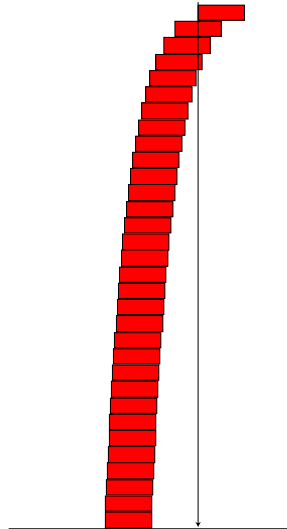


Analysis I für Ingenieure
(Mit Beispielmateriale aus der Elektrotechnik)



Information. Für die erfolgreiche Teilnahme an diesem Modul erhalten Sie

8 Leistungspunkte nach ECTS.

Entsprechend erwarten wir von durchschnittlich begabten und vorgebildeten Studierenden folgenden Arbeitsaufwand:

Vorlesung	4h/Woche
Übung	2h/Woche
Häusliche Nacharbeit und Hausaufgaben	8h/Woche
Klausurvorbereitung	30h

Inhaltsverzeichnis

1	Vorbereitungen	7
1.1	Mengen, Abbildungen	7
2	Zahlbereiche	13
2.1	Natürliche Zahlen, Vollständige Induktion	13
2.2	Ganze, rationale und reelle Zahlen	20
2.3	Komplexe Zahlen	25
3	Grenzwerte und Stetigkeit	32
3.1	Zahlenfolgen und Konvergenz	33
3.2	Konvergenzbeweise	37
3.3	Stetigkeit von Funktionen	41
4	Elementare Funktionen I	47
4.1	Polynome	47
4.2	Rationale Funktionen	51
5	Differentiation	56
5.1	Die Ableitung	56
5.2	Extremwerte, Mittelwertsatz und Konsequenzen	61
5.3	Höhere Ableitungen	64
5.4	Das Taylorpolynom	65
6	Elementare Funktionen II	70
6.1	Trigonometrische Funktionen I	71
6.2	Trigonometrische Funktionen II	76
6.3	Exponentialfunktion und Logarithmus	81
6.4	Allgemeine Potenz und Hyperbelfunktionen	87
7	Weitere Anwendungen der Differentiation	92
8	Integration	96
8.1	Das bestimmte Integral	96
8.2	Das unbestimmte Integral	100
8.3	Integrationsregeln	104
8.4	Integration komplexer und rationaler Funktionen	110
8.5	Uneigentliche Integrale	113
9	Fourieranalysis	116

9.1	Reelle Fourieranalysis	116
9.2	Anwendungsbeispiele	120
9.3	Approximation im quadratischen Mittel	127
9.4	Komplexe Fourieranalysis	130
10	Unendliche Reihen	132
10.1	Reihen mit konstanten Gliedern	132
10.2	Weitere Konvergenzkriterien	137
10.3	Funktionenreihen	141
11	Anhang	148
11.1	Inverse der RSA-Kodierung	148
11.2	Komplexe Zahlen in der Elektrotechnik: Beispiele	149
11.3	Konvergenz numerisch	151
11.4	Lineare Differentialgleichungen mit konstanten Koeffizienten	152
11.5	Regel von Bernoulli/de l' Hospital	155
11.6	Zum Integralbegriff	158
11.7	Von der Summe zum Integral	159
11.8	Numerische Integration.	160
11.9	Was sind Orthogonalitätsrelationen?	161

Literatur

Als Lehrbücher zu dieser Veranstaltung werden empfohlen:

- **Meyberg, Vachenaue:** **Höhere Mathematik 1**, Springer Verlag
- **G. Bärwolff, G. Seifert:** **Höhere Mathematik für Naturwissenschaftler und Ingenieure**, Spektrum Akademischer Verlag

Farbig unterlegt finden Sie Beispiel aus den Ingenieur Anwendungen, oft mit expliziten Hinweisen auf Ingenieurskripten des Grundstudiums:

Werkstoffe I

Werkstoffe und Bauelemente der Elektrotechnik I,
Skript TUB, Institut für Werkstoffe der Elektrotechnik

Werkstoffe II

Werkstoffe und Bauelemente der Elektrotechnik II,
Skript TUB, Institut für Werkstoffe der Elektrotechnik

