 = (\cos \varphi, \sin \varphi)$ und $\vec{a}_2 = (-\sin \varphi, \cos \varphi)$ der Matrix A stehen senkrecht aufeinander und haben die Länge 1, denn $\cos^2 \varphi + \sin^2 \varphi = 1$. Es gilt $\det A = 1$.

Jede orthogonale Matrix A mit $\det A = 1$ ist eine Drehung und damit Produkt von Drehungen in einer Ebene $\subset \mathbb{R}^n$. Die Spiegelung an den Hyperebene $E_j := \{x_j = 0\} \subset \mathbb{R}^n$ ist definiert durch $\sigma_j : \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit $\sigma_j(\vec{e}_i) = \vec{e}_i$ für alle $i \neq j$ und $\sigma_j(\vec{e}_j) = -\vec{e}_j$ und auch eine orthogonale Abbildung. Es gilt $\det(\sigma_j) = -1$ für alle j . Jede orthogonale Abbildung ist Produkt von Spiegelungen und Drehungen.

5. Normalformen einer quadratischen Matrix

Die Diagonalmatrix

$$D = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{pmatrix} \quad \text{mit } \lambda_1, \dots, \lambda_n \in \mathbb{R}$$

ist genau dann orthogonal, wenn $\lambda_j = \pm 1$ für alle $j = 1, \dots, n$. Falls $\lambda_j \neq \pm 1$ für ein j , so verändert D die Länge des Vektors $\vec{e}_j$, d.h. $\|D\vec{e}_j\| = \|\lambda_j \vec{e}_j\| = |\lambda_j| \|\vec{e}_j\| = |\lambda_j|$. Falls $\lambda_1 \dots \lambda_n \neq 1$, so ändert D das Volumen. Falls nicht alle $|\lambda_j|$ gleich sind, also $|\lambda_j| \neq |\lambda_i|$ für $i \neq j$, so ändert die Abbildung D die Winkel zwischen Vektoren.

Auch wenn eine Diagonalmatrix D in den meisten Fällen nicht orthogonal ist, kann man die Abbildung D gut verstehen und für jeden Vektor $\vec{v} \in \mathbb{R}^n$ leicht $D\vec{v}$ berechnen. Deshalb möchte für eine gegebene lineare Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine Basis des $\mathbb{R}^n$ finden, bezüglich der die Matrix der Abbildung T eine Diagonalmatrix ist (siehe Abschnitt 5.1).

Es gibt aber auch Matrizen, die nicht diagonalisierbar sind, zum Beispiel

$$A = \begin{pmatrix} 1 & a \\ 0 & 1 \end{pmatrix} \quad \text{mit } a \in \mathbb{R}.$$

Für alle $a \in \mathbb{R}$ gilt $\det A = 1$. Falls $a \neq 0$, also $A \neq E$, so ändert A Längen und Winkel zwischen Vektoren. Auch für weder über den reellen noch über den komplexen Zahlen diagonalisierbare Matrizen existieren Normalformen (siehe 5.3).

5.1. Diagonalisierbare Matrizen. Eine lineare Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt **diagonalisierbar**, falls die Matrix der Abbildung T bezüglich einer Basis $\{\vec{v}_1, \dots, \vec{v}_n\}$ Diagonalform besitzt, also für alle $j = 1, \dots, n$ reelle Zahlen $\lambda_j \in \mathbb{R}$ mit $T(\vec{v}_j) = \lambda_j \vec{v}_j$ existieren.

Ein Vektor $\vec{v} \in \mathbb{R}^n$ heißt **Eigenvektor** der Abbildung T zum Eigenwert $\lambda \in \mathbb{R}$, falls $T(\vec{v}) = \lambda \vec{v}$. Der Eigenraum E_λ der Abbildung T zum Eigenwert $\lambda \in \mathbb{R}$ besteht aus allen Eigenvektoren der Abbildung T zum Eigenwert λ , also

$$E_\lambda := \{\vec{v} \in \mathbb{R}^n : T(\vec{v}) = \lambda \vec{v}\}.$$

Es seien A die Matrix der Abbildung T bezüglich der Standardbasis und $\lambda \in \mathbb{R}$ fest. Es gilt genau dann $T(\vec{v}) = \lambda \vec{v}$, wenn $\vec{v}$ Lösung des linearen Gleichungssystems $(A - \lambda E)\vec{v} = \vec{0}$ ist.

SATZ 111. *Der Eigenraum E_λ einer linearen Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit der Matrix A ist ein Untervektorraum des $\mathbb{R}^n$ mit*

$$\dim E_\lambda = n - \text{rang}(A - \lambda E).$$

Eine reelle Zahl $\lambda \in \mathbb{R}$ heißt **Eigenwert** der Matrix A , falls $\dim E_\lambda \geq 1$, also der Eigenraum E_λ nicht nur den Nullvektor enthält. Da

$$\dim E_\lambda \geq 1 \Leftrightarrow \text{rang}(A - \lambda E) < n \Leftrightarrow \det(A - \lambda E) = 0,$$

sind die Eigenwerte einer Matrix A gerade die (reellen) Nullstellen des Polynoms $\det(A - \lambda E)$. Für eine gegebene Matrix A ist $\det(A - \lambda E)$ ein Polynom vom Grad n in der Variablen λ . Es heißt **charakteristisches Polynom** der Matrix A . Wir betrachten die Zerlegung des charakteristischen Polynoms in Linearfaktoren

$$\det(A - \lambda E) = (\lambda - \lambda_1)^{k_1} (\lambda - \lambda_2)^{k_2} \dots (\lambda - \lambda_m)^{k_m}$$

mit paarweise verschiedenen Nullstellen $\lambda_j \in \mathbb{C}$ und Multiplizitäten k_j .

SATZ 112. *Die Matrix ist genau dann diagonalisierbar, wenn für jede Nullstellen λ_j ihres charakteristischen Polynoms gilt $\lambda_j \in \mathbb{R}$ und $k_j = \dim E_{\lambda_j}$.*

In dem Fall existiert eine Basis des $\mathbb{R}^n$ aus Eigenvektoren $\vec{v}_1, \dots, \vec{v}_n$ und für die Matrix $C := (\vec{v}_1 \vec{v}_2 \dots \vec{v}_n)$ gilt

$$C^{-1}AC = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix},$$

wobei λ_j der Eigenwert des Eigenvektors $\vec{v}_j$ ist.

FOLGERUNG 38. *Besitzt eine $n \times n$ -Matrix n paarweise verschiedene Nullstellen, so ist sie diagonalisierbar.*

SATZ 113. *Wenn A symmetrisch ist, also $A^T = A$ gilt, dann ist A diagonalisierbar und es existiert eine Orthonormalbasis des $\mathbb{R}^n$ aus Eigenvektoren.*

BEMERKUNG 87. Wenn eine Basis eine Orthonormalbasis ist und nur aus Eigenvektoren besteht, dann ist die Matrix C in Satz 112 eine orthogonale Matrix und es gilt $C^T = C^{-1}$. $\square$

BEISPIEL 212. Die Matrix

$$A = \begin{pmatrix} 2 & 0 \\ 3 & -1 \end{pmatrix}$$

ist nicht symmetrisch, denn $A^T \neq A$. Es gilt

$$A - \lambda E = \begin{pmatrix} 2 - \lambda & 0 \\ 3 & -1 - \lambda \end{pmatrix}$$

und das charakteristische Polynom

$$\det(A - \lambda) = (2 - \lambda)(-1 - \lambda) - 0 \cdot 3 = (2 - \lambda)(-1 - \lambda)$$

hat die einfachen reellen Nullstellen $\lambda_1 = 2$ und $\lambda_2 = -1$. Damit ist die Matrix A diagonalisierbar und die Eigenräume sind die Lösungsmengen der linearen Gleichungssysteme $(A - \lambda_1)\vec{x} = \vec{0}$ und $(A - \lambda_2)\vec{x} = \vec{0}$.

Für den Eigenraum E_2 zum Eigenwert $\lambda_1 = 2$ gilt

$$A - \lambda_1 E = \begin{pmatrix} 2 - 2 & 0 \\ 3 & -1 - 2 \end{pmatrix} \sim \begin{pmatrix} 0 & 0 \\ 3 & -3 \end{pmatrix} \sim \begin{pmatrix} 0 & 0 \\ 1 & -1 \end{pmatrix}$$

und damit $E_2 = \{(t, t) : t \in \mathbb{R}\} = \mathbb{R}(1, 1)$. Der Eigenraum E_2 wird als vom Vektor $\vec{v}_1 := (1, 1)$ aufgespannt.

Für den Eigenraum E_{-1} zum Eigenwert $\lambda_2 = -1$ gilt

$$A - \lambda_2 E = \begin{pmatrix} 2 + 1 & 0 \\ 3 & -1 + 1 \end{pmatrix} \sim \begin{pmatrix} 3 & 0 \\ 3 & 0 \end{pmatrix} \sim \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$$

und damit $E_{-1} = \{(0, t) : t \in \mathbb{R}\} = \mathbb{R}(0, 1)$. Der Eigenraum E_{-1} wird als vom Vektor $\vec{v}_2 := (0, 1)$ aufgespannt.

Die Eigenräume stehen nicht senkrecht aufeinander. Mit $C := (\vec{v}_1 \vec{v}_2)$ erhält man

$$C = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix}, (C|E) = \left(\begin{array}{cc|cc} 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 \end{array} \right) \sim \left(\begin{array}{cc|cc} 1 & 0 & 1 & 0 \\ 0 & 1 & -1 & 1 \end{array} \right) = (E|C^{-1})$$

und

$$C^{-1}AC = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} 2 & 0 \\ 3 & -1 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} 2 & 0 \\ 2 & -1 \end{pmatrix} = \begin{pmatrix} 2 & 0 \\ 0 & -1 \end{pmatrix} = D.$$

BEISPIEL 213. Die Matrix

$$A = \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$$

ist symmetrisch, denn $A^T = A$ und damit diagonalisierbar. Es gilt

$$A - \lambda E = \begin{pmatrix} 1 - \lambda & 2 \\ 2 & 1 - \lambda \end{pmatrix}$$

und das charakteristische Polynom

$$\det(A - \lambda) = (1 - \lambda)(1 - \lambda) - 2^2 = \lambda^2 - 2\lambda - 3 = (\lambda - 3)(\lambda + 1)$$

hat die einfachen reellen Nullstellen $\lambda_1 = 3$ und $\lambda_2 = -1$. Die Eigenräume sind die Lösungsmengen der linearen Gleichungssysteme $(A - \lambda_1)\vec{x} = \vec{0}$ und $(A - \lambda_2)\vec{x} = \vec{0}$.

Für den Eigenraum E_2 zum Eigenwert $\lambda_1 = 3$ gilt

$$A - \lambda_1 E = \begin{pmatrix} 1-3 & 2 \\ 2 & 1-3 \end{pmatrix} \sim \begin{pmatrix} -2 & 2 \\ 2 & -2 \end{pmatrix} \sim \begin{pmatrix} -1 & 1 \\ 0 & 0 \end{pmatrix}$$

und damit $E_2 = \{(t, t) : t \in \mathbb{R}\} = \mathbb{R}(1, 1)$. Der Eigenraum E_2 wird als vom Vektor $\vec{v}_1 := (1, 1)$ aufgespannt.

Für den Eigenraum E_{-1} zum Eigenwert $\lambda_2 = -1$ gilt

$$A - \lambda_2 E = \begin{pmatrix} 1+1 & 2 \\ 2 & 1+1 \end{pmatrix} \sim \begin{pmatrix} 2 & 2 \\ 2 & 2 \end{pmatrix} \sim \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix}$$

und damit $E_{-1} = \{(t, -t) : t \in \mathbb{R}\} = \mathbb{R}(1, -1)$. Der Eigenraum E_{-1} wird als vom Vektor $\vec{v}_2 := (1, -1)$ aufgespannt.

Die Eigenräume stehen senkrecht aufeinander. Wählt man normierte Eigenvektoren, so erhält man eine orthogonale Matrix C . Mit

$$C := \begin{pmatrix} \frac{1}{\|\vec{v}_1\|} \vec{v}_1 & \frac{1}{\|\vec{v}_2\|} \vec{v}_2 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \text{ und } C^{-1} = C^T = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$$

erhält man

$$\begin{aligned} C^{-1}AC &= \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix} \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} = \frac{1}{2} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \begin{pmatrix} 3 & -1 \\ 3 & 1 \end{pmatrix} \\ &= \frac{1}{2} \begin{pmatrix} 6 & 0 \\ 0 & -2 \end{pmatrix} = \begin{pmatrix} 3 & 0 \\ 0 & -1 \end{pmatrix} = D. \end{aligned}$$

BEISPIEL 214. Die Matrix

$$A = \begin{pmatrix} 2 & 3 \\ 0 & 2 \end{pmatrix}$$

ist nicht symmetrisch, denn $A^T \neq A$. Es gilt

$$A - \lambda E = \begin{pmatrix} 2-\lambda & 3 \\ 0 & 2-\lambda \end{pmatrix}$$

und das charakteristische Polynom

$$\det(A - \lambda) = (2 - \lambda)(2 - \lambda) - 0 \cdot 3 = (2 - \lambda)^2$$

hat die doppelte reelle Nullstelle $\lambda_1 = 2$, also $k_1 = 2$. Der Eigenraum ist die Lösungsmenge des linearen Gleichungssystems $(A - \lambda_1)\vec{x} = \vec{0}$. Für den Eigenraum E_2 zum Eigenwert $\lambda_1 = 2$ gilt

$$A - \lambda_1 E = \begin{pmatrix} 2-2 & 3 \\ 0 & 2-2 \end{pmatrix} \sim \begin{pmatrix} 0 & 3 \\ 0 & 0 \end{pmatrix} \sim \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$$

und damit $E_2 = \{(t, 0) : t \in \mathbb{R}\} = \mathbb{R}(1, 0)$. Der Eigenraum E_2 wird als vom Vektor $\vec{v}_1 := (1, 0)$ aufgespannt. Es gilt $\dim E_2 = 1 < 2 = k_1$. Da die Vielfachheit des Eigenwertes $\lambda_1 = 2$ kleiner als die Dimension des Eigenwertes ist, ist die Matrix A nicht diagonalisierbar.

BEMERKUNG 88. Wenn $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ eine diagonalisierbare lineare Abbildung mit einem doppelten Eigenwert $\lambda \in \mathbb{R}$ ist, dann gilt $T(\vec{v}) = \lambda\vec{v}$ für alle $\vec{v} \in \mathbb{R}^2$. Also ist die Matrix der Abbildung T bezüglich jeder beliebigen Basis eine Diagonalmatrix. $\square$

5.2. Komplexe Eigenwerte. Die Drehung um den Winkel α in $\mathbb{R}^2$

$$(159) \quad D_\alpha := \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}$$

hat das charakteristische Polynom

$$\begin{aligned} \det(D_\alpha - \lambda E) &= \begin{vmatrix} \cos \alpha - \lambda & -\sin \alpha \\ \sin \alpha & \cos \alpha - \lambda \end{vmatrix} = (\cos \alpha - \lambda)^2 - \sin \alpha(-\sin \alpha) \\ &= \cos^2 \alpha - 2\lambda \cos \alpha + \lambda^2 + \sin^2 \alpha = \lambda^2 - 2\lambda \cos \alpha + 1 \end{aligned}$$

hat die Nullstellen

$$\lambda_{1,2} = \cos \alpha \pm \sqrt{\cos^2 \alpha - 1} = \cos \alpha \pm \sqrt{-\sin^2 \alpha} = \cos \alpha \pm i \sin \alpha = e^{\pm i\alpha}.$$

Die Eigenwerte der Drehung D_α sind genau dann reell, wenn $\alpha = k\pi$ für ein $k \in \mathbb{Z}$. Für $\alpha = 0$ bzw. $\alpha = \pi$ erhält man

$$D_0 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = E \text{ und } D_\pi = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} = -E.$$

Falls α kein ganzzahliges Vielfaches von π ist, so sind die beiden Eigenwerte der Matrix D_α nicht reell.

Identifiziert man den reellen Vektorraum $\mathbb{R}^2$ mit $\mathbb{C}$, also $(x, y) = x + iy = z$, so gilt

$$\begin{aligned} D_\alpha(z) &= \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} x \cos \alpha - y \sin \alpha \\ x \sin \alpha + y \cos \alpha \end{pmatrix} \\ &= x \cos \alpha - y \sin \alpha + i(x \sin \alpha + y \cos \alpha) \\ &= (x + iy)(\cos \alpha + i \sin \alpha) = ze^{i\alpha}. \end{aligned}$$

Die Drehung D_α multipliziert die komplexe Zahl z mit dem (komplexen) Streckungsfaktor $e^{i\alpha}$.

SATZ 114. Ist $\lambda = re^{i\alpha} \notin \mathbb{R}$ ein Eigenwert der linearen Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit Vielfachheit 1, so existiert eine Basis $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n\}$ des $\mathbb{R}^n$ bezüglich der die Matrix A der Abbildung T die folgende Form hat

$$\begin{pmatrix} rD_\alpha & 0 \\ 0 & B \end{pmatrix},$$

wobei D_α die 2×2 -Matrix aus Gleichung 159 ist und B eine $(n-2) \times (n-2)$ -Matrix ist.

5.3. Jordansche Normalform. Es sei A eine $n \times n$ -Matrix A mit paarweise verschiedenen, reellen Eigenwerten λ_j , also

$$\det(A - \lambda E) = (\lambda - \lambda_1)^{k_1} \dots (\lambda - \lambda_m)^{k_m}.$$

Die Vektorräume

$$\hat{E}_{\lambda_j} := \{\vec{v} \in \mathbb{R}^n : (A - \lambda_j E)^n \vec{v} = 0\}$$

enthalten die Eigenräume E_{λ_j} , also $E_{\lambda_j} \subset \hat{E}_{\lambda_j}$ für alle j .

SATZ 115. Für alle $j \neq k$ gilt $A(\hat{E}_{\lambda_j}) \subset \hat{E}_{\lambda_j}$, $\hat{E}_{\lambda_j} \cap \hat{E}_{\lambda_k} = \{\vec{0}\}$ und $\dim \hat{E}_{\lambda_j} = k_j$ und der Vektorraum $\mathbb{R}^n$ wird von den Vektorräumen $\hat{E}_{\lambda_j}$ erzeugt, d.h.

$$\mathbb{R}^n = \hat{E}_{\lambda_1} \oplus \dots \oplus \hat{E}_{\lambda_m},$$

und existiert eine Basis bezüglich der die Matrix der linearen Abbildung folgende Form hat

$$J = \begin{pmatrix} J_1 & 0 & \cdots & 0 \\ 0 & J_2 & & 0 \\ \vdots & & \ddots & \vdots \\ 0 & 0 & \cdots & J_k \end{pmatrix} \quad \text{mit } J_i = \begin{pmatrix} \lambda_i & 1 & 0 & \cdots & 0 \\ 0 & \lambda_i & 1 & & 0 \\ 0 & 0 & \lambda_i & \ddots & \vdots \\ \vdots & \vdots & & \ddots & 1 \\ 0 & 0 & 0 & \cdots & \lambda_i \end{pmatrix} \quad \text{für alle } i.$$

Die Matrix J heißt **Jordansche Normalform**. Die Untermatrizen J_i heißen Jordankästchen. Es können mehrere Jordankästchen zum gleichen Eigenwert λ auftauchen.

BEISPIEL 215. Die Matrix

$$A = \begin{pmatrix} 2 & 3 \\ 0 & 2 \end{pmatrix}$$

hat den doppelten Eigenwert 2 und ist nicht diagonalisierbar (siehe Beispiel 214). Es gilt $E_2 = \{(t, 0) : t \in \mathbb{R}\} = \mathbb{R}\vec{v}_1$ mit $\vec{v}_1 = (1, 0)$ und $\hat{E}_2 = \mathbb{R}^2$, denn

$$(A - 2E)^2 = \begin{pmatrix} 0 & 3 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 3 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

Wir suchen einen Vektor $\vec{v}_2$ mit $(A - 2E)\vec{v}_2 = \vec{v}_1$. Es ist eine Lösung des linearen Gleichungssystems

$$\left(\begin{array}{cc|c} 0 & 3 & 1 \\ 0 & 0 & 0 \end{array} \right),$$

also zum Beispiel $\vec{v}_2 = (0, 1/3)$. Es gilt $A\vec{v}_1 = 2\vec{v}_1$ und $A\vec{v}_2 = \vec{v}_1 + 2\vec{v}_2$. Bezüglich der Basis $\{\vec{v}_1, \vec{v}_2\}$ besitzt die lineare Abbildung A die Matrix

$$\begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}.$$

BEISPIEL 216. Die Matrix

$$A = \begin{pmatrix} 1 & 1 \\ -1 & 3 \end{pmatrix}$$

ist nicht symmetrisch, denn $A^T \neq A$. Es gilt

$$A - \lambda E = \begin{pmatrix} 1 - \lambda & 1 \\ -1 & 3 - \lambda \end{pmatrix}$$

und das charakteristische Polynom

$$\det(A - \lambda) = (1 - \lambda)(3 - \lambda) - 1 \cdot (-1) = \lambda^2 - 4\lambda + 4 = (\lambda - 2)^2$$

hat die doppelte reelle Nullstelle $\lambda_1 = 2$, also $k_1 = 2$. Der Eigenraum ist die Lösungsmenge des linearen Gleichungssystems $(A - \lambda_1)\vec{x} = \vec{0}$. Für den Eigenraum E_2 zum Eigenwert $\lambda_1 = 2$ gilt

$$A - \lambda_1 E = \begin{pmatrix} 1 - 2 & 1 \\ -1 & 3 - 2 \end{pmatrix} \sim \begin{pmatrix} -1 & 1 \\ -1 & 1 \end{pmatrix} \sim \begin{pmatrix} 1 & -1 \\ 0 & 0 \end{pmatrix}$$

und damit $E_2 = \{(t, t) : t \in \mathbb{R}\} = \mathbb{R}(1, 1)$. Der Eigenraum E_2 wird als vom Vektor $\vec{v}_1 := (1, 1)$ aufgespannt. Es gilt $\dim E_2 = 1 < 2 = k_1$. Da die Vielfachheit des Eigenwertes $\lambda_1 = 2$ kleiner als die Dimension des Eigenwertes ist, ist die Matrix A nicht diagonalisierbar.

Es $\hat{E}_2 = \mathbb{R}^2$, denn

$$(A - 2E)^2 = \begin{pmatrix} -1 & 1 \\ -1 & 1 \end{pmatrix} \begin{pmatrix} -1 & 1 \\ -1 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

Wir suchen einen Vektor $\vec{v}_2$ mit $(A - 2E)\vec{v}_2 = \vec{v}_1$. Es ist eine Lösung des linearen Gleichungssystems

$$\left(\begin{array}{cc|c} -1 & 1 & 1 \\ -1 & 1 & 1 \end{array} \right) \sim \left(\begin{array}{cc|c} 1 & -1 & -1 \\ 0 & 0 & 0 \end{array} \right),$$

also zum Beispiel $\vec{v}_2 = (-1, 0)$. Es gilt $A\vec{v}_1 = 2\vec{v}_1$ und $A\vec{v}_2 = \vec{v}_1 + 2\vec{v}_2$. Bezüglich der Basis $\{\vec{v}_1, \vec{v}_2\}$ besitzt die lineare Abbildung A die Matrix

$$\begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}.$$

6. Homogene, lineare DGL mit konstanten Koeffizienten

Eine homogene, lineare Differentialgleichung mit **konstanten Koeffizienten** n -ter Ordnung ist eine Gleichung der Form

$$(160) \quad y^{(n)} + a_{n-1}y^{(n-1)} + \dots + a_1y' + a_0y = 0 \text{ mit } a_j \in \mathbb{R}, a_n = 1.$$

SATZ 116. Wenn y eine Lösung der Differentialgleichung 160, dann ist auch y' eine Lösung dieser Differentialgleichung.