Regelungstechnik I

King: Regelungstechnik I,
Skript TUB

Müller: Mechanik II

W.H. Müller: Mechanik II,
Skript TUB SS 2002

Energie-, Impuls- und Stofftransport

Lehrbuch *Baehr/Stephan: Wärme- und Stoffübertragung*
zur VL Auracher: Energie-, Impuls- und Stofftransport

Verfahrenstechnik I

Kraume: Verfahrenstechnik I,
Skript TUB

Signale und Systeme

Noll: Signale und Systeme,
Skript 1999

Einführung in die Netzwerktheorie

Naunin: Einführung in die Netzwerktheorie,
Skript TUB 1998

Zu einzelnen Veranstaltungen existieren mehrere, vielleicht auch neuere Skripten und Lehrbücher. Wir begnügen uns meistens mit *einem* Zitat, auch wenn sich der betreffende Sachverhalt in der Regel an mehreren Stellen findet.

Das Griechische Alphabet

α	A	Alpha	ι	I	Iota	ρ	P	Rho
β	B	Beta	κ	K	Kappa	σ	Σ	Sigma
γ	Γ	Gamma	λ	Λ	Lambda	τ	T	Tau
δ	Δ	Delta	μ	M	My	υ	Y	Ypsilon
ϵ, ε	E	Epsilon	ν	N	Ny	ϕ, φ	Φ	Phi
ζ	Z	Zeta	ξ	Ξ	Xi	χ	X	Chibert
η	H	Eta	o	O	Omikron	ψ	Ψ	Psi
θ, ϑ	Θ	Theta	π	Π	Pi	ω	Ω	Omega

1 Vorbereitungen

- Wir erinnern uns an die Begriffe und Symbole der Mengen-Sprache.
- Wir erinnern uns an den Begriff der *Abbildung*, der den Begriff der *Funktion* auf den Fall erweitert, wo die Zuordnung nicht Zahlen in Zahlen überführt.
- Wir lernen spezielle Eigenschaften von Abbildungen ...
- ... und verschiedene Beispiel von Abbildungen kennen.

1.1 Mengen. Abbildungen

MENGEN.

Wir geben hier eine kurze Einführung in die Sprache der Mengenlehre und die elementaren Operationen mit Mengen.

Ist A eine Menge und a ein Element von A , so schreiben wir

$$a \in A$$

und sagen auch “ a ist in A ”, “ a liegt in A ”, “ a in A ” oder “ a Element A ”. Gelegentlich wird auch $A \ni a$ verwendet. Ist a nicht Element von A , so schreibt man $a \notin A$.

Mengen kann man beschreiben, indem man ihre Elemente angibt. Dafür verwendet man üblicherweise geschweifte Klammern:

$$\begin{aligned} \mathbb{N} &:= \{0, 1, 2, 3, \dots\} && \text{ist die Menge der natürlichen Zahlen} \\ G &:= \{0, 2, 4, 6, \dots\} && \text{ist die Menge der geraden natürlichen Zahlen.} \end{aligned}$$

Das schreibt man auch so:

$$G = \{n \in \mathbb{N} \mid n \text{ gerade} \}.$$

Zum Beispiel ist

$$\mathbb{R} = \{x \mid x \text{ ist eine reelle Zahl} \}$$

die Menge der reellen Zahlen, oder

$$\sin \in \{f \mid f \text{ reelle Funktion und } f'' = -f\}.$$

Ein nützliches Ding ist die *leere Menge* $\emptyset$, die Menge, die überhaupt kein Element enthält. Die „nächst größeren“ Mengen sind die einelementigen, die nur ein einziges Element enthalten, z.B. $\{7\}$. Wir unterscheiden meistens sorgfältig zwischen der Zahl 7 und der Menge, die mit dieser Zahl als einzigem Element gebildet ist. Wieviele Elemente hat $\{\emptyset\}$?

Die Symbole $\infty = +\infty$ und $-\infty$ sind keine reellen Zahlen:

$$+\infty \notin \mathbb{R}, \quad -\infty \notin \mathbb{R}.$$

Sind A, B zwei Mengen, so nennt man B eine *Teilmenge* von A , wenn jedes Element von B auch in A liegt. Notation:

$$B \subset A \quad \text{oder} \quad A \supset B.$$

Beachten Sie, dass mit dieser Definition $A \subset A$. Was soll wohl $B \not\subset A$ bedeuten?