BEWEIS. Es gilt $\sum_{j=0}^n a_j(y')^{(j)} = \sum_{j=0}^n a_j y^{(j+1)} = (\sum_{j=0}^n a_j y^{(j)})' = 0$. $\square$

6.1. Das charakteristische Polynom. Eine Funktion $y(x) = e^{\lambda x}$ mit $\lambda \in \mathbb{R}$ ist genau dann eine Lösung der Differentialgleichung 160, wenn

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

da $(e^{\lambda x})^{(j)} = \lambda^j e^{\lambda x}$. Das Polynom $\sum_{j=0}^n a_j \lambda^j$ heißt **charakteristisches Polynom** der Differentialgleichung 160.

Das charakteristische Polynom der Differentialgleichung 160 stimmt mit dem charakteristischen Polynom der Matrix

$$A = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ 0 & 0 & 1 & & 0 \\ \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & 0 & \cdots & & 1 \\ -a_0 & -a_1 & -a_2 & \cdots & -a_{n-1} \end{pmatrix}$$

überein. Eine Funktion $y : \mathbb{R} \rightarrow \mathbb{R}$ ist genau dann Lösung der Differentialgleichung 160, wenn eine vektorwertige Funktion $\vec{y}$ Lösung einer linearen Differentialgleichung mit konstanten Koeffizienten erster Ordnung ist:

$$\vec{y}' = A\vec{y} \text{ mit } \vec{y} := \begin{pmatrix} y \\ y' \\ \vdots \\ y^{(n-1)} \end{pmatrix}$$

Eine Funktion $\vec{y} = \vec{c}e^{\lambda x}$ ist genau dann Lösung der Differentialgleichung $\vec{y}' = A\vec{y}$, wenn λ ein Eigenwert und $\vec{c}$ ein Eigenvektor der Matrix A ist, denn $(\vec{c}e^{\lambda x})' = \lambda \vec{c}e^{\lambda x}$.

BEMERKUNG 89. Für jede reelle Zahl $\lambda \in \mathbb{R}$ gilt $\text{rang}(A - \lambda E) \geq n - 1$. Darum ist jeder Eigenraum der Matrix A höchstens eindimensional. $\square$

6.2. Verschiedene, reelle Nullstellen - Exponentialfunktionen.

SATZ 117. Wenn das charakteristische Polynom der Differentialgleichung 160 n paarweise verschiedene, reelle Nullstellen $\lambda_1, \dots, \lambda_n$ besitzt, dann bilden die Funktionen $y_i(x) := e^{\lambda_i x}$, $i = 1, \dots, n$, eine Basis des Vektorraumes aller Lösungen.

BEWEIS. Die Funktionen y_i sind Lösungen der Differentialgleichung, da

$$\sum_{j=0}^n a_n y_i^{(j)} = \sum_{j=0}^n a_j \lambda_i^j e^{\lambda_i x} = e^{\lambda_i x} \sum_{j=0}^n a_j \lambda_i^j = 0.$$

Die Funktionen y_i , $i = 1, \dots, n$, sind linear unabhängig, weil

$$\det C = \det \left(y_i^{(j)}(0) \right)_{i=1, \dots, n; j=0, \dots, n-1} = \left| \lambda_i^j \right| = \prod_{k>i} (\lambda_k - \lambda_i) \neq 0$$

und damit die eindeutige Lösung jeder Anfangswertaufgabe $y^{(j)}(0) = \eta_j$ für $j = 0, \dots, n-1$ Linearkombination der Funktionen $y_i(x) = e^{\lambda_i x}$ ist. $\square$

BEISPIEL 217. Das charakteristische Polynom der homogenen, linearen Differentialgleichung zweiter Ordnung $y'' - 2y' - 3y = 0$ ist $\lambda^2 - 2\lambda - 3 = (\lambda - 3)(\lambda + 1)$. Es hat die reellen Nullstellen $\lambda_1 = 3$ und $\lambda_2 = -1$. Die allgemeine Lösung der Differentialgleichung $y'' - 2y' - 3y = 0$ ist also $y(x) = c_1 e^{3x} + c_2 e^{-x}$ mit $c_1, c_2 \in \mathbb{R}$.

6.3. Komplexe Nullstellen - Winkelfunktionen.

SATZ 118. Wenn das charakteristische Polynom der Differentialgleichung 160 die komplexe Nullstelle $\lambda = \alpha + i\beta$ besitzt, dann sind die Funktionen

$$\Re(e^{\lambda x}) = e^\alpha \cos \beta \text{ und } \Im(e^{\lambda x}) = e^\alpha \sin \beta$$

Lösungen der Differentialgleichung 160.

BEWEIS. Für die Ableitung einer Funktion $f : \mathbb{R} \rightarrow \mathbb{C}$ mit $f = u + iv$ und $u, v : \mathbb{R} \rightarrow \mathbb{R}$ gilt $f'(x) = u'(x) + iv'(x)$. $\square$

BEMERKUNG 90. Wenn $\lambda = \alpha + i\beta \notin \mathbb{R}$, also $\beta \neq 0$, dann ist auch $\bar{\lambda} = \alpha - i\beta$ eine Nullstelle des charakteristischen Polynoms. Aus dem Nullstellenpaar $\lambda, \bar{\lambda}$ erhält man zwei Lösungen der Differentialgleichung. $\square$

BEISPIEL 218. Das charakteristische Polynom der homogenen, linearen Differentialgleichung zweiter Ordnung $y'' + 4y = 0$ ist $\lambda^2 + 4$. Es hat die komplexen Nullstellen $\lambda = \pm 2i$. Die allgemeine Lösung der Differentialgleichung $y'' + 4y = 0$ ist also $y(x) = c_1 \cos(2x) + c_2 \sin(2x)$ mit $c_1, c_2 \in \mathbb{R}$.

6.4. Mehrfache Nullstellen - Resonanzfaktoren.

SATZ 119. Wenn das charakteristische Polynom der Differentialgleichung 160 nur reelle Nullstellen $\lambda_1, \dots, \lambda_m$ mit den Vielfachheiten $k_1, \dots, k_m$ besitzt, also

$$P(\lambda) := \sum_{j=0}^n a_j \lambda^j = (\lambda - \lambda_1)^{k_1} (\lambda - \lambda_2)^{k_2} \dots (\lambda - \lambda_m)^{k_m},$$

dann bilden die Funktionen $y_{il}(x) := x^l e^{\lambda_i x}$ mit $i = 1, \dots, m$ und $l = 0, \dots, k_i - 1$ eine Basis des Vektorraumes aller Lösungen.

BEWEIS. Es gilt

$$\begin{aligned}
 y_{il}^{(j)} &= \sum_{r=0}^j \binom{j}{r} l(l-1) \dots (l-r+1) x^{l-r} \lambda_i^{j-r} e^{\lambda_i x} \\
 \sum_{j=0}^n a_j y_{il}^{(j)} &= \sum_{j=0}^n a_j \sum_{r=0}^j \binom{j}{r} l(l-1) \dots (l-r+1) x^{l-r} \lambda_i^{j-r} e^{\lambda_i x} \\
 &= \sum_{r=0}^n l(l-1) \dots (l-r+1) x^{l-r} e^{\lambda_i x} \sum_{j=r}^n \binom{j}{r} a_j \lambda_i^{j-r} \\
 &= \sum_{r=0}^n \frac{l(l-1) \dots (l-r+1)}{r!} x^{l-r} e^{\lambda_i x} \sum_{j=r}^n j(j-1) \dots (j-r+1) a_j \lambda_i^{j-r} \\
 &= \sum_{r=0}^n \frac{l(l-1) \dots (l-r+1)}{r!} x^{l-r} e^{\lambda_i x} P^{(r)}(\lambda_i) = 0,
 \end{aligned}$$

weil für $r > l$ das Produkt $l(l-1) \dots (l-r+1)$ verschwindet und für $r \leq l$ gilt $P^{(r)}(\lambda_i) = 0$, da λ_i eine k_i -fache Nullstelle des charakteristischen Polynoms $P(\lambda)$ ist und $l < k_i$.

Die lineare Unabhängigkeit der Funktionen folgt wie im Satz 117 mit Hilfe der Vandermondschen Determinante. $\square$

Betrachtet man statt der Differentialgleichung n -ter Ordnung für die reellwertige Funktion y die Differentialgleichung erster Ordnung für die vektorwertige Funktion $\vec{y}$, so erhält man eine Basis des Vektorraumes aller Lösungen mit Hilfe der Matrix e^{Ax} .

BEISPIEL 219. Das charakteristische Polynom der homogenen, linearen Differentialgleichung zweiter Ordnung $y'' - 2y' + y = 0$ ist $\lambda^2 - 2\lambda + 1 = (\lambda - 1)^2$. Es hat die doppelte reelle Nullstelle $\lambda_1 = 1$. Die allgemeine Lösung der Differentialgleichung $y'' - 2y' + y = 0$ ist also $y(x) = c_1 e^x + c_2 x e^x$ mit $c_1, c_2 \in \mathbb{R}$.

Die Matrix A der dazugehörigen vektorwertigen Differentialgleichung ist

$$A = \begin{pmatrix} 0 & 1 \\ -1 & 2 \end{pmatrix} \text{ mit } \det(A - \lambda E) = -\lambda(2 - \lambda) + 1 = (\lambda - 1)^2.$$

Es gilt $E_1 = \{(t, t) : t \in \mathbb{R}\}$, denn

$$A - E = \begin{pmatrix} -1 & 1 \\ -1 & 1 \end{pmatrix} \sim \begin{pmatrix} -1 & 1 \\ 0 & 0 \end{pmatrix}.$$

Eine Lösung der Gleichung $(A - E)\vec{x} = \vec{v}_1 = (1, 1)^T$ ist $\vec{v}_2 = (0, 1)$, denn

$$(A - E|\vec{v}_1) = \left(\begin{array}{cc|c} -1 & 1 & 1 \\ -1 & 1 & 1 \end{array} \right) \sim \left(\begin{array}{cc|c} -1 & 1 & 1 \\ 0 & 0 & 0 \end{array} \right).$$

Nun gilt

$$\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} = C^{-1}AC \text{ mit } C = \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \text{ und } C^{-1} = \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix},$$

$A^n = CD^nC^{-1}$ und

$$\begin{aligned} e^{Ax} &= \sum_{n=0}^{\infty} \frac{x^n}{n!} A^n = C \left(\sum_{n=0}^{\infty} \frac{x^n}{n!} D^n \right) C^{-1} \\ C^{-1}e^{Ax}C &= \sum_{n=0}^{\infty} \frac{x^n}{n!} D^n = \sum_{n=0}^{\infty} \frac{x^n}{n!} \begin{pmatrix} 1 & n \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} e^x & xe^x \\ 0 & e^x \end{pmatrix} \\ e^{Ax} &= \begin{pmatrix} 1 & 0 \\ 1 & 1 \end{pmatrix} \begin{pmatrix} e^x & xe^x \\ 0 & e^x \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} = \begin{pmatrix} e^x & xe^x \\ e^x & (x+1)e^x \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -1 & 1 \end{pmatrix} \\ &= \begin{pmatrix} (1-x)e^x & xe^x \\ -xe^x & (x+1)e^x \end{pmatrix} \end{aligned}$$

Die Einträge der ersten Zeile der Matrizen e^{Dx} und e^{Ax} bilden jeweils eine Basis des Vektorraumes aller Lösungen der Differentialgleichung $y'' - 2y' + y = 0$.

KAPITEL XIV

Quadratische Formen

Eine $n \times n$ -Matrix A definiert eine reellwertige Funktion $q : \mathbb{R}^n \rightarrow \mathbb{R}$ durch $\vec{x} \mapsto \vec{x}^T A \vec{x}$, d.h.

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \mapsto \begin{pmatrix} x_1 & \cdots & x_n \end{pmatrix} \begin{pmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \cdots & a_{nn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = \sum_{i,j=1}^n a_{ij} x_i x_j.$$

BEISPIEL 220. Die Matrix

$$A = \begin{pmatrix} 2 & 0 \\ 0 & -3 \end{pmatrix}$$

definiert die Funktion $q : \mathbb{R}^2 \rightarrow \mathbb{R}$,

$$q(x_1, x_2) = \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 2 & 0 \\ 0 & -3 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 2x_1 \\ -3x_2 \end{pmatrix} = 2x_1^2 - 3x_2^2.$$

BEISPIEL 221. Die Matrix

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

definiert die Funktion $q : \mathbb{R}^2 \rightarrow \mathbb{R}$,

$$\begin{aligned} q(x_1, x_2) &= \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 3 & 2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \\ &= \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} x_2 \\ 3x_1 + 2x_2 \end{pmatrix} = x_1x_2 + 3x_1x_2 + 2x_2^2 = 4x_1x_2 + 2x_2^2. \end{aligned}$$

Auch Matrix

$$A = \begin{pmatrix} 0 & 2 \\ 2 & 2 \end{pmatrix}$$

definiert die Funktion $q : \mathbb{R}^2 \rightarrow \mathbb{R}$,

$$\begin{aligned} q(x_1, x_2) &= \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 0 & 2 \\ 2 & 2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \\ &= \begin{pmatrix} x_1 & x_2 \end{pmatrix} \begin{pmatrix} 2x_2 \\ 2x_1 + 2x_2 \end{pmatrix} = 2x_1x_2 + 2x_1x_2 + 2x_2^2 = 4x_1x_2 + 2x_2^2. \end{aligned}$$

Umgekehrt läßt sich jedes homogene Polynom zweiten Grades in den Variablen $x_1, \dots, x_n$, d.h. eine Funktion der Form

$$q(x_1, \dots, x_n) = \sum_{i=1}^n a_{ii} x_i^2 + \sum_{j>i} 2a_{ij} x_i x_j$$

mit reellen Koeffizienten a_{ij} mit Hilfe einer symmetrischen Matrix A als $q(\vec{x}) = \vec{x}^T A \vec{x}$ schreiben. Dazu teilt man für $i < j$ den Koeffizienten $2a_{ij}$ vor $x_i x_j$ auf, also $2a_{ij} = a_{ij} + a_{ji}$ mit $a_{ji} := a_{ij}$ und erhält

$$\vec{x}^T A \vec{x} = \sum_{i=1}^n \sum_{j=1}^n a_{ij} x_i x_j = q(\vec{x}).$$

1. Extremwerte quadratischer Formen auf der Kugel

Wir suchen die Extremwerte einer quadratischen Form $q(\vec{x})$ unter der Nebenbedingung $\|\vec{x}\|^2 = x_1^2 + \dots + x_n^2 = 1$.

Treten im homogenen quadratischen Polynom q keine gemischten Terme auf, also $q(x_1, \dots, x_n) = \lambda_1 x_1^2 + \dots + \lambda_n x_n^2$ für $\lambda_i \in \mathbb{R}$, so gilt

$$\max\{q(\vec{x}) : \|\vec{x}\|^2 = 1\} = \max\{\lambda_1, \dots, \lambda_n\}$$

und

$$\min\{q(\vec{x}) : \|\vec{x}\|^2 = 1\} = \min\{\lambda_1, \dots, \lambda_n\},$$

denn für $\lambda_k := \max\{\lambda_1, \dots, \lambda_n\}$ gilt

$$\begin{aligned} q(\vec{x}) &= q(\vec{x}) - \lambda_k + \lambda_k = q(\vec{x}) - \lambda_k(x_1^2 + \dots + x_n^2) + \lambda_k = \lambda_k + \sum_{i=1}^n (\lambda_i - \lambda_k) x_i^2 \\ &= \lambda_k + \sum_{i \neq k} (\lambda_i - \lambda_k) x_i^2 \leq \lambda_k = q(\vec{e}_k) \end{aligned}$$

und für $\lambda_l := \min\{\lambda_1, \dots, \lambda_n\}$ gilt

$$\begin{aligned} q(\vec{x}) &= q(\vec{x}) - \lambda_l + \lambda_l = q(\vec{x}) - \lambda_l(x_1^2 + \dots + x_n^2) + \lambda_l = \lambda_l + \sum_{i=1}^n (\lambda_i - \lambda_l) x_i^2 \\ &= \lambda_l + \sum_{i \neq l} (\lambda_i - \lambda_l) x_i^2 \geq \lambda_l = q(\vec{e}_l). \end{aligned}$$

Die Funktion q wächst in Richtung der x_k -Achse am stärksten und in Richtung der x_l -Achse am schwächsten.

Betrachtet man ein homogenes quadratisches Polynom $q(\vec{x}) = \vec{x}^T A \vec{x}$ mit einer symmetrischen Matrix A , so ist die Matrix diagonalisierbar, d.h. es existiert eine Orthonormalbasis $\{\vec{v}_1, \dots, \vec{v}_n\} \subset \mathbb{R}^n$ aus Eigenvektoren, also $A\vec{v}_j = \lambda_j \vec{v}_j$ für Eigenwerte $\lambda_j \in \mathbb{R}$. Wie ermittelt man die Koordinaten eines Vektors $\vec{x} \in \mathbb{R}^n$ bezüglich der Basis $\{\vec{v}_1, \dots, \vec{v}_n\}$ aus den Koordinaten

bezüglich der Standardbasis? Mit Hilfe der Matrix $C = (\vec{v}_1 \cdots \vec{v}_n) = (c_{ij})$, also $\vec{v}_j = \sum_{i=1}^n c_{ij} \vec{e}_i$, erhält man

$$\begin{aligned} \vec{x} &= \sum_{i=1}^n x_i \vec{e}_i = \sum_{j=1}^n \xi_j \vec{v}_j = \sum_{j=1}^n \xi_j \sum_{i=1}^n c_{ij} \vec{e}_i = \sum_{i=1}^n \left(\sum_{j=1}^n \xi_j c_{ij} \right) \vec{e}_i \\ &= \sum_{i=1}^n \left(\sum_{j=1}^n c_{ij} \xi_j \right) \vec{e}_i. \end{aligned}$$

Dies bedeutet

$$\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = C \begin{pmatrix} \xi_1 \\ \vdots \\ \xi_n \end{pmatrix} = C \vec{\xi}, \quad \vec{x}^T = \vec{\xi}^T C^T \quad \text{und} \quad q(\vec{\xi}) = \vec{\xi}^T C^T A C \vec{\xi} = \sum_{i=1}^n \lambda_i \xi_i^2.$$

Da C eine orthogonale Matrix ist, und damit die Länge eines Vektors nicht ändert, gilt der folgende

SATZ 120. *Es seien λ der größte und μ der kleinste Eigenwert der symmetrischen Matrix A . Für die quadratische Form $q : \mathbb{R}^n \rightarrow \mathbb{R}$ mit $q(\vec{x}) = \vec{x}^T A \vec{x}$ gilt*

$$(161) \quad \max\{q(\vec{x}) : \|\vec{x}\| = 1\} = \lambda \quad \text{und} \quad \min\{q(\vec{x}) : \|\vec{x}\| = 1\} = \mu.$$

Für $\vec{x}$ mit $\|\vec{x}\| = 1$ gilt genau dann $q(\vec{x}) = \lambda$, wenn $\vec{x}$ eine Eigenvektor zum Eigenwert λ ist, also $\vec{x} \in E_\lambda$.

Für $\vec{x}$ mit $\|\vec{x}\| = 1$ gilt genau dann $q(\vec{x}) = \mu$, wenn $\vec{x}$ eine Eigenvektor zum Eigenwert μ ist, also $\vec{x} \in E_\mu$.

BEISPIEL 222. Zu der quadratischen Form $q(x_1, x_2, x_3) = x_1^2 + 4x_1x_3$ gehört die quadratische Matrix

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

mit dem charakteristischen Polynom

$$\det(A - \lambda E) = \begin{vmatrix} 1 - \lambda & 0 & 2 \\ 0 & -\lambda & 0 \\ 2 & 0 & -\lambda \end{vmatrix} = (-\lambda)((1 - \lambda)(-\lambda) - 4) = -\lambda(\lambda^2 - \lambda - 4)$$

und den Eigenwerten $\lambda_1 = 0$ und

$$\lambda_{2,3} = \frac{1}{2} \pm \sqrt{\frac{1}{4} + 4} = \frac{1}{2}(1 \pm \sqrt{17}).$$

Da $\lambda_2 = \frac{1}{2}(1 + \sqrt{17}) > 0 > \lambda_3 = \frac{1}{2}(1 - \sqrt{17})$, ist $\frac{1}{2}(1 + \sqrt{17})$ das Maximum und $\frac{1}{2}(1 - \sqrt{17})$ das Minimum der Funktion q unter der Nebenbedingung $x_1^2 + x_2^2 + x_3^2 = 1$.

BEISPIEL 223. Zu der quadratischen Form

$$q(x_1, x_2, x_3) = 7x_1^2 + 7x_2^2 - 2x_3^2 - 10x_1x_2 + 8x_1x_3 + 8x_2x_3$$

gehört die symmetrische quadratische Matrix

$$A = \begin{pmatrix} 7 & -5 & 4 \\ -5 & 7 & 4 \\ 4 & 4 & -2 \end{pmatrix}$$

mit dem charakteristischen Polynom

$$\begin{aligned} \det(A - \lambda E) &= \begin{vmatrix} 7-\lambda & -5 & 4 \\ -5 & 7-\lambda & 4 \\ 4 & 4 & -2-\lambda \end{vmatrix} = \begin{vmatrix} 7-\lambda & -5 & 4 \\ -12+\lambda & 12-\lambda & 0 \\ 4 & 4 & -2-\lambda \end{vmatrix} \\ &= \begin{vmatrix} 7-\lambda & 2-\lambda & 4 \\ -12+\lambda & 0 & 0 \\ 4 & 8 & -2-\lambda \end{vmatrix} \\ &= (-1)(-12+\lambda)((2-\lambda)(-2-\lambda) - 32) = -(\lambda-12)(\lambda^2-36) \\ &= -(\lambda-12)(\lambda-6)(\lambda+6) \end{aligned}$$

und den Eigenwerten $\lambda_1 = 12$, $\lambda_2 = 6$ und $\lambda_3 = -6$. Damit ist $\lambda_1 = 12$ das Maximum und $\lambda_3 = -6$ das Minimum der Funktion q unter der Nebenbedingung $x_1^2 + x_2^2 + x_3^2 = 1$.