Beispiel 1 (Intervalle). Ein *Intervall* ist eine Teilmenge $I \subset \mathbb{R}$ der reellen Zahlen von einer der folgenden Formen, wobei $a, b \in \mathbb{R}$ mit $a \leq b$:

$$\begin{aligned}]a, b[&:= \{x \in \mathbb{R} \mid a < x < b\} && \text{(offenes Intervall)} \\ [a, b[&:= \{x \in \mathbb{R} \mid a \leq x < b\} && \text{(halboffenes Intervall)} \\]a, b] &:= \{x \in \mathbb{R} \mid a < x \leq b\} && \text{(halboffenes Intervall)} \\ [a, b] &:= \{x \in \mathbb{R} \mid a \leq x \leq b\} && \text{(kompaktes Intervall)} \end{aligned}$$

Weiter bezeichnet man auch folgende Mengen als Intervalle:

$$\begin{aligned}]a, +\infty[&:= \{x \in \mathbb{R} \mid a < x\} && \text{(offenes Intervall)} \\ [a, +\infty[&:= \{x \in \mathbb{R} \mid a \leq x\} && \text{(halboffenes Intervall)} \\]-\infty, b[&:= \{x \in \mathbb{R} \mid x < b\} && \text{(offenes Intervall)} \\]-\infty, b] &:= \{x \in \mathbb{R} \mid x \leq b\} && \text{(halboffenes Intervall)} \end{aligned}$$

und

$$]-\infty, +\infty[:= \mathbb{R}.$$

□

Aus zwei Mengen A, B bildet man neue Mengen, die *Vereinigung(smenge)* und den *Durchschnitt* (den die Schulmathematiker „Schnittmenge“ nennen, um keine Verwechslung z.B. mit der Durchschnittsnote einer Klausur zu provozieren):

$$\begin{aligned} A \cup B &:= \{x \mid x \in A \text{ oder } x \in B\} && \text{„}A \text{ vereinigt } B\text{“} \\ A \cap B &:= \{x \mid x \in A \text{ und } x \in B\} && \text{„}A \text{ geschnitten } B\text{“} \end{aligned}$$

Hier – wie üblicherweise in der Mathematik – ist „oder“ nicht ausschließend (nicht „entweder ... oder“). Diese Operationen kann man auch mit mehr als zwei Mengen durchführen, sogar mit einer ganzen Familie von Mengen A_i , wobei der *Index* i eine gewisse Indexmenge I durchläuft. Man schreibt dann

$$\begin{aligned} \bigcup_{i \in I} A_i &:= \{x \mid x \in A_i \text{ für wenigstens ein } i \in I\}, \\ \bigcap_{i \in I} A_i &:= \{x \mid x \in A_i \text{ für alle } i \in I\}. \end{aligned}$$

Auch die Differenzmenge $A \setminus B := \{a \in A \mid a \notin B\}$ kommt vor, z.B. ist

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

Beispiel 2. Wir bezeichnen mit A_i die Menge der natürlichen Zahlen, die durch i teilbar sind. Dabei sei i aus der Menge $P := \mathbb{N} \setminus \{0\}$ der natürlichen Zahlen ohne die Null, und sei $J := P \cap G = \{2, 4, 6, \dots\}$ die Menge der positiven geraden Zahlen. Dann gilt

$$A_2 \cap A_3 = A_6, \quad \bigcup_{j \in J} A_j = A_2, \quad \bigcap_{j \in J} A_j = \{0\}.$$

□

Überlegen Sie sich, dass

$$A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$$

gilt. Das sieht aus wie das Distributivgesetz

$$a(b + c) = ab + ac.$$

Tatsächlich kann man eine ganze Reihe solcher Regeln herleiten und dann mit Mengenoperationen nach diesen Regeln einfach rechnen, ohne jedesmal inhaltlich zu überlegen, was das für die Elemente bedeutet. („Boolesche Algebra“ nennt man solche Rechnungen oder, genauer, einen Bereich mit Verknüpfungen $\cap$ und $\cup$, in dem man nach diesen Regeln rechnen kann.)