Der Eigenraum E_{12} ist $E_{12} = \{(t, -t, 0) : t \in \mathbb{R}\}$, denn

$$A - 12E = \begin{pmatrix} -5 & -5 & 4 \\ -5 & -5 & 4 \\ 4 & 4 & -14 \end{pmatrix} \sim \begin{pmatrix} 10 & 10 & -8 \\ 0 & 0 & 0 \\ 2 & 2 & -7 \end{pmatrix} \sim \begin{pmatrix} 5 & 5 & -4 \\ 0 & 0 & 0 \\ 0 & 0 & 33 \end{pmatrix} \sim \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Es gibt im Eigenraum E_{12} zwei Vektoren der Länge 1. Daraus folgt

$$12 = q\left(\frac{1}{2}\sqrt{2}, -\frac{1}{2}\sqrt{2}, 0\right) = q\left(-\frac{1}{2}\sqrt{2}, \frac{1}{2}\sqrt{2}, 0\right).$$

Der Eigenraum E_{-6} ist $E_{-6} = \{(t, t, -2t) : t \in \mathbb{R}\}$, denn

$$A - 12E = \begin{pmatrix} 13 & -5 & 4 \\ -5 & 13 & 4 \\ 4 & 4 & 4 \end{pmatrix} \sim \begin{pmatrix} 9 & -9 & 0 \\ -9 & 9 & 0 \\ 1 & 1 & 1 \end{pmatrix} \sim \begin{pmatrix} 0 & 0 & 0 \\ -1 & 1 & 0 \\ 1 & 1 & 1 \end{pmatrix}.$$

Es gibt im Eigenraum E_{-6} zwei Vektoren der Länge 1. Daraus folgt

$$-6 = q\left(\frac{1}{6}\sqrt{6}, \frac{1}{6}\sqrt{6}, -\frac{1}{3}\sqrt{6}\right) = q\left(-\frac{1}{6}\sqrt{6}, -\frac{1}{6}\sqrt{6}, \frac{1}{3}\sqrt{6}\right).$$

2. Hauptachsentransformation in $\mathbb{R}^2$, Kegelschnitte

Eine **quadratische Gleichung** in zwei Variablen, x und y , mit Koeffizienten $a, b, c, d, e, f \in \mathbb{R}$, $(a, b, c) \neq (0, 0, 0)$, ist eine Gleichung der Form

$$(162) \quad 0 = ax^2 + 2bxy + cy^2 + dx + dy + f =: q(x, y).$$

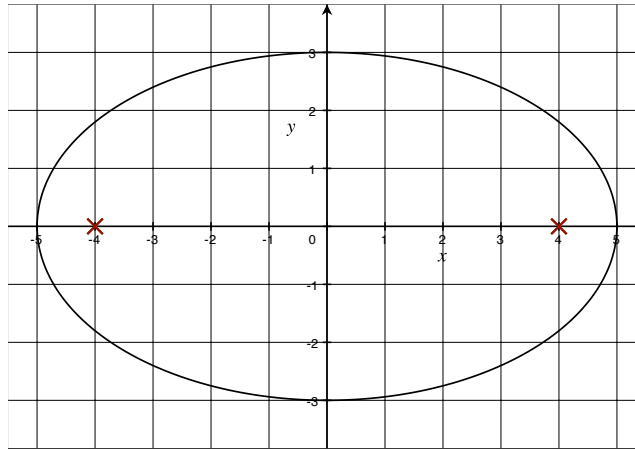


ABBILDUNG 1

BEMERKUNG 91. Falls $a = b = c = 0$, so ist Gleichung 162 eine lineare Gleichung. $\square$

BEISPIEL 224. Für $d = e = b = 0$, $a = c = 1$ und $f = -r^2$ erhält man die Gleichung $x^2 + y^2 = r^2$. Die Lösungsmenge ist die Menge aller Vektoren (x, y) mit Länge r , also ein Kreis mit Radius r um $(0, 0)$.

BEISPIEL 225. Für $d = e = b = 0$ und $f = 1$, $a = 1/25$ und $b = 1/9$ erhält man die Gleichung

$$\left(\frac{x}{5}\right)^2 + \left(\frac{y}{3}\right)^2 = 1.$$

Die Lösungsmenge ist eine Ellipse mit Mittelpunkt $(0, 0)$, den Halbachsenlängen 5 und 3 und den Brennpunkten $(4, 0)$ und $(-4, 0)$ (siehe Abbildung 1).

BEISPIEL 226. Für $d = e = b = 0$, $c = -1$, $a = 1$ und $f = -1$ erhält man die Gleichung $x^2 - y^2 = 1$. Die Lösungsmenge besteht aus zwei Hyperbelarmen, $y = \pm\sqrt{x^2 - 1}$ (siehe Abbildung 2).

BEISPIEL 227. Für $d = b = c = f = 0$, $e = -1$ und $a = 1$ erhält man die Gleichung $x^2 - y = 0$, also die Parabel $y = x^2$.

BEISPIEL 228. Für $d = e = b = c = 0$, $a = 1$ und $f = -1$ erhält man die Gleichung $x^2 = 1$, also $x = \pm 1$ und $y \in \mathbb{R}$ beliebig. Die Lösungsmenge besteht aus zwei parallelen Geraden.

BEISPIEL 229. Für $d = e = b = c = f = 0$ und $a = 1$ erhält man die Gleichung $x^2 = 0$. Die Lösungsmenge enthält nur den Punkt $(0, 0)$.

BEISPIEL 230. Für $d = e = b = c = 0$ und $a = f = 1$ erhält man die Gleichung $x^2 = -1$. Die Lösungsmenge ist leer.

Wir werden sehen, dass die Lösungsmenge der Gleichung 162 immer eine Ellipse, eine Hyperbel, eine Parabel, zwei Geraden, ein Punkt oder die leere Menge ist. Allerdings kann sie im Vergleich zu den Beispielen 224 bis 229 gestreckt, gestaucht, gedreht oder verschoben, also nicht im Punkt $(0,0)$ zentriert, sein.

Gleichung 162 lässt sich mit Hilfe einer symmetrischen 2×2 -Matrix A und eines Vektors $\vec{v} \in \mathbb{R}$ schreiben als

$$0 = q(x, y) = \vec{x}^T A \vec{x} + \vec{v}^T \vec{x} + f \text{ mit } A = \begin{pmatrix} a & b \\ b & c \end{pmatrix}, \vec{v} = \begin{pmatrix} d \\ d \end{pmatrix}, \vec{x} = \begin{pmatrix} x \\ y \end{pmatrix}.$$

Zuerst wird im Abschnitt 2.1 eine Orthonormalbasis aus Eigenvektoren der Matrix A bestimmt. Drückt man die Funktion q in den Koordinaten bezüglich dieser Basis aus, so treten keine gemischten quadratischen Terme mehr auf. Von den Vorzeichen der Eigenwerte der Matrix A hängt ab, ob die Lösungsmenge der Gleichung $q(x, y) = 0$ eine Ellipse (Abschnitt 2.2), eine Hyperbel (Abschnitt 2.3) oder eine Parabel (Abschnitt 2.4) ist. In jedem der drei Fälle kann durch eine Verschiebung bzw. quadratisch Ergänzung die Lösungsmenge zentriert werden.

2.1. Hauptachsen der Matrix A . Das charakteristische Polynom der Matrix A ist

$$\det(A - \lambda E) = \begin{vmatrix} a - \lambda & b \\ b & c - \lambda \end{vmatrix} = (a - \lambda)(c - \lambda) - b^2 = \lambda^2 - (a + c)\lambda + ac - b^2.$$

Es hat die reellen Nullstellen

$$\lambda_{1,2} = \frac{a + c}{2} \pm \sqrt{\frac{(a + c)^2}{4} - ac + b^2} = \frac{a + c \pm \sqrt{(a - c)^2 + 4b^2}}{2}$$

und die Eigenräume

$$\begin{aligned} E_{\lambda_1} &= \text{span } \vec{w}_1 \text{ mit } \vec{w}_1 = \left(c - a - \sqrt{(a - c)^2 + 4b^2}, -2b \right), \\ E_{\lambda_2} &= \text{span } \vec{w}_2 \text{ mit } \vec{w}_2 = \left(c - a + \sqrt{(a - c)^2 + 4b^2}, -2b \right). \end{aligned}$$

Die Geraden $\mathbb{R}\vec{w}_1$ und $\mathbb{R}\vec{w}_2$ sind die **Hauptachsen** (Symmetrieachsen) der quadratischen Form $\vec{x}^T A \vec{x}$.

Normiert man die Vektoren $\vec{w}_1$ und $\vec{w}_2$, so erhält man eine Orthonormalbasis $\{\vec{v}_1, \vec{v}_2\}$ des $\mathbb{R}^2$. Die Beziehung zwischen den Koordinaten eines Vektors $\vec{x} \in \mathbb{R}^2$ mit $\vec{x} = x\vec{e}_1 + y\vec{e}_2 = \xi\vec{v}_1 + \eta\vec{v}_2$ bezüglich der $\{\vec{v}_1, \vec{v}_2\}$ bzw. der Standardbasis lässt sich mit Hilfe der Matrix $C := (\vec{v}_1 \vec{v}_2)$ beschreiben:

$$(163) \quad \begin{pmatrix} x \\ y \end{pmatrix} = C \begin{pmatrix} \xi \\ \eta \end{pmatrix}$$

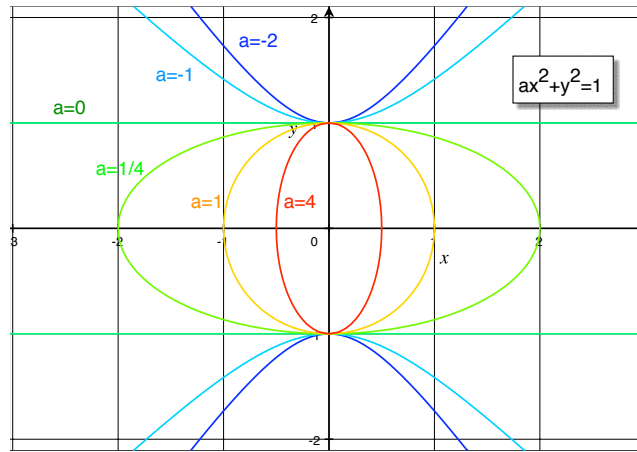


ABBILDUNG 2

In Koordinaten (ξ, η) bezüglich einer Orthonormalbasis aus Eigenvektoren hat die quadratische Funktion q folgende Form:

$$\begin{aligned} q(x, y) &= \vec{x}^T A \vec{x} + \vec{v}^T \vec{x} + f = (\xi, \eta) C^T A C \begin{pmatrix} \xi \\ \eta \end{pmatrix} + \vec{v} C \begin{pmatrix} \xi \\ \eta \end{pmatrix} + f \\ &= (\xi, \eta) \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \begin{pmatrix} \xi \\ \eta \end{pmatrix} + \vec{w} \begin{pmatrix} \xi \\ \eta \end{pmatrix} + f \text{ mit } \vec{w} := \vec{v} C = \begin{pmatrix} \alpha \\ \beta \end{pmatrix} \\ &= \lambda_1 \xi^2 + \lambda_2 \eta^2 + \alpha \xi + \beta \eta + f \end{aligned}$$

Es gilt

$$(164) \quad \lambda_1 \lambda_2 = \frac{1}{4} ((a+c)^2 - ((a-c)^2 + 4b^2)) = b^2 - ac.$$

BEISPIEL 231. Die Hauptachsen der quadratischen Gleichung $ax^2 + y^2 = 1$ sind die x - und die y -Achse. Abbildung 2 zeigt, wie die Verkleinerung des Koeffizienten vor x^2 die Ellipse ($a > 0$) in Richtung der x -Achse streckt bis aus der Ellipse zwei parallele Geraden ($a = 0$) und dann zwei Hyperbelarme ($a < 0$) werden.

BEISPIEL 232. Die Hauptachsen hängen von den Koeffizienten a , b und c ab. Hält man a und c fest, zum Beispiel $a = 1/4$ und $c = 1$, so kann man in Abbildung 3 beobachten, wie die Änderung des Parameters b die Hauptachsen dreht und die Eigenwerte beeinflusst. Es können Ellipsen, parallele Geraden oder Hyperbeln als Lösungsmenge auftreten.

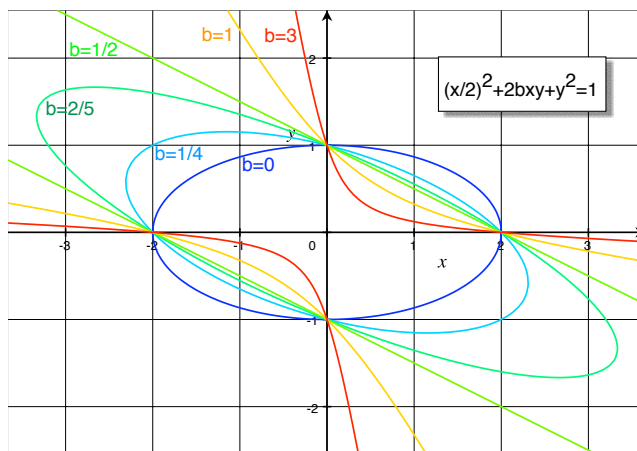


ABBILDUNG 3

2.2. Ellipsen. Falls $\lambda_1 \lambda_2 > 0$, wenn also beide Eigenwerte das gleiche Vorzeichen haben, dann gilt

$$\begin{aligned}
 q(\xi, \eta) &= \lambda_1 \lambda_2 \left(\frac{\xi^2}{\lambda_2} + \frac{\eta^2}{\lambda_1} + \frac{\alpha}{\lambda_1 \lambda_2} \xi + \frac{\beta}{\lambda_1 \lambda_2} \eta + \frac{f}{\lambda_1 \lambda_2} \right) \\
 \frac{q(\xi, \eta)}{\lambda_1 \lambda_2} &= \left(\frac{\xi}{\sqrt{\pm \lambda_2}} \pm \frac{\alpha}{2\lambda_1 \sqrt{\pm \lambda_2}} \right)^2 + \left(\frac{\eta}{\sqrt{\pm \lambda_1}} \pm \frac{\beta}{2\sqrt{\pm \lambda_1 \lambda_2}} \right)^2 + \gamma \\
 &= \left(\frac{\xi \pm \frac{\alpha}{2\lambda_1}}{\sqrt{\pm \lambda_2}} \right)^2 + \left(\frac{\eta \pm \frac{\beta}{2\lambda_2}}{\sqrt{\pm \lambda_1}} \right)^2 + \gamma \\
 \text{mit } \gamma &:= \frac{4f\lambda_1\lambda_2 \mp \lambda_2\alpha^2 \mp \lambda_1\beta^2}{4\lambda_1^2\lambda_2^2}
 \end{aligned}$$

Die Gleichung $q(x, y) = 0$ besitzt genau dann eine Lösung, wenn $\gamma \leq 0$. Die Lösungsmenge ist in (ξ, η) -Koordinaten eine Ellipse mit Mittelpunkt

$$\left(\mp \frac{\alpha}{2\lambda_1}, \mp \frac{\beta}{2\lambda_2} \right)$$

und Hauptachsen der Länge $\sqrt{-\gamma\lambda_2}$ und $\sqrt{-\gamma\lambda_1}$. Für $\gamma = 0$ entartet die Ellipse zu einem Punkt.

BEISPIEL 233. Die quadratische Gleichung $x^2 - xy + y^2 + dx = 1$ besitzt zwei positive Eigenwerte. Abbildung 4 zeigt wie die Veränderung des Parameters d im linearen Term dx die Ellipse in x - und y -Richtung verschiebt und die Längen der Hauptachsen beeinflusst.

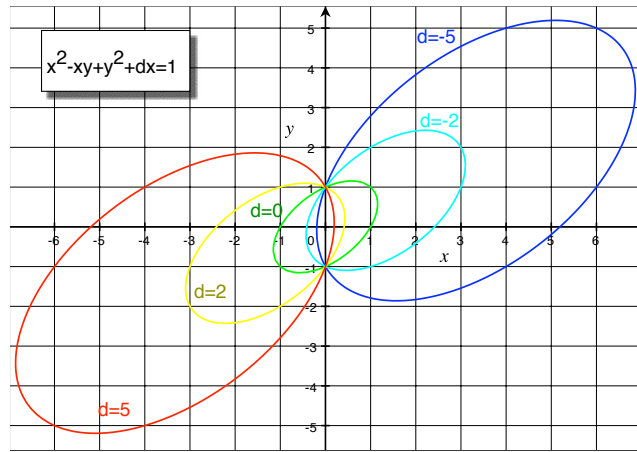


ABBILDUNG 4

2.3. Hyperbeln. Falls $\lambda_1 \lambda_2 < 0$, wenn also beide Eigenwerte verschiedene Vorzeichen haben, dann gilt

$$\begin{aligned}
 q(\xi, \eta) &= \lambda_1 \lambda_2 \left(\frac{\xi^2}{\lambda_2} - \frac{\eta^2}{-\lambda_1} + \frac{\alpha}{(\lambda_1) \lambda_2} \xi - \frac{\beta}{(-\lambda_1) \lambda_2} \eta + \frac{f}{\lambda_1 \lambda_2} \right) \\
 \frac{q(\xi, \eta)}{\lambda_1 \lambda_2} &= \pm \left(\frac{\xi}{\sqrt{\pm \lambda_2}} + \frac{\alpha}{2\lambda_1 \sqrt{\pm \lambda_2}} \right)^2 \mp \left(\frac{\eta}{\sqrt{\mp \lambda_1}} + \frac{\beta}{2\sqrt{\mp \lambda_1} \lambda_2} \right)^2 + \gamma \\
 &= \pm \left(\frac{\xi + \frac{\alpha}{2\lambda_1}}{\sqrt{\pm \lambda_2}} \right)^2 \mp \left(\frac{\eta + \frac{\beta}{2\lambda_2}}{\sqrt{\mp \lambda_1}} \right)^2 + \gamma \\
 \text{mit } \gamma &:= \frac{4f\lambda_1 \lambda_2 - \lambda_2 \alpha^2 - \lambda_1 \beta^2}{4\lambda_1^2 \lambda_2^2}
 \end{aligned}$$

Die Gleichung $q(x, y) = 0$ besitzt für alle $\gamma \in \mathbb{R}$ eine Lösung. Die Lösungsmenge besteht in (ξ, η) -Koordinaten aus zwei Hyperbelarmen mit Mittelpunkt

$$M := \left(-\frac{\alpha}{2\lambda_1}, -\frac{\beta}{2\lambda_2} \right).$$

Für $\gamma = 0$ entarten die beiden Hyperbelarme zu zwei Geraden, die sich im Punkt M schneiden.

2.4. Parabeln. Falls $\lambda_1 \lambda_2 = 0$, also $\lambda_1 = 0$ oder $\lambda_2 = 0$, so kann man durch Vertauschung der Variablen und Vorzeichenwechsel in q erreichen,

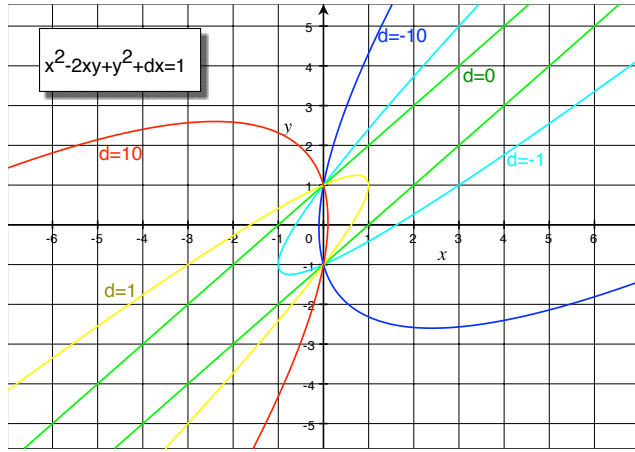


ABBILDUNG 5

dass $\lambda_1 = 0$ und $\lambda_2 > 0$. Dann folgt

$$q(\xi, \eta) = \lambda_2^2 \left(\frac{\eta^2}{\lambda_2} + \frac{\beta}{\lambda_2^2} \eta + \frac{\alpha}{\lambda_2^2} \xi + \frac{f}{\lambda_2^2} \right)$$

$$\frac{q(\xi, \eta)}{\lambda_2^2} = \left(\frac{\eta + \frac{\beta}{2\lambda_2}}{\sqrt{\lambda_2}} \right)^2 + \frac{4\alpha\xi + 4f - \beta^2}{4\lambda_2^2}.$$

Die Lösungsmenge der Gleichung $q(x, y) = 0$ ist in (ξ, η) -Koordinaten für $\alpha \neq 0$ eine Parabel mit Scheitelpunkt

$$\left(\frac{\beta^2 - 4f}{4\alpha}, -\frac{\beta}{2\lambda_2} \right).$$

Für $\alpha = 0$ entartet die Lösungsmenge zu zwei parallelen Geraden.

BEISPIEL 234. Ein Eigenwert der quadratischen Gleichung $x^2 - 2xy + y^2 + dx = 1$ ist Null. Deshalb erhält man für $d = 0$ als Lösungsmenge zwei zu einer Hauptachse parallele Geraden. Abbildung 5 zeigt wie die Veränderung des Parameters d im linearen Term dx die parallelen Geraden zu unterschiedlich stark geöffneten Parabeln deformiert und dabei den Scheitelpunkt in x - und y -Richtung verschiebt.

Fourierreihen und Orthonormalsysteme

1. Fourierreihen

1.1. Ziel. Man möchte sich regelmäßig wiederholende (periodische) Funktionen durch natürliche/schöne periodische Funktionen (Sinus- und Kosinusfunktionen) approximieren (annähern). Zu einer gegebenen periodischen Funktion $f(x)$ suchen wir eine Reihenentwicklung

$$\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$$

mit reellen Koeffizienten $a_n, b_n \in \mathbb{R}$.

1.2. Periodische Funktionen.

DEFINITION 1. Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **L -periodisch**, ($L \in \mathbb{R}$, $L > 0$), wenn $f(x + L) = f(x)$ für alle $x \in \mathbb{R}$.

BEISPIEL 235. Die Funktionen $\sin x$ und $\cos x$ sind 2π -periodisch, denn $\sin(x + 2\pi) = \sin x \cos(2\pi) + \sin(2\pi) \cos x = \sin x$ für alle $x \in \mathbb{R}$, da $\sin(2\pi) = 0$ und $\cos 2\pi = 1$.

BEISPIEL 236. Die Funktionen $\sin kx$ und $\cos kx$ sind $2\pi/k$ -periodisch für jedes $k \neq 0$.

LEMMA 28. Eine Funktion f ist genau dann L -periodisch, wenn die Funktion

$$g : x \mapsto f\left(\frac{L}{2\pi}x\right)$$

eine 2π -periodische Funktion ist.

BEWEIS. Wenn $f(x + L) = f(x)$ für alle $x \in \mathbb{R}$, dann

$$g(x + 2\pi) = f\left(\frac{L}{2\pi}(x + 2\pi)\right) = f\left(\frac{L}{2\pi}x + L\right) = f\left(\frac{L}{2\pi}x\right) = g(x)$$

für alle $x \in \mathbb{R}$. Wenn $g(x + 2\pi) = g(x)$ für alle $x \in \mathbb{R}$, dann

$$f(x + L) = f\left(\frac{L}{2\pi}\left(x \frac{2\pi}{L} + 2\pi\right)\right) = g\left(x \frac{2\pi}{L} + 2\pi\right) = g\left(x \frac{2\pi}{L}\right) = f(x)$$

für alle $x \in \mathbb{R}$. □

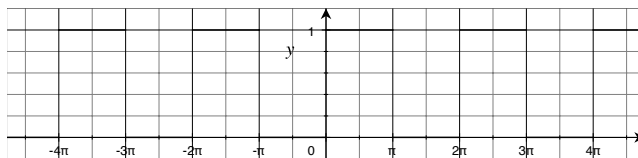


ABBILDUNG 1. 2π -periodische Fortsetzung von $f(x) = x$ für $x \in [-\pi, \pi)$

Dies bedeutet, dass man durch eine einfache Koordinatentransformation, $x = 2\pi y/L$, eine L -periodische Funktion in eine 2π -periodische Funktion umwandeln kann.

Jede L -periodische Funktion ist für alle $k \in \mathbb{N}$ auch kL -periodisch. Deshalb ist es sinnvoll, die kleinste Periode einer periodischen Funktion zu bestimmen.

BEISPIEL 237 (Zweiweg-gleichgerichteter Sinus). Die kleinste Periode der Funktion $|\sin x|$ ist π , denn $|\sin(x + \pi)| = |\sin x \cos \pi + \sin \pi \cos x| = |-\sin x| = |\sin x|$, da $\sin \pi = 0$ und $\cos \pi = -1$. $L = \pi$ ist die kleinste Periode der Funktion $|\sin x|$, weil die Menge $\pi\mathbb{Z}$ der Nullstellen der Funktion $|\sin x|$ keine kleinere Periode besitzt.

BEISPIEL 238. Die kleinste Periode der Funktion $f(x) = \sin(3x) \cos(2x)$ ist 2π . Da $\sin(3x)$ und $\cos(2x)$ die Periode 2π haben, ist auch f eine 2π -periodische Funktion. Gibt es noch eine kleinere Periode? Wir betrachten die Nullstellen von f . Es gilt $f(x) = 0$ genau dann, wenn $x = k\pi/3$ oder $x = \pi/4 + k\pi/2$. Die Nullstellen haben die kleinste Periode π . Wir prüfen $f(x + \pi) = \sin(3(x + \pi)) \cos(2(x + \pi)) = \sin(3x + 3\pi) \cos(2x + 2\pi) = -\sin(3x) \cos(2x) \neq f(x)$.

DEFINITION 2. Die **L -periodische Fortsetzung** einer auf einem Intervall $[a, b)$ oder $(a, b]$ der Länge L , d.h. $b - a = L$, definierten Funktion f ist gegeben durch $x + kL \mapsto f(x)$ für $x \in [a, b)$ und $k \in \mathbb{Z}$.

BEISPIEL 239 (Rechteckimpuls). Der Rechteckimpuls

$$f : [-\pi, \pi) \rightarrow \mathbb{R}, \quad f(x) = \begin{cases} 1 & x \geq 0 \\ 0 & x < 0 \end{cases}$$

hat die 2π -periodische Fortsetzung (wie in Abbildung 1)

$$f(x + 2k\pi) = \begin{cases} 1 & 0 \leq x < \pi, k \in \mathbb{Z} \\ 0 & -\pi \leq x < 0, k \in \mathbb{Z} \end{cases}.$$

BEISPIEL 240 (Sägezahnkurve). Die Funktion $f(x) : [-\pi, \pi) \rightarrow \mathbb{R}$, $f(x) = x$ hat die 2π -periodische Fortsetzung $f(x + 2k\pi) = x$ für $x \in [-\pi, \pi)$, $k \in \mathbb{Z}$ (Abbildung 2).

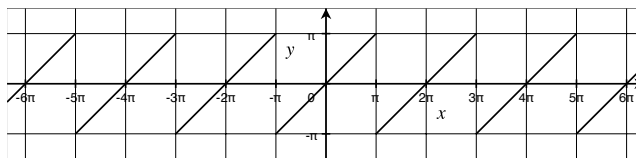


ABBILDUNG 2. 2π -periodische Fortsetzung von $f(x) = x$ für $x \in [-\pi, \pi)$

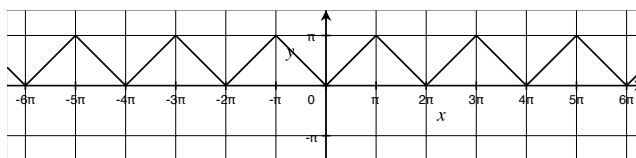


ABBILDUNG 3. 2π -periodische Fortsetzung von $f(x) = |x|$ für $x \in [-\pi, \pi]$

BEISPIEL 241 (Dreieckskurve). Die 2π -periodische Fortsetzung der Funktion $f(x) : [-\pi, \pi) \rightarrow \mathbb{R}$, $f(x) = |x|$ ist $f(x + 2k\pi) = |x|$ für $x \in [-\pi, \pi)$, $k \in \mathbb{Z}$ (Abbildung 3).

1.3. Warum Fourierreihen? Wir kennen bereits die Taylorreihenentwicklung als eine Methode, um eine Funktionen f zu approximieren. Hier sollen Taylorreihen und Fourierreihen verglichen werden.

Um eine Funktion f in eine Taylorreihe $\sum_{n=0}^{\infty} a_n(x - x_0)^n$ um einen Punkt x_0 entwickeln zu können, muss f in x_0 glatt sein. Die Koeffizienten a_n hängen nur von den Funktionswerten $f(x)$ in der Nähe von x_0 ab, da $a_n = f^{(n)}(x_0)/n!$. Es gibt einen Konvergenzradius R , so dass die Taylorreihe für alle x mit $|x - x_0| < R$ konvergiert. Auch wenn der Konvergenzradius R positiv ist, kann es passieren, dass die Funktion f nur für $x = x_0$ durch ihre Taylorreihe dargestellt wird. Taylorreihen sind insbesondere sehr hilfreich, wenn man eine Funktion lokal verstehen will und die auftretenden Restglieder (Fehler) gut abschätzen kann. Die Grundfunktionen x^n sind algebraische Funktionen. Ihre Funktionswerte lassen sich leicht berechnen.

Weder der unstetige Rechteckimpuls und die unstetige Sägezahnkurve noch die in $x = 0$ nicht differenzierbare Dreieckskurve lassen sich auf dem Intervall $[-\pi, \pi]$ gut durch eine Taylorreihe annähern. Eine Taylorentwicklung wäre um $x = 0$ bzw. $x = \pm\pi$ gar nicht möglich. Die Taylorentwicklung dieser Funktionen um einen glatten Punkt würde die Funktionen immer nur auf einem Teilintervall darstellen. Zum Beispiel ist $1 + (x - 1)$ die Taylorreihe der Dreieckskurve um $x = 1$. Sie stimmt nur auf $[0, \pi]$ mit der Dreieckskurve überein.

Um eine Funktion f in eine Fourierreihe $a_0/2 + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$ entwickeln zu können, muss f integrierbar und periodisch sein. Die Koeffizienten a_n, b_n hängen von fast allen Funktionswerten ab, da sie

durch Integration bestimmt werden. Wenn die Funktion f sogar stückweise stetig differenzierbar ist, so konvergiert ihre Fourierreihe fast alle x gegen $f(x)$. Wir benutzen Fourierreihen, um global eine möglichst gute Näherung zu finden. Die Grundfunktionen $\sin(nx)$ und $\cos(nx)$ sind transzendente Funktionen, d.h. ihre Funktionswerte lassen sich schwer berechnen. Aber als Schwingungen sind sie technisch gut realisierbar und ihre Ableitungseigenschaften sind auch bequem.

1.4. Fourierkoeffizienten und Orthogonalitätsrelationen.

SATZ 121. Wenn die Funktionenreihe

$$\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$$

gleichmäßig gegen eine Funktion $f(x)$ konvergiert, dann gilt für die Koeffizienten

$$(165) \quad a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos(nx) dx, \quad b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx.$$

BEWEIS. Wir betrachten das Integral $\int_{-\pi}^{\pi} f(x) \cos(kx) dx$ mit $k \in \mathbb{N}$. Da die Funktionenreihe gleichmäßig konvergiert, können wir Summation und Integration vertauschen.

$$\begin{aligned} \int_{-\pi}^{\pi} f(x) \cos(kx) dx &= \int_{-\pi}^{\pi} \left(\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx) \right) \cos(kx) dx \\ &= \frac{a_0}{2} \int_{-\pi}^{\pi} \cos(kx) dx \\ &\quad + \sum_{n=1}^{\infty} a_n \int_{-\pi}^{\pi} \cos(nx) \cos(kx) dx + b_n \int_{-\pi}^{\pi} \sin(nx) \cos(kx) dx \\ &= \begin{cases} \frac{a_0}{2} 2\pi & (k=0) \\ a_k \pi & (k \neq 0) \end{cases} = \pi a_k, \end{aligned}$$

denn

$$\int_{-\pi}^{\pi} \cos(kx) dx = \frac{1}{k} [\sin(kx)]_{-\pi}^{\pi} = 0,$$

für alle $k \in \mathbb{N}$

$$\begin{aligned} \int_{-\pi}^{\pi} \sin(nx) \cos(kx) dx &= \int_{-\pi}^0 \sin(nx) \cos(kx) dx + \int_0^{\pi} \sin(nx) \cos(kx) dx \\ &= \int_{\pi}^0 -\sin(-ny) \cos(-ky) dy + \int_0^{\pi} \sin(nx) \cos(kx) dx \\ &= -\int_0^{\pi} \sin(ny) \cos(ky) dy + \int_0^{\pi} \sin(nx) \cos(kx) dx = 0 \end{aligned}$$

mit der Substitution $y = -x$, $-dy = dx$ im ersten Summanden,

$$\begin{aligned}
 \int_{-\pi}^{\pi} \cos(nx) \cos(kx) dx &= \frac{1}{2} \int_{-\pi}^{\pi} \cos((n+k)x) + \cos((n-k)x) dx \\
 &= \frac{1}{2} \int_{-\pi}^{\pi} \cos((n+k)x) dx + \frac{1}{2} \int_{-\pi}^{\pi} \cos((n-k)x) dx \\
 &= \frac{1}{2} \frac{1}{n+k} [\sin((n+k)x)]_{-\pi}^{\pi} \\
 &\quad + \frac{1}{2} \begin{cases} 2\pi & (n=k) \\ \frac{1}{2} \frac{1}{n-k} [\sin((n-k)x)]_{-\pi}^{\pi} & (n \neq k) \end{cases} \\
 &= \begin{cases} \pi & (n=k) \\ 0 & (n \neq k) \end{cases},
 \end{aligned}$$

da $k, n \in \mathbb{N}$ und $n > 0$.

Nun betrachten wir das Integral $\int_{-\pi}^{\pi} f(x) \sin(kx) dx$ mit $k \in \mathbb{N}$, $k \neq 0$, um die Koeffizienten b_n zu bestimmen.

$$\begin{aligned}
 \int_{-\pi}^{\pi} f(x) \sin(kx) dx &= \int_{-\pi}^{\pi} \left(\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx) \right) \sin(kx) dx \\
 &= \frac{a_0}{2} \int_{-\pi}^{\pi} \sin(kx) dx \\
 &\quad + \sum_{n=1}^{\infty} a_n \int_{-\pi}^{\pi} \cos(nx) \sin(kx) dx + b_n \int_{-\pi}^{\pi} \sin(nx) \sin(kx) dx \\
 &= \pi b_k,
 \end{aligned}$$

da $\int_{-\pi}^{\pi} \cos(nx) \sin(kx) dx = 0$ für alle $n \in \mathbb{N}$, $\int_{-\pi}^{\pi} \sin(kx) dx = 0$ für alle $k \neq 0$ und

$$\begin{aligned}
 \int_{-\pi}^{\pi} \sin(nx) \sin(kx) dx &= \frac{1}{2} \int_{-\pi}^{\pi} -\cos((n+k)x) + \cos((n-k)x) dx \\
 &= -\frac{1}{2} \int_{-\pi}^{\pi} \cos((n+k)x) dx + \frac{1}{2} \int_{-\pi}^{\pi} \cos((n-k)x) dx \\
 &= -\frac{1}{2} \frac{1}{n+k} [\sin((n+k)x)]_{-\pi}^{\pi} \\
 &\quad + \frac{1}{2} \begin{cases} 2\pi & (n=k) \\ \frac{1}{2} \frac{1}{n-k} [\sin((n-k)x)]_{-\pi}^{\pi} & (n \neq k) \end{cases} \\
 &= \begin{cases} \pi & (n=k) \\ 0 & (n \neq k) \end{cases},
 \end{aligned}$$

da $k, n \in \mathbb{N}$ und $n > 0$. □

Die Beziehungen

$$(166) \quad \frac{1}{\pi} \int_{-\pi}^{\pi} (\sin(nx))^2 dx = \frac{1}{\pi} \int_{-\pi}^{\pi} (\cos(nx))^2 dx = 1 \quad (n \neq 0),$$

$$(167) \quad \frac{1}{\pi} \int_{-\pi}^{\pi} \cos(kx) \sin(nx) dx = 0 \quad (n, k \in \mathbb{N}, n > 0 \text{ oder } k > 0)$$

heißen **Orthogonalitätsrelationen**.

DEFINITION 3. Die **Fourierreihe** einer stückweise stetigen, 2π -periodischen Funktion $f(x)$ ist die Reihe $a_0/2 + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$ mit den **Fourierkoeffizienten**

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos(nx) dx, \quad b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx.$$

Wir schreiben

$$f(x) \sim \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$$

um anzudeuten, dass es sich um die Fourierreihe der Funktion f handelt.

Warnung: Bisher wissen wir weder ob die Fourierreihe konvergiert, noch ob ihr Grenzwert, falls er existiert, mit der Funktion f übereinstimmt!

BEISPIEL 242 (Polynome in $\sin x$ und $\cos x$). Die konstante Funktion $f(x) \equiv c \in \mathbb{R}$ ist bereits eine Fourierreihe ($a_0 = 2c$). Die Fourierreihe der Funktion $f(x) = \sin^2 x$ erhalten wir mit Hilfe der Additionstheoreme. Da $\cos 2x = \cos^2 x - \sin^2 x$, gilt

$$\sin^2 x = \frac{1}{2}(1 - \cos(2x)) = \frac{1}{2} - \frac{1}{2} \cos(2x).$$

BEISPIEL 243 (Rechteckimpuls). Wir berechnen die Fourierreihe der 2π -periodischen Fortsetzung der Funktion

$$f: [-\pi, \pi) \rightarrow \mathbb{R}, \quad f(x) = \begin{cases} 1 & (x \geq 0) \\ 0 & (x < 0) \end{cases}.$$

Es gilt

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos(nx) dx = \frac{1}{\pi} \int_0^{\pi} \cos(nx) dx = \begin{cases} 1 & (n = 0) \\ \frac{1}{\pi n} [\sin(nx)]_0^{\pi} = 0 & (n \neq 0) \end{cases}$$

und

$$\begin{aligned} b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx = \frac{1}{\pi} \int_0^{\pi} \sin(nx) dx = -\frac{1}{\pi n} [\cos(nx)]_0^{\pi} = -\frac{1}{\pi n} ((-1)^n - 1) \\ &= \begin{cases} 0 & (n \text{ gerade}) \\ \frac{2}{n\pi} & (n \text{ ungerade}) \end{cases}. \end{aligned}$$

damit ist die Fourierreihe von f die Reihe

$$\chi_{[0,\pi)}(x) \sim \frac{1}{2} + \sum_{m=0}^{\infty} \frac{2}{\pi} \frac{\sin((2m+1)x)}{2m+1} \text{ auf } [-\pi, \pi].$$

1.5. Gerade und ungerade Funktionen. Wenn man die Symmetrien der Funktion beachtet, kann man ihre Fourierreihe schneller berechnen.

DEFINITION 4. Eine auf $\mathbb{R}$ definierte Funktion f heißt **gerade**, wenn $f(x) = f(-x)$ für alle $x \in \mathbb{R}$. Sie heißt **ungerade**, wenn $f(-x) = -f(x)$ für alle $x \in \mathbb{R}$.

An den Symmetrien des Graphen einer Funktion kann man leicht erkennen, ob die Funktion gerade oder ungerade ist. Der Graph einer geraden Funktion ist symmetrisch bezüglich Spiegelung an der y -Achse. Der Graph einer ungeraden Funktion ist symmetrisch bezüglich Drehung um den Ursprung um 180° .

BEISPIEL 244. Gerade 2π -periodische Funktionen sind $\cos(nx)$, Dreieckskurve, $|\sin x|$.

BEISPIEL 245. Ungerade 2π -periodische Funktionen sind $\sin(nx)$, Sägezahnkurve und die 2π -periodische Fortsetzung der Funktion

$$f : [-\pi, \pi) \rightarrow \mathbb{R}, \quad f(x) = 2\chi_{[0,\pi)} - 1 = \begin{cases} 1 & (x \geq 0) \\ -1 & (x < 0) \end{cases}.$$

LEMMA 29. Wenn f eine ungerade integrierbare Funktion ist, dann gilt

$$\int_{-a}^a f(x) dx = 0$$

für alle $a > 0$.

BEWEIS. Es gilt

$$\begin{aligned} \int_{-a}^a f(x) dx &= \int_{-a}^0 f(x) dx + \int_0^a f(x) dx = \int_a^0 -f(-y) dy + \int_0^a f(x) dx \\ &= \int_a^0 f(y) dy + \int_0^a f(x) dx = - \int_0^a f(y) dy + \int_0^a f(x) dx = 0 \end{aligned}$$

mit $x = -y$, $dx = -dy$, da $f(-y) = -f(y)$ und sich bei der Vertauschung der Integrationsgrenzen das Vorzeichen des Integrals ändert. $\square$

FOLGERUNG 39. Es sei $\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$ die Fourierreihe einer stückweise stetigen, 2π -periodischen Funktion. Wenn f ungerade ist, dann $a_n = 0$ für alle $n \in \mathbb{N}$. Wenn f gerade ist, dann $b_n = 0$ für alle $n \in \mathbb{N}$.

BEWEIS. Wenn f ungerade ist, dann ist $f(x) \cos(nx)$ eine ungerade Funktion. Wenn f gerade ist, dann ist $f(x) \sin(nx)$ eine ungerade Funktion. $\square$

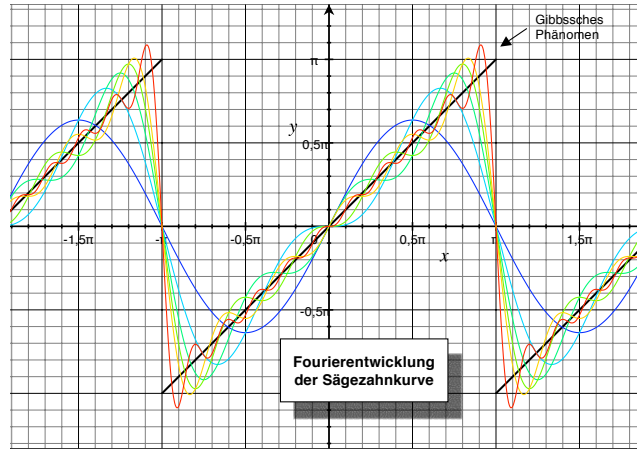


ABBILDUNG 4. Fourierapproximation der 2π -periodischen Fortsetzung von $f(x) = x$ für $x \in (-\pi, \pi]$

BEISPIEL 246 (Fourierreihe der Sägezahnkurve). Die Funktion $f : (-\pi, \pi) \rightarrow \mathbb{R}$, $f(x) = x$ ist ungerade. Darum gilt für die Fourierkoeffizienten $a_n = 0$ und

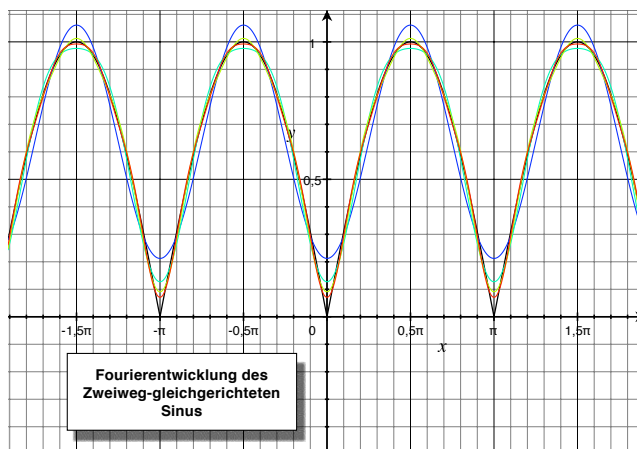
$$\begin{aligned} b_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} x \sin(nx) dx = \frac{1}{\pi} \left(\left[-\frac{x}{n} \cos(nx) \right]_{-\pi}^{\pi} + \frac{1}{n} \int_{-\pi}^{\pi} \cos(nx) dx \right) \\ &= \frac{1}{\pi} \left((-1)^{n+1} \frac{2\pi}{n} + 0 \right) = 2 \frac{(-1)^{n+1}}{n}, \end{aligned}$$

wegen der Orthogonalitätsrelation $\int_{-\pi}^{\pi} \cos(nx) dx = 0$. Die Fourierreihe der 2π -periodischen Fortsetzung von $f(x) = x$ ist also die Reihe

$$x \sim 2 \sum_{n=1}^{\infty} (-1)^{n+1} \frac{\sin(nx)}{n} \text{ auf } (-\pi, \pi).$$

BEISPIEL 247 (Fourierreihe des Zweiweg-gleichgerichteten Sinus). Da die Funktion $f(x) = |\sin x|$ gerade ist, gilt für die Fourierkoeffizienten $b_n = 0$ und

$$\begin{aligned} a_n &= \frac{1}{\pi} \int_{-\pi}^{\pi} |\sin x| \cos(nx) dx = \frac{2}{\pi} \int_0^{\pi} |\sin x| \cos(nx) dx = \frac{2}{\pi} \int_0^{\pi} \sin x \cos(nx) dx \\ &= \frac{2}{\pi} \frac{1}{n^2 - 1} [\cos x \cos(nx) + n \sin x \sin(nx)]_0^{\pi} \\ &= \frac{2}{\pi} \frac{1}{n^2 - 1} ((-1)^{n+1} - 1) = \begin{cases} 0 & n \text{ ungerade} \\ -\frac{4}{\pi} \frac{1}{n^2 - 1} & n \text{ gerade} \end{cases}, \end{aligned}$$

ABBILDUNG 5. Fourierapproximation von $|\sin x|$

da

$$\begin{aligned} \int \sin x \cos(nx) dx &= -\cos x \cos(nx) - n \int \cos x \sin(nx) dx \\ &= -\cos x \cos(nx) - n \sin x \sin(nx) + n^2 \int \sin x \cos(nx) dx \\ &= \frac{1}{n^2 - 1} (\cos x \cos(nx) + n \sin x \sin(nx)). \end{aligned}$$

Die Fourierreihe von $|\sin x|$ ist

$$|\sin x| \sim \frac{2}{\pi} - \frac{4}{\pi} \sum_{m=1}^{\infty} \frac{1}{4m^2 - 1} \cos(2mx).$$

1.6. Rechenregeln - erster Teil. Wenn f und g zwei integrierbare, 2π -periodische Funktionen mit den Fourierreihen $f(x) \sim \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$ und $g(x) \sim \frac{c_0}{2} + \sum_{n=1}^{\infty} c_n \cos(nx) + d_n \sin(nx)$ sind, dann ist für alle $\lambda, \mu \in \mathbb{R}$ die Fourierreihe der Funktion $\lambda f(x) + \mu g(x)$ gerade

$$\lambda f(x) + \mu g(x) \sim \frac{\lambda a_0 + \mu c_0}{2} + \sum_{n=1}^{\infty} (\lambda a_n + \mu c_n) \cos(nx) + (\lambda b_n + \mu d_n) \sin(nx).$$

BEISPIEL 248. Die Fourierreihe der 2π -periodischen Fortsetzung der Funktion

$$g : [-\pi, \pi) \rightarrow \mathbb{R}, \quad g(x) = \begin{cases} h & (x \geq 0) \\ -h & (x < 0) \end{cases}$$

mit $h \geq 0$ ist

$$g(x) \sim \frac{4h}{\pi} \sum_{m=0}^{\infty} \frac{\sin((2m+1)x)}{2m+1},$$

denn $g(x) = 2h(f(x) - 1/2)$ mit dem Rechteckimpuls f aus Beispiel 243.

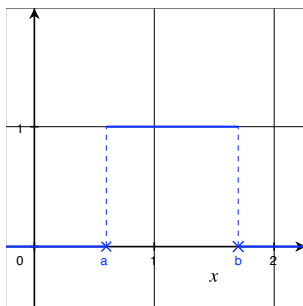


ABBILDUNG 6. Fourierapproximation des symmetrischen Rechteckimpuls

1.7. Darstellungssatz. In diesem Abschnitt zeigen wir, dass die Fourierreihe einer stückweise differenzierbaren, 2π -periodischen Funktion fast überall punktweise gegen die Ausgangsfunktion konvergiert und diese Konvergenz auf jedem abgeschlossenen Intervall ohne Sprungstellen sogar gleichmäßig ist.