Wir schließen mit zwei weiteren Begriffen, von denen vor allem der zweite oft vorkommt: Die Gesamtheit der Teilmengen einer festen Menge A bildet eine Menge $\mathcal{P}(A)$, die sogenannte *Potenzmenge* von A :

$$\mathcal{P}(A) := \{B \mid B \subset A\}.$$

Also gilt $B \in \mathcal{P}(A)$ genau dann, wenn $B \subset A$. Schließlich sei für Mengen A, B die Produktmenge $A \times B$ definiert als die Menge der geordneten Paare (a, b) von Elementen aus A bzw B :

$$A \times B := \{(a, b) \mid a \in A \text{ und } b \in B\}.$$

Allgemeiner definiert man

$$A_1 \times \dots \times A_n := \{(a_1, \dots, a_n) \mid a_i \in A_i\}$$

als die Menge der geordneten n -tupel.

Beispiel 3. Seien $A = \{1, 2\}, B = \{1, 2, 3\}$. Dann ist

$$\begin{aligned}\mathcal{P}(A) &= \{\emptyset, \{1\}, \{2\}, \{1, 2\}\} \\ A \times B &= \{(1, 1), (1, 2), (1, 3), (2, 1), (2, 2), (2, 3)\}.\end{aligned}$$

□

Beispiel 4. An manchen elektronischen Geräten befinden sich einige sog. Dip-Schalter, die einzeln ein- bzw. ausgeschaltet werden können. Der Schaltzustand ist charakterisiert durch die genaue Angabe aller eingeschalteten Dip-Schalter. Sind die einzelnen Schalter durchnummeriert von 1 bis n , so ist der Schaltzustand durch die Nummern der eingeschalteten Schalter gegeben also eine Teilmenge von $\{1, \dots, n\}$. Die Menge aller möglichen Dip-Schaltereinstellungen ist dann gerade durch die Potenzmenge von $\{1, 2, \dots, n\}$ gegeben. Es gibt übrigens 2^n solcher Teilmengen, angefangen von der leeren Menge (kein Schalter eingeschaltet) bis zur gesamten Menge (alle Schalter eingeschaltet).

ABBILDUNGEN

Definition 5. Seien A, B Mengen. Eine Abbildung f von A in (oder nach) B ist eine Vorschrift, die *jedem* Element $x \in A$ *genau ein* Element $y = f(x) \in B$ zuordnet.

Man nennt A den *Definitionsbereich* und B den *Wertebereich* von f .

Notationen:

$$\begin{aligned}f &: A \rightarrow B \\ f &: A \rightarrow B, \quad x \mapsto f(x) \\ f &: A \ni x \mapsto f(x) \in B \\ f &: x \mapsto f(x), \quad \text{wenn } A \text{ und } B \text{ klar sind.}\end{aligned}$$

Man nennt Abbildungen auch *Funktionen*, vor allem dann, wenn die Werte Zahlen sind.

Beispiel 6. Die Funktion

$$\sin : \mathbb{R} \rightarrow [-1, 1].$$

□

Beispiel 7. Sehr gebräuchlich ist es, Funktionen in der Form

$$f(x) := \sin \frac{1}{x}, \quad x > 0$$

anzugeben, obwohl man eigentlich penibel zwischen der Funktion f und dem Funktionswert $f(x)$ unterscheiden muss. Hier ist der Definitionsbereich $A =]0, +\infty[$, den Wertebereich B muss man raten (Könnte z.B. $= \mathbb{R}$ sein. Oder was sonst?), und die Abbildungsvorschrift ist in obiger Notation gegeben durch $x \mapsto \sin \frac{1}{x}$. □

Beispiel 8. Sei $n = 2773$ und seien $A = B = \{0, 1, 2, \dots, 2772\}$. Definiere $f : A \rightarrow A$ für $x \in A$ folgendermaßen:

Teile x^{17} durch n . Der Rest, wieder eine Zahl in A , sei $f(x)$. Von dieser Form sind Codierungsfunktionen bei der sogenannten RSA-Codierung.

Beispiel 9. An einer Schaltung wird in einem Zeitintervall von t_0 bis t_1 die zeitabhängige Spannung $U(t)$ gemessen. Dann ist $A = \{t \mid t_0 \leq t \leq t_1\}$, $B = \mathbb{R}$ die Menge der reellen Zahlen und $U : A \rightarrow \mathbb{R}$.