DEFINITION 5. Eine Funktion $f[a, b] \rightarrow \mathbb{R}$ heißt **stückweise stetig** bzw. **stückweise stetig differenzierbar**, wenn f bis auf höchstens endlich viele Stellen in $[a, b]$ stetig bzw. stetig differenzierbar ist und in den Ausnahmestellen jeweils alle im abgeschlossenen Intervall $[a, b]$ möglichen einseitigen Grenzwerte von f bzw. f und f' existieren.

Natürlich sind alle stetigen bzw. stetig differenzierbaren Funktionen auch stückweise stetig bzw. stückweise stetig differenzierbar. Der Rechteckimpuls und die Sägezahnkurve sind Beispiele für nicht stetige, aber stückweise stetige und stückweise stetig differenzierbare Funktionen. Die Dreieckskurve, $|\sin x|$ und die 2π -periodische Fortsetzung von $f(x) = x^2$ mit $x \in [-\pi, \pi]$ sind Beispiele für stetige, aber nur stückweise stetig differenzierbare Funktionen.

Stückweise stetige Funktionen f auf einem abgeschlossenen Intervall $[a, b]$ sind beschränkt. Die Funktionswerte an den endlich vielen Sprungstellen haben auf die Berechnung des Integrals $\int_a^b f(x) dx$ keinen Einfluss.

BEISPIEL 249 (fast überall stetige, nicht stückweise stetige Funktionen). Die Funktionen $f, g : [0, 1] \rightarrow \mathbb{R}$ mit

$$f(x) = \begin{cases} \sin \frac{1}{x} & (x \neq 0) \\ 0 & (x = 0) \end{cases}, \quad g(x) = \begin{cases} \frac{1}{x} & (x \neq 0) \\ 0 & (x = 0) \end{cases}$$

sind auf dem Intervall $(0, 1]$ stetig und in $x = 0$ unstetig. Sie sind aber nicht stückweise stetig, weil der einseitige Grenzwert $\lim_{x \rightarrow 0} f(x)$ nicht existiert und $\lim_{x \rightarrow 0+} g(x) = \infty$.

LEMMA 30. Für alle $x \neq 2l\pi$ mit $l \in \mathbb{Z}$ gilt

$$\sum_{k=1}^n \cos(kx) = -\frac{1}{2} + \frac{\sin(nx + x/2)}{2 \sin(x/2)}.$$

BEWEIS. Es gilt

$$\begin{aligned} \sum_{k=0}^n \cos(kx) &= \sum_{k=0}^n \Re(e^{ikx}) = \Re\left(\sum_{k=0}^n e^{ikx}\right) = \Re\left(\sum_{k=0}^n (e^{ix})^k\right) = \Re\left(\frac{e^{ix(n+1)} - 1}{e^{ix} - 1}\right) \\ &= \Re\left(\frac{e^{i(xn+x/2)} - e^{-ix/2}}{e^{ix/2} - e^{-ix/2}}\right) = \Re\left(\frac{e^{i(xn+x/2)} - e^{-ix/2}}{2i \sin(x/2)}\right) \\ &= \frac{\sin(nx + x/2) + \sin(x/2)}{2 \sin(x/2)} = \frac{1}{2} + \frac{\sin(nx + x/2)}{2 \sin(x/2)}. \end{aligned}$$

□

LEMMA 31. Wenn $f : \mathbb{R} \rightarrow \mathbb{R}$ integrierbar und 2π -periodisch, dann gilt für das Fourierpolynom

$$T_n(x) := \frac{a_0}{2} + \sum_{k=1}^n a_k \cos(kx) + b_k \sin(kx) = \frac{2}{\pi} \int_0^{\pi/2} \frac{f(x+2t) + f(x-2t)}{2} K_n(t) dt$$

mit dem Kern

$$K_n(t) = \begin{cases} \frac{\sin(2nt+t)}{\sin t} & (t \neq 0) \\ 2n+1 & (t = 0) \end{cases}$$

BEWEIS. Wir erhalten

$$\begin{aligned} T_n(x) &= \frac{a_0}{2} + \sum_{k=1}^n a_k \cos(kx) + b_k \sin(kx) \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t) dt + \frac{1}{\pi} \sum_{k=1}^n \int_{-\pi}^{\pi} f(t) \cos(kt) \cos(kx) dt + \int_{-\pi}^{\pi} f(t) \sin(kt) \sin(kx) dt \\ &= \frac{1}{\pi} \int_{-\pi}^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(kt) \cos(kx) + \sin(kt) \sin(kx) \right) f(t) dt \\ &= \frac{1}{\pi} \int_{-\pi}^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos k(t-x) \right) f(t) dt \\ &= \frac{1}{\pi} \int_{x-\pi}^{x+\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(k(t-x)) \right) f(t) dt \\ &= \frac{1}{\pi} \int_{-\pi}^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(ks) \right) f(s+x) ds \\ &= \frac{1}{\pi} \int_{-\pi}^0 \left(\frac{1}{2} + \sum_{k=1}^n \cos(ks) \right) f(s+x) ds + \frac{1}{\pi} \int_0^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(ks) \right) f(s+x) ds \end{aligned}$$

$$\begin{aligned}
&= -\frac{1}{\pi} \int_{\pi}^0 \left(\frac{1}{2} + \sum_{k=1}^n \cos(-kt) \right) f(-t+x) dt \\
&\quad + \frac{1}{\pi} \int_0^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(kt) \right) f(t+x) dt \\
&= \frac{1}{\pi} \int_0^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(kt) \right) f(x-t) dt + \frac{1}{\pi} \int_0^{\pi} \left(\frac{1}{2} + \sum_{k=1}^n \cos(kt) \right) dt \\
&= \frac{1}{\pi} \int_0^{\pi} (f(t+x) + f(x-t)) \left(\frac{1}{2} + \sum_{k=1}^n \cos(kt) \right) dt \\
&= \frac{2}{\pi} \int_0^{\pi/2} (f(x+2s) + f(x-2s)) \left(\frac{1}{2} + \sum_{k=1}^n \cos(2ks) \right) ds \\
&= \frac{2}{\pi} \int_0^{\pi/2} (f(x+2s) + f(x-2s)) \frac{\sin(2ns+s)}{2 \sin s} ds
\end{aligned}$$

indem wir

- die Formeln für die Fourierkoeffizienten einsetzen,
- Summation und Integration vertauschen,
- $f(t)$ ausklammern,
- das Additionstheorem für $\cos(kt - kx)$ anwenden,
- die Integrationsgrenzen verschieben, weil der Integrand 2π -periodisch ist,
- die Substitution $s = t - x$, $ds = dt$ durchführen,
- das Integrationsintervall $[-\pi, \pi] = [-\pi, 0] \cup [0, \pi]$ zerlegen,
- im ersten Summanden $t = -s$ und $dt = -ds$ substituieren und im zweiten Summanden die Integrationsvariable s durch t ersetzen,
- ausnutzen, dass Kosinus eine gerade Funktion ist, und dass sich bei der Vertauschung der Integrationsgrenzen das Vorzeichen des Integrals umkehrt,
- Summe und Integration vertauschen,
- $t = 2s$ und $dt = 2ds$ substituieren,
- das Lemma 30 mit $x = 2s$ anwenden.

□

FOLGERUNG 40. *Es gilt*

$$\frac{2}{\pi} \int_0^{\pi/2} K_n(t) dt = 1.$$

BEWEIS. Wir betrachten die konstante Funktion $f(x) \equiv 1$ und wenden Lemma 31 an:

$$T_n(x) = 1 = \frac{2}{\pi} \int_0^{\pi/2} \frac{1+1}{2} K_n(t) dt = \frac{2}{\pi} \int_0^{\pi/2} K_n(t) dt.$$

□

LEMMA 32. Wenn f stückweise stetig im Intervall $[a, b]$, dann gilt

$$\lim_{n \rightarrow \infty} \int_a^b f(t) \sin(nt) dt = 0.$$

BEWEIS. Wir zerlegen das Intervall $[a, b]$ in endlich viele Intervalle, in denen f stetig ist. So reicht es, die Aussage für stetige Funktionen f zu beweisen. Zu jedem $\delta > 0$ und jedem $x \in [a, b]$ gibt es ein $\varepsilon_x > 0$ mit $|f(y) - f(x)| < \delta$ für alle y mit $|y - x| < \varepsilon_x$. Da das Intervall $[a, b]$ kompakt ist, gibt es endlich viele x_j , so dass die Vereinigung der ε_{x_j} -Umgebungen von x_j das Intervall $[a, b]$ überdecken. Damit erhalten wir

$$\begin{aligned} \int_a^b f(t) \sin(nt) dt &= \int_a^b f(t) \sin(nt) dt = \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} f(t) \sin(nt) dt \\ &= \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} (f(t) - f(x_j) + f(x_j)) \sin(nt) dt \\ &= \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} (f(t) - f(x_j)) \sin(nt) dt + \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} f(x_j) \sin(nt) dt. \end{aligned}$$

Mit der Dreiecksungleichung ergibt sich

$$\begin{aligned} \left| \int_a^b f(t) \sin(nt) dt \right| &\leq \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} |f(t) - f(x_j)| |\sin(nt)| dt \\ &\quad + \sum_j \left| \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} f(x_j) \sin(nt) dt \right| \\ &\leq \delta \sum_j \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} dt + \sum_j f(x_j) \left| \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} \sin(nt) dt \right| \\ &\leq \delta(b - a) + M \sum_j \left| \int_{x_j - \varepsilon_j}^{x_j + \varepsilon_j} \sin(nt) dt \right|, \end{aligned}$$

da $|f(t) - f(x_j)| < \delta$ auf dem jeweiligen Integrationsintegral, $|\sin t| \leq 1$ und mit $M = \max\{f(x) : x \in [a, b]\}$.

Auf jedem abgeschlossenen Intervall $[c, d]$, z.B. $[c, d] = [x_j - \varepsilon_j, x_j + \varepsilon_j]$, gilt

$$\left| \int_c^d \sin(nt) dt \right| = \left| -\frac{1}{n} [\cos(nt)]_c^d \right| \leq \frac{1}{n},$$

da $|\cos x| \leq 1$. Damit gilt

$$\left| \int_a^b f(t) \sin(nt) dt \right| \leq \delta(b - a) + \frac{M}{n}$$

und

$$\lim_{n \rightarrow \infty} \int_a^b f(t) \sin(nt) dt = \lim_{n \rightarrow \infty} \frac{M}{n} = 0,$$

da $\delta(b-a)$ beliebig klein gewählt werden kann. $\square$

SATZ 122 (Punktweise und gleichmäßige Konvergenz der Fourierreihe).
Für jede stückweise stetige, stückweise stetig differenzierbare, 2π -periodische Funktion f konvergiert die Fourierreihe von f punktweise gegen

$$\begin{cases} f(x) & , \text{ falls } f \text{ stetig in } x \\ \frac{1}{2}(f(x^+) + f(x^-)) & , \text{ falls } x \text{ Sprungstelle von } f \end{cases}$$

Die Fourierreihe konvergiert auf jedem abgeschlossenen Intervall ohne Sprungstellen gleichmäßig.

Hier bezeichnen $f(x^+)$ und $f(x^-)$ den rechts- bzw. linksseitigen Grenzwert von f an der Stelle x .

BEWEIS. Wir betrachten die Differenz zwischen Fourierpolynom $T_n(x) = \frac{a_0}{2} + \sum_{k=1}^n a_k \cos(kx) + b_k \sin(kx)$ und behauptetem Grenzwert $(f(x^+) + f(x^-))/2$:

$$T_n(x) - \frac{f(x^+) + f(x^-)}{2} = \frac{2}{\pi} \int_0^{\frac{\pi}{2}} \frac{f(x+2t) + f(x-2t) - f(x^+) - f(x^-)}{2} K_n(t) dt.$$

Der Integrand

$$\begin{aligned} & \frac{f(x+2t) + f(x-2t) - f(x^+) - f(x^-)}{2} K_n(t) \\ &= \left(\frac{f(x+2t) - f(x^+)}{2t} + \frac{f(x-2t) - f(x^-)}{2t} \right) \frac{t}{\sin t} \sin((2n+1)t) \end{aligned}$$

ist stetig auf dem Intervall $[0, \pi/2]$. Für $t \neq 0$ ist dies klar, für $t = 0$ folgt dies aus

$$\lim_{t \rightarrow 0} \frac{t}{\sin t} = 1, \text{ und } \lim_{t \rightarrow 0} \frac{f(x \pm 2t) - f(x^\pm)}{2t} = f'(x^\pm).$$

Mit Lemma 32 erhalten wir die punktweise Konvergenz

$$\lim_{n \rightarrow \infty} T_n(x) - \frac{f(x^+) + f(x^-)}{2} = 0.$$

Aus dem Beweis des Lemmas 32 folgt, dass die Konvergenz auf jedem abgeschlossenen Intervall ohne Sprungstellen von f gleichmäßig ist. $\square$

BEISPIEL 250 (Fourierreihe des Sägezahns). Die Fourierreihe der 2π -periodischen Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$, $f(x+2k\pi) = x$ mit $k \in \mathbb{Z}$ und $x \in [-\pi, \pi)$ konvergiert punktweise gegen die Funktion

$$x + 2k\pi \mapsto \begin{cases} 0 & (x = -\pi, k \in \mathbb{Z}) \\ x & (x \in (-\pi, \pi), k \in \mathbb{Z}) \end{cases}$$

Das bedeutet, dass für $x = (2k+1)\pi$ mit $k \in \mathbb{Z}$ die Fourierreihe von f nicht gegen $f(x)$ konvergiert, sondern gegen den Mittelwert des rechts- und linksseitigen Grenzwertes von f im Punkt $(2k+1)\pi$. Aber in $x = \pi/2$ konvergiert

die Fourierreihe gegen $f(x)$ und wir erhalten eine Reihenentwicklung für π durch

$$\begin{aligned} f(\pi/2) &= \frac{\pi}{2} = 2 \sum_{m=0}^{\infty} (-1)^{n+1} \frac{\sin(n\pi/2)}{n} = 2 \sum_{m=0}^{\infty} (-1)^{2m+1+1} \frac{(-1)^m}{2m+1} \\ \frac{\pi}{4} &= \sum_{m=0}^{\infty} \frac{(-1)^m}{2m+1}. \end{aligned}$$

1.8. Integration von Fourierreihen. Da die Konvergenz der Fourierreihe einer stückweise stetigen und stückweise differenzierbaren Funktion auf jedem Stetigkeitsintervall gleichmäßig ist, kann man die Fourierreihe auf diesen Intervallen gliedweise integrieren.

BEISPIEL 251 (Fourierreihe der Dreiecksfunktion). Aus Beispiel 248 wissen wir, dass die Fourierreihe

$$\frac{4}{\pi} \sum_{m=0}^{\infty} \frac{\sin(2m+1)x}{2m+1} = -\chi_{(-\pi,0)} + \chi_{(0,\pi)},$$

also punktweise gegen die 2π -periodische Funktion g mit

$$g(x+2\pi k) = \begin{cases} 1 & , x \in (0, \pi), k \in \mathbb{Z} \\ -1 & , x \in (-\pi, 0), k \in \mathbb{Z} \\ 0 & , x \in \{-\pi, 0\}, k \in \mathbb{Z} \end{cases}$$

konvergiert. Da die Konvergenz auf jedem abgeschlossenen Intervall, das keine ganzzahligen Vielfachen von π enthält, gleichmäßig ist, erhalten wir durch gliedweise Integration die Fourierreihe von $\int g(x) dx = f(x)$, wobei $f(x)$ die Dreieckskurve aus Beispiel 241 ist. Nur den Fourierkoeffizienten a_0 der Fourierreihe von f können wir nicht durch Integration der Fourierreihe von g bestimmen. Wir berechnen

$$a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} |x| dx = \frac{1}{\pi} \pi^2 = \pi.$$

Also

$$f(x) = |x| = \frac{\pi}{2} - \frac{4}{\pi} \sum_{m=0}^{\infty} \frac{\cos((2m+1)x)}{(2m+1)^2} \text{ auf } [-\pi, \pi]$$

und die Fourierreihe der Dreieckskurve f konvergiert auf ganz $\mathbb{R}$ gleichmäßig gegen f , da f stetig ist.

BEISPIEL 252 (Fourierreihe der 2π -periodischen Fortsetzung von x^2). Die Fourierreihe der 2π -periodischen Fortsetzung von $f : [-\pi, \pi] \rightarrow \mathbb{R}$, $f(x) = x^2$, erhalten wir durch Integration des Doppelten der Sägezahnkurve, da $(x^2)' = 2x$. Also

$$f(x) = \frac{1}{2\pi} \int_{-\pi}^{\pi} x^2 dx + 2 \int 2 \sum_{n=1}^{\infty} (-1)^{n+1} \frac{\sin(nx)}{n} dx = \frac{\pi^2}{3} + 4 \sum_{n=1}^{\infty} (-1)^n \frac{\cos(nx)}{n^2}.$$

Warnung: Auch auf einem abgeschlossenen Stetigkeitsintervall gilt

$$f'(x) \sim \left(\sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx) \right)' = \sum_{n=1}^{\infty} -a_n n \sin(nx) + b_n n \cos(nx)$$

nur, wenn f stetig ist. Zum Beispiel gilt zwar $x' = 1$ auf dem Intervall $(-\pi, \pi)$ aber die Ableitung der Fourierreihe der Sägezahnkurve ist

$$2 \sum_{n=1}^{\infty} (-1)^{n+1} \left(\frac{\sin nx}{n} \right)' = 2 \sum_{n=1}^{\infty} (-1)^{n+1} \cos(nx) \not\sim 1,$$

also nicht die Fourierreihe der konstanten Funktion $f(x) \equiv 1$.

SATZ 123. Wenn f L -periodisch, $f'(x)$ stückweise stetig und

$$f'(x) \sim \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\omega x) + b_n \sin(n\omega x),$$

dann

$$f(x) \sim \frac{1}{L} \int_{-L/2}^{L/2} f(x) dx + \frac{1}{\omega} \sum_{n=1}^{\infty} \frac{a_n + a_0(-1)^{n+1}}{n} \sin(n\omega x) - \frac{b_n}{n} \cos(n\omega x).$$

BEWEIS. Der Summand $\frac{1}{L} \int_{-L/2}^{L/2} f(x) dx$ ist der konstante Anteil der Fourierreihe der Funktion f . Außerdem gilt

$$\frac{a_n}{n\omega} \sin(n\omega x) = \int a_n \cos(n\omega x) dx, \quad -\frac{b_n}{n\omega} \cos(n\omega x) = \int b_n \sin(n\omega x) dx$$

und $a_0 x/2 = \int a_0/2 dx$ mit $L = 2\pi/\omega$. Die Fourierreihe der L -periodischen Fortsetzung von $x \mapsto x$ für $-L/2 \leq x < L/2$ ist

$$\frac{2}{\omega} \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} \sin(n\omega x).$$

□

1.9. Komplexe Darstellung der Fourierreihe. Wir betrachten eine komplexwertige, 2π -periodische Funktion f , also $f : \mathbb{R} \rightarrow \mathbb{C}$ und $f(x + 2\pi k) = f(x)$ für alle $x \in \mathbb{R}$ und $k \in \mathbb{Z}$. Da $f = \Re(f) + i\Im(f)$, erweitern wir die Definition einer Fourierreihe für komplexwertige Funktionen in natürlicher Weise.

DEFINITION 6. Die **Fourierreihe** einer 2π -periodischen, komplexwertigen, integrierbaren Funktion f ist

$$\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$$

mit den **Fourierkoeffizienten** $a_n, b_n \in \mathbb{C}$ gegeben durch

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos(nx) dx, \quad b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx.$$

Alle Sätze über Fourierreihen reellwertiger Funktionen, insbesondere die Konvergenzaussagen, übertragen sich auf Fourierreihen komplexwertiger Funktionen. Wir können die Fourierreihe auch mit Hilfe der Grundfunktionen e^{inx} schreiben.

SATZ 124. *Es sei*

$$f \sim \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$$

die Fourierreihe einer 2π -periodischen, integrierbaren Funktion $f : \mathbb{R} \rightarrow \mathbb{C}$. Dann gilt

$$a_n \cos(nx) + b_n \sin(nx) = c_n e^{inx} + c_{-n} e^{-inx} \text{ mit } c_k = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-ikx} dx \quad \forall k \in \mathbb{Z}.$$

BEWEIS. Es gilt

$$\begin{aligned} \cos(nx) &= \Re(e^{inx}) = \frac{1}{2}(e^{inx} + e^{-inx}) \\ \sin(nx) &= \Im(e^{inx}) = \frac{1}{2i}(e^{inx} - e^{-inx}). \end{aligned}$$

Also erhalten wir

$$\begin{aligned} a_n \cos(nx) + b_n \sin(nx) &= \frac{a_n}{2}(e^{inx} + e^{-inx}) - \frac{ib_n}{2}(e^{inx} - e^{-inx}) \\ &= \frac{a_n - ib_n}{2} e^{inx} + \frac{a_n + ib_n}{2} e^{-inx} \\ \frac{a_n - ib_n}{2} &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x)(\cos(nx) - i \sin(nx)) dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx \\ \frac{a_n + ib_n}{2} &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x)(\cos(nx) + i \sin(nx)) dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{inx} dx \end{aligned}$$

□

DEFINITION 7. Die **komplexe Darstellung** der Fourierreihe einer 2π -periodischen, integrierbaren Funktion f ist die Reihe

$$\sum_{n=-\infty}^{\infty} c_n e^{inx} = c_0 + \sum_{n=1}^{\infty} (c_n e^{inx} + c_{-n} e^{-inx}) \text{ mit } c_k = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-ikx} dx \quad \forall k \in \mathbb{Z}.$$

Es gelten für alle $k \in \mathbb{N}$ die Beziehungen

$$(168) \quad c_k = \frac{a_k - ib_k}{2}, \quad c_{-k} = \frac{a_k + ib_k}{2}, \quad a_k = 2\Re(c_k), \quad b_k = -2\Im(c_k).$$

BEISPIEL 253. Wir berechnen die komplexe und die reelle Form der Fourierreihe der 2π -periodischen Fortsetzung der Funktion $f : (-\pi, \pi) \rightarrow \mathbb{R}$,

$f(x) = e^{\alpha x}$, mit $\alpha \in \mathbb{R} \setminus \{0\}$. Da f auf dem Intervall $(-\pi, \pi)$ stetig ist, gilt dort $e^{\alpha x} = \sum_{n=-\infty}^{\infty} c_n e^{inx}$ mit

$$\begin{aligned} c_n &= \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{\alpha x} e^{-inx} dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{(\alpha - in)x} dx = \frac{1}{2\pi} \left[\frac{1}{\alpha - in} e^{(\alpha - in)x} \right]_{-\pi}^{\pi} \\ &= \frac{1}{2\pi(\alpha - in)} (e^{(\alpha - in)\pi} - e^{-(\alpha - in)\pi}) \\ &= \frac{1}{2\pi(\alpha - in)} (-1)^n (e^{\alpha\pi} - e^{-\alpha\pi}) = \frac{(\alpha + in)(-1)^n}{\pi(\alpha^2 + n^2)} \sinh(\alpha\pi), \end{aligned}$$

also

$$e^{\alpha x} = \sum_{n=-\infty}^{\infty} \frac{(\alpha + in)(-1)^n}{\pi(\alpha^2 + n^2)} \sinh(\alpha\pi) e^{inx}$$

und

$$e^{\alpha x} = \frac{\sinh(\alpha\pi)}{\pi\alpha} + \sum_{n=1}^{\infty} \frac{2\alpha \sinh(\alpha\pi)(-1)^n}{\pi(\alpha^2 + n^2)} \cos(nx) + \frac{2n \sinh(\alpha\pi)(-1)^{n+1}}{\pi(\alpha^2 + n^2)} \sin(nx).$$

Die **Orthogonalitätsrelationen** lassen sich mit Hilfe der Grundfunktionen e^{inx} besonders elegant schreiben. Es gilt

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} e^{inx} \overline{e^{ikx}} dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{i(n-k)x} dx = \begin{cases} 1 & (n = k) \\ \frac{1}{2\pi} \frac{1}{i(n-k)} [e^{i(n-k)x}]_{-\pi}^{\pi} = 0 & (n \neq k) \end{cases},$$

da $e^{i\pi} = e^{-i\pi}$.

SATZ 125 (Besselsche Ungleichung). Für jede 2π -periodische, stückweise stetige Funktion f mit der Fourierreihe

$$f \sim \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx) = \sum_{n=-\infty}^{\infty} c_n e^{inx}$$

gilt für jedes $N \in \mathbb{N}$ die Besselsche Ungleichung

$$\frac{|a_0|^2}{2} + \sum_{n=1}^N (|a_n|^2 + |b_n|^2) = 2 \sum_{n=-\infty}^N |c_n|^2 \leq \frac{1}{\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx.$$

BEWEIS. Wir betrachten die endlichen Partialsummen der Fourierreihe $S_N(x) = \sum_{n=-N}^N c_n e^{inx}$. Es gilt

$$\begin{aligned} \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \overline{S_N(x)} dx &= \sum_{n=-N}^N \frac{\overline{c_n}}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx = \sum_{n=-N}^N \overline{c_n} c_n = \sum_{n=-N}^N |c_n|^2 \\ \frac{1}{2\pi} \int_{-\pi}^{\pi} S_N(x) \overline{S_N(x)} dx &= \sum_{n=-N}^N c_n \overline{c_n} = \sum_{n=-N}^N |c_n|^2 \end{aligned}$$

und damit

$$\begin{aligned} 0 &\leq \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x) - S_N(x)|^2 dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} (f(x) - S_N(x)) \overline{(f(x) - S_N(x))} dx \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx - S_N(x) \overline{f(x)} dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx - \sum_{n=-N}^N |c_n|^2 \end{aligned}$$

für alle $N \in \mathbb{N}$. $\square$

1.10. Rechenregeln - zweiter Teil. Mit Hilfe der komplexen Darstellung der Fourierreihe läßt sich die Berechnung der Fourierreihe einer Funktion vereinfachen. Es sei $f \sim c_0 + \sum_{n=1}^{\infty} c_n e^{inx} + c_{-n} e^{-inx}$ die Fourierreihe einer 2π -periodischen, integrierbaren Funktion f .

1.10.1. *Konjugation.* Die Funktion $g : \mathbb{R} \rightarrow \mathbb{C}$, $g(t) = \overline{f(t)}$, ist ebenfalls 2π -periodisch und integrierbar und hat die Fourierreihe

$$\overline{f(t)} \sim \overline{c_0} + \sum_{n=1}^{\infty} \overline{c_{-n}} e^{int} + \overline{c_n} e^{-int} = \sum_{n=-\infty}^{\infty} \overline{c_{-n}} e^{int},$$

denn $\overline{f(t)e^{int}} = \overline{f(t)}e^{-int}$.

1.10.2. *Zeitumkehr.* Die Funktion $g : \mathbb{R} \rightarrow \mathbb{C}$, $g(t) = f(-t)$, ist ebenfalls 2π -periodisch und integrierbar und hat die Fourierreihe

$$f(-t) \sim c_0 + \sum_{n=1}^{\infty} c_{-n} e^{int} + c_n e^{-int} = \sum_{n=-\infty}^{\infty} c_{-n} e^{int},$$

denn

$$\int_{-\pi}^{\pi} f(t) e^{-int} dt = - \int_{\pi}^{-\pi} f(-x) e^{inx} dx = \int_{-\pi}^{\pi} f(-x) e^{inx} dx.$$

1.10.3. *Verschiebung im Zeitbereich.* Die Funktion $g : \mathbb{R} \rightarrow \mathbb{C}$, $g(t) = f(t+a)$ mit $a \in \mathbb{R}$, ist ebenfalls 2π -periodisch und integrierbar und hat die Fourierreihe

$$f(t+a) \sim \sum_{n=-\infty}^{\infty} (e^{ina} c_n) e^{inx},$$

denn

$$\int_{-\pi}^{\pi} f(t+a) e^{-int} dt = \int_{-\pi+a}^{\pi+a} f(x) e^{-in(x-a)} dx = e^{ina} \int_{-\pi}^{\pi} f(x) e^{-inx} dx$$

und der Integrand ist 2π -periodisch.

BEISPIEL 254 (verschobener Rechteckimpuls). Für eine reelle Zahl $a \in (0, \pi)$ wollen wir die Fourierreihe der Funktion $g : \mathbb{R} \rightarrow \mathbb{R}$

$$g(t+2\pi k) = \begin{cases} 1 & \text{für } t \in [-a, \pi-a), k \in \mathbb{Z} \\ 0 & \text{für } t \in [-\pi, -a) \cup [\pi-a, \pi), k \in \mathbb{Z} \end{cases},$$

also die 2π -periodische Fortsetzung der charakteristischen Funktion $\chi_{-[a, \pi-a)}$ auf dem Intervall $[-\pi, \pi)$ berechnen. Es gilt $g(t) = f(t+a)$, wobei f der

Rechteckimpuls aus Beispiel 239 und Beispiel 243 ist. Also $f(t) = 1$ für $x \in [0, \pi)$ und $f(t) = 0$ für $x \in [-\pi, 0)$ mit der Fourierreihe

$$f(t) \sim \frac{1}{2} + \sum_{m=0}^{\infty} \frac{2}{\pi} \frac{\sin((2m+1)x)}{2m+1} = \frac{1}{2} + \frac{i}{\pi} \sum_{m=0}^{\infty} \frac{e^{-i(2m+1)x} - e^{i(2m+1)x}}{(2m+1)}.$$

Mit der Formel für die Verschiebung im Zeitbereich erhalten wir die komplexe Form

$$g(t) = f(t+a) \sim \frac{1}{2} + \frac{i}{\pi} \sum_{m=0}^{\infty} \frac{e^{-i(2m+1)a} e^{-i(2m+1)x} - e^{i(2m+1)a} e^{i(2m+1)x}}{(2m+1)}$$

oder mit Hilfe der Additionstheoreme die reelle Form

$$\begin{aligned} g(t) = f(t+a) &\sim \frac{1}{2} + \sum_{m=0}^{\infty} \frac{2}{\pi} \frac{\sin((2m+1)(x+a))}{2m+1} \\ &= \frac{1}{2} + \frac{2}{\pi} \sum_{m=0}^{\infty} \frac{\sin((2m+1)a) \cos((2m+1)x) + \cos((2m+1)a) \sin((2m+1)a)}{2m+1}. \end{aligned}$$

1.10.4. *Verschiebung im Frequenzbereich.* Die Funktion $g : \mathbb{R} \rightarrow \mathbb{C}$, $g(t) = e^{ikt} f(t)$ mit $k \in \mathbb{Z}$, ist ebenfalls 2π -periodisch und integrierbar und hat die Fourierreihe

$$e^{ikt} f(t) \sim \sum_{n=-\infty}^{\infty} c_{n-k} e^{inx},$$

denn

$$\int_{-\pi}^{\pi} e^{ikt} f(t) e^{-int} dt = \int_{-\pi}^{\pi} f(t) e^{-i(n-k)t} dt.$$

1.11. Das periodische Faltungsprodukt.

DEFINITION 8. Das **Faltungsprodukt** $f * g$ zweier stückweise stetiger, 2π -periodischer Funktionen f und g ist die 2π -periodische Funktion

$$(f * g)(t) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t-x) g(x) dx.$$

LEMMA 33. Das periodische Faltungsprodukt ist assoziativ, distributiv und kommutativ, d.h. für stückweise stetige, 2π -periodische Funktionen f, g, h gilt

$$f * (g * h) = (f * g) * h, \quad f * (g + h) = (f * g) + (f * h), \quad f * g = g * f.$$

BEWEIS. Es gilt

$$\begin{aligned} (f * g)(t) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t-x) g(x) dx = -\frac{1}{2\pi} \int_{t+\pi}^{t-\pi} f(y) g(t-y) dy \\ &= -\frac{1}{2\pi} \int_{\pi}^{-\pi} f(y) g(t-y) dy = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(y) g(t-y) dy = (g * f)(t) \end{aligned}$$

mit der Koordinatentransformation $y = t - x$, da f und g 2π -periodische Funktionen sind.

Das Distributivgesetz $f * (g + h) = (f * g) + (f * h)$ folgt direkt aus der Additivität des Integrals $\int f(t - x)(g(x) + h(x)) dx = \int f(t - x)g(x) dx + \int f(t - x)h(x) dx$.

Zum Beweis des Assoziativgesetzes benötigen wir einen Satz über die Vertauschbarkeit der Integrationsreihenfolge aus der Integralrechnung mehrerer Veränderlicher. Damit gilt

$$\begin{aligned}
 f * (g * h)(t) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t - x)(g * h)(x) dx \\
 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t - x) \frac{1}{2\pi} \int_{-\pi}^{\pi} g(x - y)h(y) dy dx \\
 &= \frac{1}{4\pi^2} \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} f(t - x)g(x - y)h(y) dy dx \\
 &= \frac{1}{4\pi^2} \int_{-\pi}^{\pi} \int_{-\pi}^{\pi} f(t - x)g(x - y)h(y) dx dy \\
 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} h(y) \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t - y - (x - y))g(x - y) dx dy \\
 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} h(y) \frac{1}{2\pi} \int_{-\pi-y}^{\pi-y} f(t - y - \tilde{x})g(\tilde{x}) d\tilde{x} dy \\
 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} h(y) \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t - y - \tilde{x})g(\tilde{x}) d\tilde{x} dy \\
 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} h(y) \frac{1}{2\pi} (f * g)(t - y) dy = ((f * g) * h)(t)
 \end{aligned}$$

□

SATZ 126. *Das periodische Faltungsprodukt glättet, d.h. wenn f eine (stückweise) stetige und g eine k -mal (stückweise) stetig differenzierbare Funktion ist, dann ist auch $f * g$ eine k -mal stetig differenzierbare Funktion.*

BEWEIS. Wenn f (stückweise) stetig ist, so ist f auf dem Intervall $[-\pi, \pi]$ beschränkt durch eine Konstante S . Auch jede der Ableitungen $g^{(j)}$ ist auf dem Intervall $[-\pi, \pi]$ durch eine Konstante $M^{(j)}$ beschränkt. Also gilt

$$\begin{aligned}
 |(f * g)(t + h) - (f * g)(t)| &= \frac{1}{2\pi} \left| \int_{-\pi}^{\pi} (g(t + h - x) - g(t - x))f(x) dx \right| \\
 &\leq \frac{S}{2\pi} \int_{-\pi}^{\pi} |g(t + h - x) - g(t - x)| dx \\
 &= h \frac{S}{2\pi} \int_{-\pi}^{\pi} |g'(t - x + \xi_1 h)| dx
 \end{aligned}$$

usw. nach dem Mittelwertsatz für g . Das Integral über ein kleines Intervall um eine Sprungstelle y_0 von $g^{(j)}$ schätzt man durch

$$\int_{y_0-h}^{y_0+h} |g^{(j)}(y + \xi_1 \dots \xi_j h) - g^{(j)}(y)| dy \leq hM^{(j)}$$

ab. Auf den restlichen Intervallen ist die zu integrierende Funktion stetig und man hat gleichmäßige Konvergenz, also

$$\lim_{h \rightarrow 0} \int_{-\pi}^{\pi} |g(t+h-x) - g(t-x)| dx = 0, \quad \lim_{h \rightarrow 0} \int_{-\pi}^{\pi} |g'(t+\xi h-x) - g'(t-x)| dx = 0 \dots$$

□

SATZ 127 (Fourierreihe des Faltungsprodukts). *Wenn f und g stückweise stetige, 2π -periodische Funktionen mit den Fourierreihen*

$$f \sim \sum_{k=-\infty}^{\infty} c_k e^{ikx} \quad \text{und} \quad g \sim \sum_{k=-\infty}^{\infty} d_k e^{ikx}$$

*sind, dann hat das Faltungsprodukt $f * g$ die Fourierreihe*

$$f * g \sim \sum_{k=-\infty}^{\infty} c_k d_k e^{ikx}.$$

BEWEIS. Es gilt

$$\begin{aligned} \frac{1}{2\pi} \int_{-\pi}^{\pi} (f * g)(x) e^{-ikx} dx &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1}{2\pi} \left(\int_{-\pi}^{\pi} f(t-x) g(x) dx \right) e^{-ikt} dt \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} g(x) e^{-inx} \left(\frac{1}{2\pi} \int_{-\pi}^{\pi} f(t-x) e^{-int} e^{ikx} dt \right) dx \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} g(x) e^{-ikx} \left(\frac{1}{2\pi} \int_{-\pi}^{\pi} f(t-x) e^{-ik(t-x)} dt \right) dx \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} g(x) e^{-ikx} c_k dx = \frac{c_k}{2\pi} \int_{-\pi}^{\pi} g(x) e^{-ikx} dx = c_k d_k. \end{aligned}$$

□

1.12. Vollständigkeit und Eindeutigkeit.

SATZ 128 (Vollständigkeitssatz). *Es sei $f: \mathbb{R} \rightarrow \mathbb{C}$ eine 2π -periodische, stückweise stetige Funktion. Wenn f in den Sprungstellen die Mittelwertbedingung erfüllt,*

$$f(x) = \frac{1}{2}(f(x^+) + f(x^-)) \quad \forall x \in \mathbb{R},$$

und alle Fourierkoeffizienten von f verschwinden,

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx = 0 \quad \forall n \in \mathbb{Z},$$

dann gilt $f(x) = 0$ für alle $t \in \mathbb{R}$.

BEWEIS. Natürlich gilt der Satz für stückweise stetig differenzierbare Funktionen f , weil dann die Fourierreihe von f punktweise gegen f konvergiert.

Angenommen $f(x_0) > 0$ in einer Stetigkeitsstelle x_0 . Dann falten wir f mit einer „cut-off“-Funktion φ . D.h. φ ist glatt, 2π -periodisch, $\varphi(x_0) = 1$ und verschwindet außerhalb einer kleinen Umgebung U_0 von x_0 , auf der f positiv ist. Nach der Formel für die Fourierkoeffizienten des Faltungsproduktes (Satz 127) verschwinden alle Fourierkoeffizienten von $f * \varphi$. Dies ist ein Widerspruch, da die Funktion $f * \varphi$ stetig differenzierbar ist, also durch ihre Fourierreihe dargestellt wird, und

$$(f * \varphi)(x_0) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x_0 - x) \varphi(x) dx = \frac{1}{2\pi} \int_{U_0} f(x_0 - x) \varphi(x) dx > 0.$$

□

FOLGERUNG 41 (Eindeutigkeitssatz). *Haben zwei stückweise stetige Funktionen dieselben Fourierkoeffizienten und erfüllen die Mittelwertbedingung an allen Sprungstellen, so sind sie identisch.*

BEWEIS. Die Funktion $f - g$ erfüllt die Voraussetzungen des Vollständigkeitsatzes. □

FOLGERUNG 42 (Konvergenz der Fourierreihe des Faltungsproduktes). *Es seien f und g zwei stückweise stetige, 2π -periodische Funktionen mit den Fourierreihen $f(x) \sim \sum_{n=-\infty}^{\infty} c_n e^{inx}$ bzw. $g(x) \sim \sum_{n=-\infty}^{\infty} d_n e^{inx}$. Dann konvergiert die Fourierreihe des Faltungsproduktes $f * g$ stets gleichmäßig und es gilt*

$$(f * g)(t) = \sum_{n=-\infty}^{\infty} c_n d_n e^{int}.$$

BEWEIS. Aus $0 \leq (|c_n| - |d_n|)^2$ folgt $2|c_n d_n| \leq |c_n|^2 + |d_n|^2$. Aus der Besselschen Ungleichung, $\sum_{n=-\infty}^{\infty} |c_n|^2 + |d_n|^2 \leq \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 + |g(x)|^2 dx$, folgt, dass die Reihe $\sum_{n=-\infty}^{\infty} c_n d_n e^{int}$ gleichmäßig konvergiert, da $|c_n d_n e^{int}| = |c_n d_n|$. Mit dem Eindeutigkeitssatz folgt dann, dass das Faltungsprodukt $f * g$ durch seine Fourierreihe dargestellt wird. □

FOLGERUNG 43 (Vollständigkeitsrelation, Parsevalsche Gleichung). *Für zwei stückweise stetige, 2π -periodische Funktionen $f, g : \mathbb{R} \rightarrow \mathbb{C}$ mit den Fourierreihen $f(x) \sim \sum_{n=-\infty}^{\infty} c_n e^{inx}$ und $g(x) \sim \sum_{n=-\infty}^{\infty} d_n e^{inx}$ gilt:*

$$(169) \quad \sum_{n=-\infty}^{\infty} c_n \overline{d_n} = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx$$

$$(170) \quad \sum_{n=-\infty}^{\infty} |c_n|^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx$$

BEWEIS. Wir betrachten die Faltung der Funktionen f und h mit $h(x) = \overline{g(-x)}$. Die Funktion h entsteht durch Zeitumkehr und Konjugation aus g .

Also hat h die Fourierreihe $h(x) \sim \sum_{n=-\infty}^{\infty} \overline{d_{-(-n)}} e^{inx} = \sum_{n=-\infty}^{\infty} \overline{d_n} e^{inx}$. Darum gilt

$$(f * h)(t) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) h(t-x) dx = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \overline{g(x-t)} dx \sim \sum_{n=-\infty}^{\infty} c_n \overline{d_n} e^{int}.$$

Die erste Gleichung folgt mit $t = 0$ und die zweite Gleichung mit $f = g$. $\square$

Die Gleichung

$$\sum_{n=-\infty}^{\infty} |c_n|^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx$$

heißt **Parsevalsche Gleichung**. Ihre reelle Form ist für reellwertige, stückweise stetige, 2π -periodische Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ mit der Fourierreihe $a_0/2 + \sum_{n=1}^{\infty} a_n \cos(nx) + b_n \sin(nx)$ ist

$$f(x) \sim \frac{|a_0|^2}{2} + \sum_{n=1}^{\infty} |a_n|^2 + |b_n|^2 = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x)^2 dx.$$

BEISPIEL 255 (Reihenentwicklung für π^2 aus der Fourierreihe der Sägezahnkurve). Auf dem Intervall $(-\pi, \pi)$ gilt

$$x = 2 \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} \sin(nx).$$

Aus der Parsevalschen Gleichung folgt

$$4 \sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{1}{\pi} \int_{-\pi}^{\pi} x^2 dx = \frac{1}{3\pi} [x^3]_{-\pi}^{\pi} = \frac{2}{3} \pi^2, \quad \text{also} \quad \sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{1}{6} \pi^2.$$

1.13. Fourierreihenansatz bei Differentialgleichungen. Wir betrachten lineare Differentialgleichungen mit konstanten Koeffizienten und periodischem inhomogenem Anteil der Form $y'' + \omega_0^2 y = f(x)$ mit $\omega_0 > 0$, wobei f eine L -periodische, stückweise stetige und stückweise glatte Funktion ist.

Da das zur homogenen Differentialgleichung $y'' + \omega_0^2 y = 0$ gehörende charakteristische Polynom $\lambda^2 + \omega^2$ die komplexen Nullstellen $\pm i\omega_0$ besitzt, hat die allgemeine Lösung der homogenen Differentialgleichung $y'' + \omega_0^2 y = 0$ die Form

$$y_h(x) = c_1 \cos(\omega_0 x) + c_2 \sin(\omega_0 x)$$

mit Konstanten $c_1, c_2 \in \mathbb{R}$.

Wir entwickeln die Funktion f in eine Fourierreihe

$$f(x) = \frac{A_0}{2} + \sum_{n=1}^{\infty} A_n \cos(n\omega x) + B_n \sin(n\omega x), \quad \omega = \frac{2\pi}{L},$$

und machen für die spezielle Lösung der inhomogenen Differentialgleichung $y'' + \omega_0 y = f(x)$ einen Fourierreihenansatz mit der Periode L :

$$y(x) = \frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\omega x) + b_n \sin(n\omega x)$$

Setzen wir diesen Ansatz in die Differentialgleichung ein und nehmen an, dass die Reihe zweimal stetig differenzierbar ist. So erhalten wir

$$\begin{aligned} 0 &= -\frac{A_0}{2} - \sum_{n=1}^{\infty} A_n \cos(n\omega x) + B_n \sin(n\omega x) \\ &\quad + \left(\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\omega x) + b_n \sin(n\omega x) \right)'' \\ &\quad + \omega_0^2 \left(\frac{a_0}{2} + \sum_{n=1}^{\infty} a_n \cos(n\omega x) + b_n \sin(n\omega x) \right) \\ &= \frac{\omega_0^2 a_0 - A_0}{2} + \sum_{n=1}^{\infty} ((\omega_0^2 - n^2 \omega^2) a_n - A_n) \cos(n\omega x) \\ &\quad + \sum_{n=1}^{\infty} ((\omega_0^2 - n^2 \omega^2) b_n - B_n) \sin(n\omega x). \end{aligned}$$

Falls $\omega_0 \neq n\omega$ für alle $n \in \mathbb{Z}$ (keine Resonanz), so können wir aus dieser Reihenentwicklung die Koeffizienten a_n und b_n bestimmen. Es gilt

$$a_0 = \frac{A_0}{\omega_0^2}, \quad a_n = \frac{A_n}{\omega_0^2 - n^2 \omega^2}, \quad b_n = \frac{B_n}{\omega_0^2 - n^2 \omega^2}.$$

Falls $\omega_0 = k\omega$ für ein $k \in \mathbb{Z}$ (Resonanz), so ändern wir den Ansatz in

$$y(x) = \frac{a_0}{2} + \sum_{n=1, n \neq k}^{\infty} a_n \cos(n\omega x) + b_n \sin(n\omega x) + a_k x \cos(k\omega x) + b_k x \sin(k\omega x)$$

und erhalten für die speziellen Koeffizienten a_k und b_k die Gleichungen

$$A_k \cos(k\omega x) + B_k \sin(k\omega x) = -2a_k k\omega \sin(k\omega x) + 2b_k k\omega \cos(k\omega x),$$

also $a_k = -B_k/(2k\omega)$ und $b_k = A_k/(2k\omega)$.

2. Orthonormalsysteme

Wir werden sehen, dass die Fourierkoeffizienten „nur“ die Koordinaten der Funktion bezüglich eines rechtwinkligen Koordinatensystems sind.

2.1. Vektorräume.

DEFINITION 9. *Ein reeller (komplexer) Vektorraum V ist eine Menge mit zwei Operationen, der Addition von Vektoren*

$$(v, w) \mapsto v + w \text{ für } v, w \in V$$

und der Multiplikation von Vektoren mit Skalaren

$$(\lambda, v) \mapsto \lambda v \text{ für } v \in V, \lambda \in \mathbb{R} (\lambda \in \mathbb{C}),$$

so dass $v + w = w + v$, $\lambda(w + v) = \lambda w + \lambda v$ und $\lambda(\mu v) = (\lambda\mu)v$ für alle $v, w \in V$ und $\lambda, \mu \in \mathbb{R}$ (bzw. $\mathbb{C}$).