Wir werden es in diesem Semester vor allem mit Abbildungen von $\mathbb{R}$ nach $\mathbb{R}$ zu tun haben, in späteren Semestern werden dann auch Abbildungen mehrdimensionaler Räume behandelt.

Beachten Sie,

- Nach unserer Definition wird einem x nur ein Wert $y = f(x)$ zugeordnet, so dass also zum Beispiel $x \mapsto \pm\sqrt{x}$ keine Funktion ist.
- Sehr wohl kann aber verschiedenen x -Werten derselbe y -Wert zugeordnet sein kann. Der krasseste Fall sind die *konstanten Abbildungen*, die jedem x dasselbe y zuordnen.
- Weiter muss die Menge $f(A) := \{f(x) \mid x \in A\}$ der tatsächlich angenommenen Werte, das sogenannte *Bild von f* , nicht unbedingt ganz B sein.
Zum Beispiel hat $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ das Bild $f(\mathbb{R}) = [0, \infty[:= \{y \in \mathbb{R} \mid y \geq 0\}$.

Die beiden letzten Bemerkungen motivieren die folgenden Begriffsbildungen.

Definition 10. Sei $f : A \rightarrow B$ eine Abbildung. $f : A \rightarrow B$ heißt *injektiv*, wenn $f(x_1) = f(x_2)$ nur für $x_1 = x_2$. (Verschiedene Argumente liefern verschiedene Werte).

Beachten Sie, dass bei dieser Definition nicht nur die “Abbildungsvorschrift” f , sondern ganz entscheidend auch der *jeweils betrachtete* Definitionsbereich eine Rolle spielt.

Beispiel 11. Die Abbildung

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto x^2$$

ist nicht injektiv, weil z.B. $f(-2) = 4 = f(2)$. Aber

$$g : [0, \infty[\rightarrow \mathbb{R}, \quad x \mapsto x^2$$

ist injektiv. Man nennt g auch die *Einschränkung* von f auf $[0, \infty[$ und schreibt dafür

$$g = f|_{[0, \infty[}.$$

□

Definition 12 (Umkehrabbildung = Inverse). Ist f injektiv, so gibt es zu jedem $y \in f(A)$ genau ein $x \in A$ mit $f(x) = y$. Man kann dann also auf dem Bild $f(A)$ eine *Umkehrabbildung* oder *Inverse* $h : B \supset f(A) \rightarrow A$ definieren, die durch

$$h(y) = x \iff f(x) = y$$

charakterisiert ist. Man bezeichnet diese Umkehrabbildung statt mit h auch mit

$$f^{-1} : f(A) \rightarrow A.$$

Beispiel 13. Die im Beispiel 8 angegebene Abbildung

$$f : \{0, 1, 2, \dots, 2772\} \rightarrow \{0, 1, 2, \dots, 2772\}$$

ist injektiv mit $f(\{0, 1, 2, \dots, 2772\}) = \{0, 1, 2, \dots, 2772\}$. Die Umkehrfunktion $h = f^{-1}$ liefert die Decodierung. h ist von derselben Gestalt wie f , allerdings mit einem anderen Exponenten als 17. Dieser ist sehr schwer zu finden, und das ist das “Geheimnis” dieser Codierungsmethode. (Vgl. z.B. das sehr hübsche Buch: *S. Singh, Geheime Botschaften, Hanser 2000* und den Anhang Abschnitt 11.1).

Definition 14 (Urbildmenge). Ist f ganz beliebig, und ist $Y \subset B$ eine Teilmenge von B , so nennt man

$$f^{-1}(Y) := \{x \in A \mid f(x) \in Y\}$$

das *Urbild* von Y .

Jetzt haben wir also zweimal dieselbe Bezeichnung, nämlich f^{-1} , für verschiedene Begriffe verwendet, was man eigentlich in der Mathematik vermeiden soll. Das ist hier nicht ganz so schlimm, weil die Argumente von f^{-1} einmal *Elemente* von B , das andere Mal *Teilmengen* von B sind. Am Argument kann man also erkennen, welches f^{-1} gemeint ist. Welches ist denn gemeint bei $f^{-1}(y)$ und welches bei $f^{-1}(\{y\})$, wenn $y \in B$? Falls f injektiv ist, ist sogar

$$\{f^{-1}(y)\} = f^{-1}(\{y\})$$

für alle y im Bild von f .