BEISPIEL 256. Der $\mathbb{R}^n$, der Raum aller n -Tupel reeller Zahlen, ist ein reeller Vektorraum. Addition und Skalarmultiplikation sind komponentenweise definiert, also

$$(x_1, \dots, x_n) + (y_1, \dots, y_n) = (x_1 + y_1, \dots, x_n + y_n), \quad \lambda(x_1, \dots, x_n) = (\lambda x_1, \dots, \lambda x_n)$$

für alle $x_j, y_j, \lambda \in \mathbb{R}$.

BEISPIEL 257. Der $\mathbb{C}^n$, der Raum aller n -Tupel komplexer Zahlen, ist ein komplexer Vektorraum. Addition und Skalarmultiplikation sind komponentenweise definiert, also

$$(x_1, \dots, x_n) + (y_1, \dots, y_n) = (x_1 + y_1, \dots, x_n + y_n), \quad \lambda(x_1, \dots, x_n) = (\lambda x_1, \dots, \lambda x_n)$$

für alle $x_j, y_j, \lambda \in \mathbb{C}$.

BEISPIEL 258 (Vektorraum stetiger Funktionen). Auch die auf einem Intervall stetigen Funktionen bilden einen Vektorraum. Die Summe zweier Funktionen f und g ist definiert durch $(f + g)(x) = f(x) + g(x)$. Die Multiplikation einer Funktion f mit einer reellen oder komplexen Zahl λ ist die Funktion $(\lambda f)(x) = \lambda f(x)$. Aus der Stetigkeit von f und g folgt die Stetigkeit von $f + g$ und λf .

Wir benutzen die Bezeichnungen $\mathcal{C}([a, b], \mathbb{R})$ und $\mathcal{C}([a, b], \mathbb{C})$ für die Vektorräume reellwertiger bzw. komplexwertiger stetiger Funktion auf einem Intervall $[a, b]$.

2.2. Normierte Vektorräume. Wir möchten die Länge von Vektoren bzw. den Abstand zweier Vektoren v und w , also die Länge der Differenz $v - w$, messen können. Dazu definieren wir eine Norm auf dem Vektorraum.

DEFINITION 10. Eine **Norm** auf einem reellen Vektorraum V ist eine Abbildung $V \rightarrow \mathbb{R}$, $v \mapsto \|v\|$, mit folgenden Eigenschaften:

- *Positivität:* $\|v\| \geq 0$ für alle $v \in V$ und $\|v\| = 0 \Leftrightarrow v = 0$
- *Ähnlichkeit:* $\|\lambda v\| = |\lambda| \|v\|$ für alle $v \in V$ und alle $\lambda \in \mathbb{R}$
- *Dreiecksungleichung:* $\|v + w\| \leq \|v\| + \|w\|$ für alle $v, w \in V$

BEISPIEL 259 (Euklidische Norm). Auf dem reellen Vektorraum $\mathbb{R}^n$ ist durch

$$\|(x_1, \dots, x_n)\| = \sqrt{x_1^2 + \dots + x_n^2}$$

eine Norm definiert. Auf dem komplexen Vektorraum $\mathbb{C}^n$ ist durch

$$\|(z_1, \dots, z_n)\| = \sqrt{|z_1|^2 + \dots + |z_n|^2} = \sqrt{z_1 \bar{z}_1 + \dots + z_n \bar{z}_n}$$

eine Norm definiert. Wir überprüfen die drei Eigenschaften Positivität, Ähnlichkeit und die Dreiecksungleichung. Sie folgen aus den entsprechenden Eigenschaften für die Betragsfunktion der reellen Zahlen.

- **Positivität:** Es sei $x = (x_1, \dots, x_n)$ ein Vektor im $\mathbb{R}^n$. Dann gilt $\|x\| \geq 0$, da $x_j^2 \geq 0$ für alle j . Außerdem gilt $\|x\| = 0$ genau dann, wenn $x_j = 0$ für alle j , also $x = 0$.
- **Ähnlichkeit:** Für $\lambda \in \mathbb{R}$ gilt

$$\begin{aligned} \|\lambda x\| &= \sqrt{|\lambda x_1|^2 + \dots + |\lambda x_n|^2} = \sqrt{\lambda^2 |x_1|^2 + \dots + \lambda^2 |x_n|^2} \\ &= \sqrt{|\lambda|^2 (|x_1|^2 + \dots + |x_n|^2)} = |\lambda| \sqrt{|x_1|^2 + \dots + |x_n|^2} = |\lambda| \|x\| \end{aligned}$$

- **Dreiecksungleichung:** Für zwei Vektoren $x, y \in \mathbb{R}^n$ mit $x = (x_1, \dots, x_n)$ und $y = (y_1, \dots, y_n)$ gilt:

$$\begin{aligned} 0 &\leq \sum_{j \neq k} (x_j y_k - x_k y_j)^2 \\ 2 \sum_{j \neq k} x_j y_j x_k y_k &\leq \sum_{j \neq k} x_j^2 y_k^2 + x_k^2 y_j^2 \\ \sum_{j,k=1}^n x_j y_j x_k y_k &\leq \sum_{j,k=1}^n x_j^2 y_k^2 \\ \sum_{j,k=1}^n x_j y_j x_k y_k &\leq (x_1^2 + \dots + x_n^2)(y_1^2 + \dots + y_n^2) \\ 2 \sum_{j=1}^n x_j y_j &\leq 2 \sqrt{x_1^2 + \dots + x_n^2} \sqrt{y_1^2 + \dots + y_n^2} \\ (x_1 + y_1)^2 + \dots + (x_n + y_n)^2 &\leq \sum_{j=1}^n x_j^2 + \sum_{j=1}^n y_j^2 + 2 \sqrt{x_1^2 + \dots + x_n^2} \sqrt{y_1^2 + \dots + y_n^2} \\ \|x + y\|^2 &\leq \|x\|^2 + \|y\|^2 + 2\|x\|\|y\| \\ \|x + y\| &\leq \|x\| + \|y\| \end{aligned}$$

BEISPIEL 260 (Maximumnorm). Auf den Vektorräumen $\mathcal{C}([a, b], \mathbb{R})$ und $\mathcal{C}([a, b], \mathbb{C})$ ist durch $\|f\|_\infty := \max\{|f(x)| : x \in [a, b]\}$ eine Norm definiert.

Auch auf den endlich dimensionalen Vektorräumen $\mathbb{R}^n$ und $\mathbb{C}^n$ kann man durch $\|(x_1, \dots, x_n)\| := \max\{x_j : j = 1, \dots, n\}$ eine Norm definieren, die der Maximumnorm auf dem unendlich dimensionalen Vektorraum $\mathcal{C}([a, b])$ entspricht.

Wir überprüfen Positivität, Ähnlichkeit und die Dreiecksungleichung für die Definition der Maximumnorm auf $\mathcal{C}([a, b], \mathbb{R})$. Für alle $f \in \mathcal{C}([a, b], \mathbb{R})$ gilt $\|f\|_\infty \geq 0$, denn $|f(x)| \geq 0$ für alle $x \in [a, b]$. Außerdem ist $\|f\|_\infty = 0$ genau dann, wenn $|f(x)| = 0$ für alle $x \in [a, b]$, also $f \equiv 0$. Die Ähnlichkeit der Maximumnorm folgt direkt aus $|\lambda f(x)| = |\lambda| |f(x)|$ für alle $\lambda \in \mathbb{R}$ und $x \in [a, b]$. Die Dreiecksungleichung gilt für die Maximumnorm, weil $\max\{|f(x) + g(x)| : x \in [a, b]\} \leq \max\{|f(x)| : x \in [a, b]\} + \max\{|g(x)| : x \in [a, b]\}$ für alle $f, g \in \mathcal{C}([a, b], \mathbb{R})$.

2.3. Umgebungen und Konvergenz. Wenn man den Begriff der Norm hat, dann kann man Umgebungen (eine Topologie) definieren und Konvergenz diskutieren. Einen Vektorraum, der mit einer Norm $|||$ ausgestattet ist, nennt man einen normierten Vektorraum.

DEFINITION 11. *Es sei $(V, |||)$ ein normierter Vektorraum, $\varepsilon > 0$ und $x \in V$. Die ε -Umgebung $U_\varepsilon(x)$ von x ist die Menge aller Punkte in V , deren Abstand von x kleiner als ε ist, also*

$$U_\varepsilon(x) = \{y \in V : ||x - y|| < \varepsilon\}.$$

DEFINITION 12. *Es sei $(V, |||)$ ein normierter Vektorraum. Eine Folge $\{f_n\}_{n \in \mathbb{N}} \subset V$ **konvergiert** gegen $f \in V$, wenn zu jedem $\varepsilon > 0$ fast alle Folgenglieder in $U_\varepsilon(f)$ liegen.*

Die ε -Umgebung bezüglich der euklidischen Norm in $\mathbb{R}^n$ ist wirklich ein offener Ball vom Radius ε . Die ε -Umgebung bezüglich der Maximumnorm in $\mathbb{R}^n$ ist ein offener Würfel mit Seitenlänge 2ε . Auf $\mathbb{R}^n$ und $\mathbb{C}^n$ führen alle Normen zu äquivalenten Konvergenzbegriffen, auch wenn die ε -Umgebungen verschieden aussehen. Zur Maximumnorm auf $\mathcal{C}([a, b], \mathbb{R})$ gehört die gleichmäßige Konvergenz von Funktionen.

LEMMA 34 (Komponentenweise Konvergenz in $\mathbb{R}^n$). *Eine Folge $\{x_k\}_{k \in \mathbb{N}} \subset \mathbb{R}^n$ konvergiert genau dann gegen $y \in \mathbb{R}^n$ (bezüglich der euklidischen Norm), wenn $x_k = (x_{k,1}, \dots, x_{k,n})$, $y = (y_1, \dots, y_n)$ und für jedes $j = 1, \dots, n$ die Folge reeller Zahlen $\{x_{k,j}\}_{k \in \mathbb{N}}$ gegen y_j konvergiert.*

BEWEIS. Es gilt $\lim_{k \rightarrow \infty} ||x_k - y|| = 0$ genau dann, wenn $\lim_{k \rightarrow \infty} (x_{k,j} - y_j)^2 = 0$ für alle j . $\square$

2.4. Skalarprodukt. In vielen Fällen wird die Norm durch ein Skalarprodukt gegeben, das uns dann sogar erlaubt, Winkel und Längen zu messen.

DEFINITION 13. *Ein **Skalarprodukt** auf einem reellen Vektorraum V ist eine Abbildung $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{R}$ mit folgenden Eigenschaften:*

- *Symmetrie:* $\langle v, w \rangle = \langle w, v \rangle$ für alle $v, w \in V$
- *Linearität:* $\langle \lambda_1 v_1 + \lambda_2 v_2, w \rangle = \lambda_1 \langle v_1, w \rangle + \lambda_2 \langle v_2, w \rangle$ für alle $v_1, v_2, w \in V$ und $\lambda_1, \lambda_2 \in \mathbb{R}$
- *Positivität:* $\langle v, v \rangle \geq 0$ für alle $v \in V$ und $\langle v, v \rangle = 0 \Leftrightarrow v = 0$.

BEISPIEL 261 (Euklidisches Skalarprodukt). Auf dem reellen Vektorraum $\mathbb{R}^n$ kennen wir das euklidische Skalarprodukt $\langle x, y \rangle = \sum_{j=1}^n x_j y_j$ für $x = (x_1, \dots, x_n)$ und $y = (y_1, \dots, y_n)$.

BEISPIEL 262 (Konvergenz im quadratischen Mittel für reellwertige Funktionen). Auf dem Funktionenraum $\mathcal{C}([a, b], \mathbb{R})$ wird durch

$$\langle f, g \rangle = \int_a^b f(x)g(x) dx$$

ein Skalarprodukt definiert. Wir überprüfen

- Symmetrie: $\langle f, g \rangle = \int_a^b f(x)g(x) dx = \int_a^b g(x)f(x) dx = \langle g, f \rangle$,
- Linearität:

$$\begin{aligned}\langle \lambda_1 f_1 + \lambda_2 f_2, g \rangle &= \int_a^b (\lambda_1 f_1(x) + \lambda_2 f_2(x))g(x) dx \\ &= \lambda_1 \int_a^b f_1(x)g(x) dx + \lambda_2 \int_a^b f_2(x)g(x) dx \\ &= \lambda_1 \langle f_1, g \rangle + \lambda_2 \langle f_2, g \rangle\end{aligned}$$

und

- Positivität: $\langle f, f \rangle = \int_a^b f(x)^2 dx \geq 0$, da $f(x) \geq 0$ für alle $x \in [a, b]$, und $\int_a^b f(x)^2 dx = 0$ genau dann, wenn $f(x) = 0$ für alle $x \in [a, b]$. Hier ist es wichtig, dass f stetig ist. Wenn $f(x)^2 > 0$ für ein x , so ist $f(x)^2 > 0$ in einer Umgebung U von x und damit $\int_a^b f(x)^2 dx \geq \int_U f(x)^2 dx > 0$.

DEFINITION 14. Ein **hermitesches Skalarprodukt** auf einem komplexen Vektorraum V ist eine Abbildung $\langle, \rangle : V \times V \rightarrow \mathbb{C}$ mit folgenden Eigenschaften:

- Symmetrie: $\langle v, w \rangle = \overline{\langle w, v \rangle}$ für alle $v, w \in V$
- Linearität: $\langle \lambda_1 v_1 + \lambda_2 v_2, w \rangle = \lambda_1 \langle v_1, w \rangle + \lambda_2 \langle v_2, w \rangle$ für alle $v_1, v_2, w \in V$ und $\lambda_1, \lambda_2 \in \mathbb{C}$
- Positivität: $\langle v, v \rangle \geq 0$ für alle $v \in V$ und $\langle v, v \rangle = 0 \Leftrightarrow v = 0$.

Aus der Symmetrie und der Linearität folgt für ein hermitesches Skalarprodukt $\langle v, \lambda w \rangle = \bar{\lambda} \langle v, w \rangle$ für alle $v, w \in V$ und $\lambda \in \mathbb{C}$. Weiterhin sichern Symmetrie und Linearität, dass $\langle v, v \rangle \in \mathbb{R}$ für alle $v \in V$ und die Positivitätsbedingung sinnvoll ist.

BEISPIEL 263 (Hermitesches Skalarprodukt auf $\mathbb{C}^n$). Auf dem komplexen Vektorraum $\mathbb{C}^n$ kennen wir das hermitesche Skalarprodukt $\langle z, w \rangle = \sum_{j=1}^n z_j \bar{w}_j$ für $z = (z_1, \dots, z_n)$ und $w = (w_1, \dots, w_n)$.

BEISPIEL 264 (Konvergenz im quadratischen Mittel für komplexwertige Funktionen). Auf dem Funktionenraum $\mathcal{C}([a, b], \mathbb{C})$ wird durch

$$\langle f, g \rangle = \int_a^b f(x) \overline{g(x)} dx$$

ein hermitesches Skalarprodukt definiert.

SATZ 129. Es sei $\langle, \rangle$ ein (hermitesches) Skalarprodukt auf einem Vektorraum V . Dann definiert $\|v\| := \sqrt{\langle v, v \rangle}$ eine Norm auf V . Außerdem gilt die Schwarzsche Ungleichung

$$|\langle v, w \rangle| \leq \|v\| \|w\| \text{ für alle } v, w \in V.$$

BEWEIS. Wir überprüfen die Eigenschaften der so definierten Norm:

- Positivität der Norm folgt direkt aus der Positivität des Skalarproduktes.

- Ähnlichkeit der Norm folgt direkt aus der Linearität des Skalarproduktes: $\|\lambda v\| = \sqrt{\langle \lambda v, \lambda v \rangle} = \sqrt{\lambda^2 \langle v, v \rangle} = |\lambda| \sqrt{\langle v, v \rangle} = |\lambda| \|v\|$
- Wir beweisen die Schwarzsche Ungleichung: Falls $w = 0$, so gilt die Ungleichung, weil $\langle v, 0 \rangle = 0$ und $\|0\| = 0$. Falls $\|w\| \neq 0$, so betrachten wir den Hilfsvektor $h := w/\|w\|$. Es gilt

$$0 \leq \|v - \langle v, h \rangle h\|^2 = \langle v - \langle v, h \rangle h, v - \langle v, h \rangle h \rangle$$

$$0 \leq \langle v, v \rangle - 2\langle v, h \rangle^2 + \langle v, h \rangle^2 \langle h, h \rangle = \langle v, v \rangle - \langle v, h \rangle^2$$

$$\langle v, h \rangle^2 \leq \langle v, v \rangle = \|v\|^2,$$

da $\langle h, h \rangle = 1$. Nun setzen wir wieder $h = w/\|w\|$ und multiplizieren die Gleichung mit $\|w\|^2$. So erhalten wir

$$\langle v, w \rangle^2 \leq \langle v, v \rangle \|w\|^2.$$

Durch Ziehen der Wurzel ergibt sich die Schwarzsche Ungleichung.

- Die Dreiecksungleichung erhalten wir mit Hilfe der Schwarzschen Ungleichung:

$$\begin{aligned} \|v + w\| &= \sqrt{\langle v + w, v + w \rangle} = \sqrt{\langle v, v \rangle + \langle w, w \rangle + \langle v, w \rangle + \langle w, v \rangle} \\ &= \sqrt{\|v\|^2 + \|w\|^2 + 2\Re(\langle v, w \rangle)} \leq \sqrt{\|v\|^2 + \|w\|^2 + 2|\langle v, w \rangle|} \\ &\leq \sqrt{\|v\|^2 + \|w\|^2 + 2\|v\|\|w\|} = \sqrt{(\|v\| + \|w\|)^2} = \|v\| + \|w\| \end{aligned}$$

□

Das Skalarprodukt $\langle f, g \rangle = \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx$ induziert also die Norm $\|f\| = \sqrt{\int_{-\pi}^{\pi} |f(x)|^2 dx}$ und eine Funktionenfolge $\{f_n\}_{n \in \mathbb{N}}$ konvergiert bezüglich dieser Norm gegen eine Funktion f , wenn $\lim_{n \rightarrow \infty} \|f_n - f\| = 0$. Da $\|f_n - f\|^2 = \int_{-\pi}^{\pi} |f_n(x) - f(x)|^2 dx$ ist die Bezeichnung „Konvergenz im quadratischen Mittel“ gerechtfertigt.

BEMERKUNG 92. Die gleichmäßige Konvergenz ist stärker als die Konvergenz im quadratischen Mittel. Wenn eine Folge $\{f_n\}_{n \in \mathbb{N}}$ stetiger Funktionen auf einem Intervall $[a, b]$ gleichmäßig gegen eine Funktion f konvergiert, so konvergiert sie auch im quadratischen Mittel gegen f , denn aus $|f_n(x) - f(x)| < \varepsilon$ für alle $x \in [a, b]$ folgt $\int_a^b (f_n(x) - f(x))^2 dx < (b - a)\varepsilon^2$. Die Umkehrung gilt nicht, wie man leicht an der Folge $\{g_n\}_{n \in \mathbb{N}} \subset \mathcal{C}([0, 1], \mathbb{R})$ mit

$$g_n(x) = \begin{cases} 1 - nx & \text{für } x \in [0, 1/n] \\ 0 & \text{für } x \in [1/n, 1] \end{cases}$$

sehen kann, die zwar im quadratischen Mittel aber nicht punktweise gegen $g(x) \equiv 0$ konvergiert.

DEFINITION 15. Es sei V ein Vektorraum mit Skalarprodukt $\langle \cdot, \cdot \rangle$. Zwei Vektoren $v, w \in V$ stehen **senkrecht** aufeinander, sind **orthogonal** zueinander, wenn $\langle v, w \rangle = 0$. Ein Vektor $v \in V$ heißt **normiert**, wenn $\|v\| = 1$.

BEISPIEL 265 (Orthogonalitätsrelationen). Wir betrachten das Skalarprodukt aus Beispiel 262. Die Vektoren $\sin(nx)$ und $\cos(mx)$ mit $m, n \in \mathbb{N}$ stehen paarweise senkrecht aufeinander als Elemente in $\mathcal{C}([-\pi, \pi], \mathbb{R})$ mit dem Skalarprodukt aus Beispiel 262. Die Elemente $\sin(nx)/\sqrt{\pi}$ und $\cos(nx)/\sqrt{\pi}$ haben für alle $n \in \mathbb{N}$ die Länge 1.

Die Vektoren e^{inx} mit $n \in \mathbb{N}$ stehen ebenfalls paarweise senkrecht aufeinander als Elemente in $\mathcal{C}([-\pi, \pi], \mathbb{C})$ mit dem Skalarprodukt aus Beispiel 264. Die Elemente $e^{inx}/\sqrt{2\pi}$ haben für alle $n \in \mathbb{Z}$ die Länge 1.

2.5. Orthonormalsysteme.

DEFINITION 16. Es sei V ein Vektorraum mit einem Skalarprodukt $\langle \cdot, \cdot \rangle$. Eine Menge von Vektoren $\{v_j\} \subset V$ heißt **Orthonormalsystem**, wenn die Vektoren normiert sind und paarweise senkrecht aufeinander stehen, also $\langle v_j, v_k \rangle = \delta_{jk}$.

BEISPIEL 266 (Standardbasis in $\mathbb{R}^n$). Die Standardeinheitsvektoren $e_1 = (1, 0, \dots, 0)$, $e_2 = (0, 1, 0, \dots, 0)$, $\dots$, $e_n = (0, \dots, 0, 1)$ bilden ein Orthonormalsystem.

BEISPIEL 267 (Harmonische Schwingungen). Bezüglich des Skalarprodukts aus Beispiel 264 sind die Mengen

$$\left\{ \frac{1}{\sqrt{2\pi}} e^{inx} \right\}_{n \in \mathbb{Z}} \quad \text{und} \quad \left\{ \frac{1}{\sqrt{2\pi}}, \frac{1}{\sqrt{\pi}} \sin(nx), \frac{1}{\sqrt{\pi}} \cos(nx) : n \in \mathbb{N} \right\}$$

Orthonormalsysteme.

LEMMA 35. Ein Orthonormalsystem ist linear unabhängig.

BEWEIS. Wenn $\sum_{j=1}^N a_j v_j = 0$ und $\{v_j : j = 1, \dots, N\}$ ist ein Orthonormalsystem, dann $0 = \langle v_k, \sum_{j=1}^N a_j v_j \rangle = a_k$ für alle $k = 1, \dots, N$. $\square$

BEMERKUNG 93. Aus jeder Menge linear unabhängiger Vektoren kann man mit Hilfe des Schmidtschen Orthonormalisierungsverfahrens ein Orthonormalsystem konstruieren. $\square$

SATZ 130 (Besselsche Ungleichung). Es sei $\{v_j : j \in \mathbb{N}\} \subset V$ ein Orthonormalsystem. Dann gilt für jedes $v \in V$ die Besselsche Ungleichung

$$\sum_{j \in \mathbb{N}} |\langle v, v_j \rangle|^2 \leq \|v\|^2$$

und $\sum_{j=1}^N \langle v, v_j \rangle v_j$ ist die beste Approximation von v im Untervektorraum, der von $\{v_1, \dots, v_N\}$ aufgespannt wird.

BEWEIS. Es gilt

$$0 \leq \left\| v - \sum_{j=1}^N \langle v, v_j \rangle v_j \right\|^2 = \left\langle v - \sum_{j=1}^N \langle v, v_j \rangle v_j, v - \sum_{j=1}^N \langle v, v_j \rangle v_j \right\rangle = \langle v, v \rangle - \sum_{j=1}^N |\langle v, v_j \rangle|^2,$$

da $\langle v, v_j \rangle = \overline{\langle v_j, v \rangle}$.

Um zu zeigen, dass $\sum_{j=1}^N \langle v, v_j \rangle v_j$ die beste Approximation von v ist, müssen wir die Differentialrechnung mehrerer Veränderlicher benutzen. Wir wollen nämlich die Funktion $f(a_1, \dots, a_N) = \|v - \sum_{j=1}^N a_j v_j\|^2$ minimieren. Eine notwendige Bedingung für ein lokales Extremum der Funktion f ist das Verschwinden des Gradienten $(-2\langle v, v_j \rangle + 2a_j)_{j=1, \dots, N}$. Dies bedeutet $a_j = \langle v, v_j \rangle$. Man überprüft, dass für diese a_j ein Maximum der Funktion f angenommen wird, da die Hessematrix von f gerade 2Id , also positiv definit ist. $\square$

BEMERKUNG 94. Der Vektor $v - \sum_{j=1}^N \langle v, v_j \rangle v_j$ ist die Projektion von v in einen Untervektorraum, der zum Orthonormalsystem $\{v_j : j \in \mathbb{N}\} \subset V$ senkrecht steht, denn $\langle v - \sum_{j=1}^N \langle v, v_j \rangle v_j, v_k \rangle = \langle v, v_k \rangle - \sum_{j=1}^N \langle v, v_j \rangle \langle v_j, v_k \rangle = 0$ für alle k , denn $\langle v_j, v_k \rangle = \delta_{jk}$. $\square$

DEFINITION 17. Ein Orthonormalsystem $\{v_j\}_{j \in J}$ heißt **vollständig**, wenn für alle $v \in V$ aus $\langle v, v_j \rangle = 0 \ \forall j$ auch $v = 0$ folgt.

BEISPIEL 268 (Harmonische Schwingungen). Die Orthonormalsysteme

$$\left\{ \frac{1}{\sqrt{2\pi}} e^{inx} \right\}_{n \in \mathbb{Z}} \quad \text{und} \quad \left\{ \frac{1}{\sqrt{2\pi}}, \frac{1}{\sqrt{\pi}} \sin(nx), \frac{1}{\sqrt{\pi}} \cos(nx) : n \in \mathbb{N} \right\}$$

sind vollständig.

BEMERKUNG 95. In unendlich dimensionalen Vektorräumen kann es passieren, dass eine unendliche Summe von Vektoren, z.B. $\sum_{j=0}^{\infty} \langle v, v_j \rangle v_j$, kein Element des Vektorraumes ist. Zum Beispiel ist jede endliche Überlagerung von Sinusschwingungen der Form

$$f_N(x) := \pi - 2 \sum_{n=1}^N \frac{\sin(nx)}{n}$$

eine stetige, 2π -periodische Funktion und damit ein Vektor in $\mathcal{C}([-\pi, \pi])$. Für jedes $x \in \mathbb{R}$ existiert der Grenzwert $\lim_{N \rightarrow \infty} f_N(x)$ und die Funktionenfolge $\{f_N\}$ konvergiert punktweise gegen die Sägezahnkurve f mit $f(x) = x$ für $x \in (-\pi, \pi)$ und $f(0) = f(\pm\pi) = 0$. Jedoch ist diese Grenzfunktion f nicht mehr stetig. Man kann den normierten Vektorraum $\mathcal{C}([-\pi, \pi], \mathbb{R})$ durch Hinzufügen der Grenzwerte solcher Cauchy-Folgen vervollständigen. $\square$

FOLGERUNG 44 (Parsevalsche Gleichung). Es sei $\{v_j : j \in \mathbb{N}\}$ ein vollständiges Orthonormalsystem. Dann gilt für jedes $v \in V$ die Parsevalsche Gleichung

$$\sum_{j=0}^{\infty} |\langle v, v_j \rangle|^2 = \|v\|^2.$$

BEWEIS. Der Vektor $v - \sum_{j \in J} \langle v, v_j \rangle v_j$ ist senkrecht zu allen v_j . Da das Orthonormalsystem vollständig ist, muss er der Nullvektor sein, also

$$\begin{aligned} 0 &= \|v - \sum_{j \in \mathbb{N}} \langle v, v_j \rangle v_j\| = \langle v - \sum_{j \in \mathbb{N}} \langle v, v_j \rangle v_j, v - \sum_{j \in J} \langle v, v_j \rangle v_j \rangle \\ &= \|v\|^2 - \sum_{j \in \mathbb{N}} \langle v, v_j \rangle \langle v_j, v \rangle - \sum_{j \in \mathbb{N}} \overline{\langle v, v_j \rangle} \langle v, v_j \rangle + \sum_{j \in \mathbb{N}} \langle v, v_j \rangle \overline{\langle v, v_j \rangle} \\ \sum_{j \in \mathbb{N}} |\langle v, v_j \rangle|^2 &= \|v\|^2, \end{aligned}$$

da $\langle v, v_j \rangle = \overline{\langle v_j, v \rangle}$. □

Bisher haben wir für den Vektorraum $\mathcal{C}([-\pi, \pi], \mathbb{C})$ mit dem hermiteschen Skalarprodukt

$$\langle f, g \rangle = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx$$

die vollständigen Orthonormalsysteme $\{\sqrt{2}, \sqrt{2} \sin(nx), \sqrt{2} \cos(nx) : n \in \mathbb{N}\}$ und $\{e^{inx} : n \in \mathbb{Z}\}$ kennen gelernt. Die Koordinaten einer Funktion $f \in \mathcal{C}([-\pi, \pi], \mathbb{C})$ bezüglich dieser Orthonormalbasis sind die Fourierkoeffizienten und man erhält die Fourierreihenentwicklung

$$f(x) = \sum_{n \in \mathbb{Z}} \langle f, e^{inx} \rangle e^{inx}.$$

2.6. Approximation stetiger Funktionen durch Polynome. Wir wollen stetige, reellwertige Funktionen durch Polynome approximieren. Der Weierstraßsche Approximationssatz sichert die Existenz einer approximierenden Folge von Polynomen, erklärt aber nicht, wie man eine solche Folge konstruiert.

SATZ 131 (Weierstraßscher Approximationssatz). *Zu jedem $f \in \mathcal{C}([a, b], \mathbb{R})$ gibt es eine Folge von Polynomen $\{p_n\}_{n \in \mathbb{N}}$, die gleichmäßig gegen f konvergiert.*

BEMERKUNG 96. Wenn die Folge $\{p_n\}_{n \in \mathbb{N}}$ gleichmäßig gegen f konvergiert, so konvergiert sie auch im quadratischen Mittel gegen f . □

Die Menge der Monome $\{x^j : j \in \mathbb{N}\} = \{1, x, x^2, \dots\}$ ist linear unabhängig. Deshalb versuchen wir, aus dieser Menge ein vollständiges Orthonormalsystem des Vektorraumes $\mathcal{C}([a, b], \mathbb{R})$ zu konstruieren. Wir können uns auf die Betrachtung des Intervalls $[-1, 1]$ beschränken, weil durch die Substitution $y = -1 + 2(x-a)/(b-a)$ ein Isomorphismus $\mathcal{C}([a, b], \mathbb{R}) \cong \mathcal{C}([-1, 1], \mathbb{R})$ definiert wird

2.6.1. *Legendresche Polynome.* Zur Darstellung elektrostatischer Potentiale benötigt man Legendresche Polynome.

Wir beginnen, mit Hilfe des Schmidtschen Orthonormalisierungsverfahrens ein Orthonormalsystem $\{\tilde{h}_n(x) : n \in \mathbb{N}\}$ in $\mathcal{C}([-1, 1], \mathbb{R})$ mit dem Skalarprodukt $\langle f, g \rangle = \int_{-1}^1 f(x)g(x) dx$ aus der Basis $\{h_n(x) = x^n : n \in \mathbb{N}\}$ zu konstruieren:

$$\begin{aligned} \|h_0\|^2 &= \int_{-1}^1 1 dx = 2, & \tilde{h}_0(x) &= \frac{1}{\sqrt{2}} \\ \langle h_1, \tilde{h}_0 \rangle &= \frac{1}{\sqrt{2}} \int_{-1}^1 x dx = 0 \\ \|h_1\|^2 &= \int_{-1}^1 x^2 dx = \frac{1}{3} [x^3]_{-1}^1 = \frac{2}{3}, & \tilde{h}_1(x) &= \sqrt{\frac{3}{2}} x \\ \langle h_2, \tilde{h}_0 \rangle &= \frac{1}{\sqrt{2}} \int_{-1}^1 x^2 dx = \frac{1}{3\sqrt{2}} [x^3]_{-1}^1 = \frac{2}{3\sqrt{2}} \\ \langle h_2, \tilde{h}_1 \rangle &= \sqrt{\frac{3}{2}} \int_{-1}^1 x^3 dx = 0 \\ \|h_2 - \frac{2}{3\sqrt{2}} \tilde{h}_0\|^2 &= \int_{-1}^1 \left(x^2 - \frac{1}{3}\right)^2 dx = \int_{-1}^1 x^4 - \frac{2}{3}x^2 + \frac{1}{9} dx = \left[\frac{x^5}{5} - \frac{2x^3}{9} + \frac{x}{9}\right]_{-1}^1 \\ &= \frac{2}{5} - \frac{2}{9} = \frac{8}{45} \\ \tilde{h}_2(x) &= \frac{3\sqrt{5}}{2\sqrt{2}} \left(x^2 - \frac{1}{3}\right) = \sqrt{\frac{5}{2}} \left(\frac{3}{2}x^2 - \frac{1}{2}\right) \end{aligned}$$

Die Funktionen $\tilde{h}_0$, $\tilde{h}_1$ und $\tilde{h}_2$ bilden ein Orthonormalsystem. Wir brechen das Schmidtsche Orthonormalisierungsverfahren hier ab und zeigen, dass die Funktionen $\tilde{h}_n$ bis auf einen Normalisierungsfaktor die Legendreschen Polynome sind.

Für $n \in \mathbb{N}$ ist das **Legendresche Polynom** L_n vom Grad n definiert durch

$$L_n(x) = \frac{1}{2^n n!} \frac{d^n}{dx^n} ((x^2 - 1)^n).$$

SATZ 132. *Die normierten Legendreschen Polynome*

$$\left\{ \frac{\sqrt{2n+1}}{\sqrt{2}} L_n : n \in \mathbb{N} \right\}$$

bilden ein vollständiges Orthonormalsystem in $\mathcal{C}([-1, 1], \mathbb{R})$ bezüglich des Skalarproduktes $\langle f, g \rangle = \int_{-1}^1 f(x)g(x)dx$.

BEWEIS. Wir zeigen zuerst, dass die Legendreschen Polynome orthogonal zueinander sind. Für $n > 0$ gilt

$$\langle L_n, 1 \rangle = \int_{-1}^1 L_n(x) dx = \frac{1}{2^n n!} \left[((x^2 - 1)^n)^{(n-1)} \right]_{-1}^1 = 0,$$

da $x = \pm 1$ Nullstellen von $((x^2 - 1)^n)^{(n-1)}$ sind. Für $n > m > 0$ erhalten wir

$$\begin{aligned}\langle L_m, L_n \rangle &= \int_{-1}^1 L_m(x) L_n(x) dx = \frac{1}{2^{n+m} m! n!} \int_{-1}^1 ((x^2 - 1)^n)^{(n)} ((x^2 - 1)^m)^{(m)} dx \\ &= \frac{1}{2^{n+m} m! n!} \left(\left[((x^2 - 1)^n)^{(n-1)} ((x^2 - 1)^m)^{(m)} \right]_{-1}^1 \right. \\ &\quad \left. - \int_{-1}^1 ((x^2 - 1)^n)^{(n-1)} ((x^2 - 1)^m)^{(m+1)} dx \right) \\ &= -\frac{(-1)^{m+1}}{2^{n+m} m! n!} \int_{-1}^1 ((x^2 - 1)^n)^{(n-m-1)} ((x^2 - 1)^m)^{(2m+1)} dx = 0,\end{aligned}$$

da $x = \pm 1$ Nullstellen von $((x^2 - 1)^n)^{(n-k)}$ sind und $((x^2 - 1)^m)^{(2m+1)} = 0$.

Durch partielle Integration läßt sich auch die Norm $\|L_m\|$ berechnen:

$$\begin{aligned}\|L_m\|^2 &= \frac{1}{2^{2m} (m!)^2} \int_{-1}^1 ((x^2 - 1)^m)^{(m)} ((x^2 - 1)^m)^{(m)} dx \\ &= \frac{1}{2^{2m} (m!)^2} (-1)^m \int_{-1}^1 ((x^2 - 1)^m)^{(2m)} ((x^2 - 1)^m) dx \\ &= \frac{(2m)!}{2^{2m} (m!)^2} (-1)^m \int_{-1}^1 (x-1)^m (x+1)^m dx \\ &= \frac{(2m)!}{2^{2m} (m!)^2} \frac{m!}{(m+1) \dots 2m} \int_{-1}^1 (x-1)^{2m} (x+1)^0 dx \\ &= \frac{1}{2^{2m}} \int_{-1}^1 (x-1)^{2m} dx = \frac{1}{2^{2m}} \frac{1}{2m+1} [(x-1)^{2m+1}]_{-1}^1 = \frac{2}{2m+1}\end{aligned}$$

Da das Legendresche Polynom L_n den Grad n hat, liegen alle Monome x^j mit $j \in \mathbb{N}$ im Vektorraum, der von den Legendreschen Polynomen erzeugt wird. Aus dem Weierstraßschen Approximationssatz folgt dann, dass die normierten Legendreschen Polynome ein vollständiges Orthonormalsystem in $\mathcal{C}([-1, 1], \mathbb{R})$ bilden. $\square$

Mit der Leibnizregel für die Ableitung erhalten wir

$$\begin{aligned}L'_{n+1}(x) &= \frac{1}{2^{n+1} (n+1)!} ((x^2 - 1)^{n+1})^{(n+2)} = \frac{1}{2^n n!} ((x^2 - 1)^n x)^{(n+1)} \\ &= \frac{1}{2^n n!} \left(x((x^2 - 1)^n)^{(n+1)} + (n+1)((x^2 - 1)^n)^{(n)} \right) \\ &= x L'_n(x) + (n+1) L_n(x)\end{aligned}$$

und

$$\begin{aligned}
 L'_{n+1}(x) &= \frac{1}{2^{n+1}(n+1)!} \left(((x^2-1)^n(x^2-1))^{(n+1)} \right)' \\
 &= \frac{1}{2^{n+1}(n+1)!} \left((x^2-1)((x^2-1)^n)^{(n+1)} + 2x(n+1)((x^2-1)^n)^{(n)} \right. \\
 &\quad \left. + n(n+1)((x^2-1)^n)^{(n-1)} \right)' \\
 &= \frac{x^2-1}{2(n+1)} L''_n(x) + \frac{(n+2)x}{n+1} L'_n(x) + \frac{n+2}{2} L_n(x).
 \end{aligned}$$

Aus diesen beiden Gleichungen für $L'_{n+1}(x)$ ergibt sich *Legendresche Differentialgleichung*

$$(171) \quad 0 = (x^2-1)L''_n(x) + 2xL'_n(x) - n(n+1)L_n(x).$$

Ähnlich können wir eine *Rekursionsformel*

$$(172) \quad (n+1)L_{n+1}(x) = (2n+1)xL_n(x) - nL_{n-1}(x)$$

für die Legendreschen Polynome beweisen:

$$\begin{aligned}
 (n+1)L_{n+1}(x) &= \frac{1}{2^{n+1}n!} ((x^2-1)^{n+1})^{(n+1)} \\
 &= \frac{1}{2^n n!} ((n+1)x(x^2-1)^n)^{(n)} \\
 &= \frac{1}{2^n n!} ((2n+1)x(x^2-1)^n - nx(x^2-1)^n)^{(n)} \\
 &= \frac{1}{2^n n!} \left((2n+1)x((x^2-1)^n)^{(n)} + (2n+1)n((x^2-1)^n)^{(n-1)} \right. \\
 &\quad \left. - n^2((x^2-1)^{n-1}2x^2)^{(n-1)} - n((x^2-1)^n)^{(n-1)} \right) \\
 &= (2n+1)xL_n(x) - nL_{n-1}(x)
 \end{aligned}$$

Abbildung 7 zeigt die Graphen der ersten vier Legendreschen Polynome.

BEISPIEL 269 (Approximation von e^x durch quadratische Polynome). Die Linearkombination $\sum_{j=0}^2 a_j \tilde{h}_j(x)$ mit $a_j = \langle f, \tilde{h}_j \rangle$ ist die beste Approximation von $f(x) = e^x$ im quadratischen Mittel durch quadratische Polynome. Wir berechnen die Koeffizienten a_0 , a_1 und a_2

$$\begin{aligned}
 \langle e^x, \tilde{h}_0 \rangle &= \int_{-1}^1 e^x \frac{1}{\sqrt{2}} dx = \frac{1}{\sqrt{2}} [e^x]_{-1}^1 = \frac{1}{\sqrt{2}} (e - e^{-1}) \\
 \langle e^x, \tilde{h}_1 \rangle &= \int_{-1}^1 e^x \sqrt{\frac{3}{2}} x dx = \sqrt{\frac{3}{2}} \int_{-1}^1 e^x x dx = \sqrt{\frac{3}{2}} \left([e^x x]_{-1}^1 - \int_{-1}^1 e^x dx \right) \\
 &= \sqrt{\frac{3}{2}} (e + e^{-1} - e^x + e^{-1}) = 2\sqrt{\frac{3}{2}} e^{-1} \\
 \langle e^x, \tilde{h}_2 \rangle &= \frac{1}{2} \sqrt{\frac{5}{2}} \int_{-1}^1 e^x (3x^2 - 1) dx = \frac{1}{2} \sqrt{\frac{5}{2}} \left([-e^x]_{-1}^1 + [e^x 3x^2]_{-1}^1 - \int_{-1}^1 6xe^x dx \right)
 \end{aligned}$$

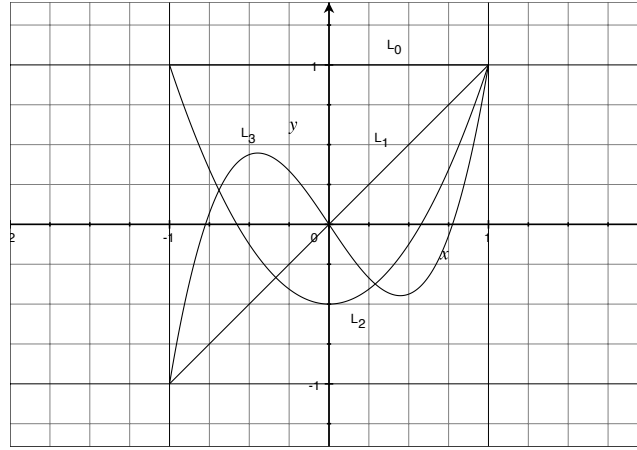


ABBILDUNG 7. Die ersten vier Legendreschen Polynome

$$\begin{aligned}
 &= \frac{1}{2} \sqrt{\frac{5}{2}} \left(-e + e^{-1} + 3e - 3e^{-1} - [6xe^x]_{-1}^1 + \int_{-1}^1 6e^x dx \right) \\
 &= \frac{1}{2} \sqrt{\frac{5}{2}} (2e - 2e^{-1} - 6e - 6e^{-1} + 6e - 6e^{-1}) = \frac{1}{2} \sqrt{\frac{5}{2}} (2e - 14e^{-1})
 \end{aligned}$$

und erhalten die quadratische Approximation

$$\begin{aligned}
 e^x &\sim \frac{1}{\sqrt{2}}(e - e^{-1}) \frac{1}{\sqrt{2}} + 2\sqrt{\frac{3}{2}}e^{-1} \sqrt{\frac{3}{2}}x + \frac{1}{2}\sqrt{\frac{5}{2}}(2e - 14e^{-1}) \frac{1}{2}\sqrt{\frac{5}{2}}(3x^2 - 1) \\
 e^x &\sim \frac{1}{2}(e - e^{-1}) + 3e^{-1}x + \frac{5}{8}(2e - 14e^{-1})(3x^2 - 1) \\
 e^x &\sim -\frac{3}{4}e + \frac{33}{4}e^{-1} + 3e^{-1}x + \frac{15}{4}(e - 7e^{-1})x^2
 \end{aligned}$$

2.6.2. Tschebyscheffsche Polynome. Tschebyscheffsche Polynome treten bei Tschebyscheff-Filtern der Elektrotechnik im Nenner der Übertragungsfunktion auf. Wir betrachten den Vektorraum $\mathcal{C}([-1, 1], \mathbb{R})$ mit dem Skalarprodukt

$$\langle f, g \rangle = \int_{-1}^1 \frac{f(x)g(x)}{\sqrt{1-x^2}} dx.$$

Die **Tschebyscheffschen Polynome** $T_n(x) = \cos(n \arccos x)$ sind auf dem Intervall $[-1, 1]$ glatte Funktionen. Sie erfüllen die Orthogonalitätsrelationen $\langle T_n, T_m \rangle = 0$ für alle $n \neq m$, die Differentialgleichung $(1-x^2)T_n''(x) - xT_n'(x) + n^2T_n(x) = 0$ und die Rekursionsformel $T_{n+1}(x) = 2xT_n(x) - T_{n-1}(x)$, also $T_0(x) \equiv 1$, $T_1(x) = x$, $T_2(x) = 2x^2 - 1$, $T_3(x) = 2x(2x^2 - 1) - x = 4x^3 - 3x$. Abbildung 9 zeigt die ersten vier Tschebyscheffschen Polynome.

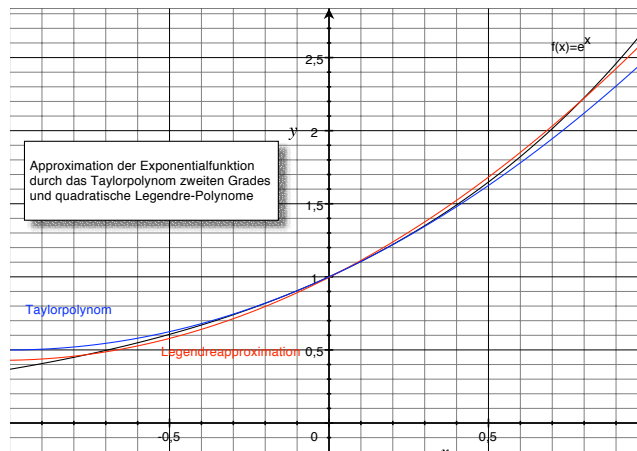


ABBILDUNG 8. Approximation der Exponentialfunktion (schwarz) durch das Taylorpolynom zweiten Grades (blau) und quadratische Legendre-Polynome (rot)

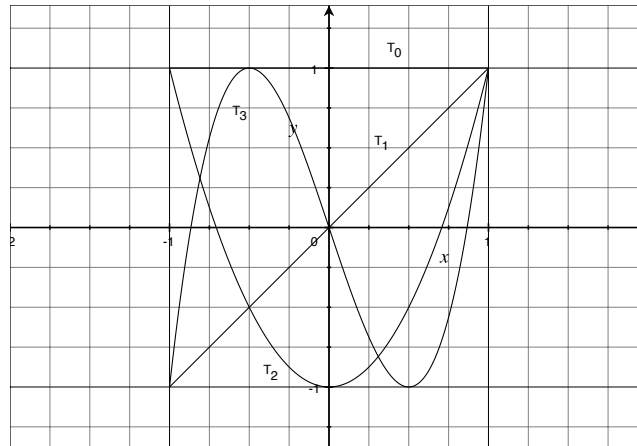


ABBILDUNG 9. Die ersten vier Tschebyscheffschen Polynome