Es gibt übrigens noch einen Kandidaten, der Anspruch auf die Bezeichnung f^{-1} erheben könnte, nämlich $1/f$. Wir geben dazu noch ein

Beispiel 15. Die Abbildung

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto x^3$$

ist injektiv mit $f(\mathbb{R}) = \mathbb{R}$. Man hat

$$f^{-1} : \mathbb{R} \rightarrow \mathbb{R}, \quad y \mapsto \sqrt[3]{y},$$

wobei $\sqrt[3]{y} < 0$, falls $y < 0$, und

$$1/f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}, \quad x \mapsto \frac{1}{x^3}.$$

Also ist zum Beispiel

$$\begin{aligned}f^{-1}(-27) &= -3 \\f^{-1}([8, \infty[) &= [2, \infty[\\f(x)^{-1} &= \frac{1}{x^3}.\end{aligned}$$

□

Zum Schluss betrachten wir die Hintereinanderschaltung von Abbildungen.

Definition 16 (Komposition). Seien

$$f : A \rightarrow B, \quad g : B \rightarrow C$$

zwei Abbildungen. Dann heißt die Abbildung

$$g \circ f : A \rightarrow C, x \mapsto g(f(x))$$

die *Komposition* von f und g . (Man liest das „ g nach f “ oder „ g Kringel f “.)

Beispiel 17. Das Volumen eines Heliumballons nimmt mit steigender Höhe zu, da der Luftdruck der umgebenden Luft mit zunehmender Höhe abnimmt. Bezeichnet man das Volumen des Ballons in der Höhe h mit $V(h)$, so wird dadurch eine Funktion $V : \mathbb{R} \rightarrow \mathbb{R}$ gegeben. Misst man zusätzlich die Höhe des Ballons in Abhängigkeit von der Zeit und bezeichnet sie zum Zeitpunkt t mit $h(t)$, so wird dadurch eine weitere Funktion $h : \mathbb{R} \rightarrow \mathbb{R}$ gegeben. Das Volumen des Ballons zum Zeitpunkt t , an dem er sich in der Höhe $h(t)$ befindet, ist dann durch $V(h(t))$ gegeben. Die Komposition $V \circ h : \mathbb{R} \rightarrow \mathbb{R}$ der beiden Funktionen gibt somit das Volumen des Ballons in direkter Abhängigkeit von der Zeit.

Falls g nur auf einer Teilmenge $G \subset B$ definiert ist, kann man $g \circ f$ genauso definieren wie oben, falls $f(A) \subset G$. Liegt aber $f(A)$ nicht ganz in G , so ist $g \circ f$ nur definiert auf der Menge aller $x \in A$, die unter f in G landen:

$$g \circ f : A \supset f^{-1}(G) \rightarrow C.$$

Beispiel 18.

$$\begin{aligned}f : \mathbb{R} \rightarrow \mathbb{R}, \quad x &\mapsto 1 - x^2, & g : [0, \infty[\rightarrow \mathbb{R}, \quad y &\mapsto \sqrt{y} \\g \circ f : [-1, 1] &\rightarrow \mathbb{R}, \quad x &\mapsto \sqrt{1 - x^2}.\end{aligned}$$

□

Beispiel 19. Ist $f : A \rightarrow B$ injektiv mit inverser Abbildung f^{-1} , so ist

$$f^{-1} \circ f : A \rightarrow A$$

die identische Abbildung $\text{Id}_A : A \rightarrow A, \quad x \mapsto x$, und

$$f \circ f^{-1} = \text{Id}_{f(A)} : f(A) \rightarrow f(A), \quad y \mapsto y.$$

□

2 Zahlbereiche

2.1 Natürliche Zahlen. Vollständige Induktion

- Die natürlichen Zahlen dienen oft zum Abzählen oder Nummerieren.
- Summen- und Produktzeichen werden eingeführt.
- Das Prinzip der vollständigen Induktion dient der Herleitung wichtiger Formeln.
- Wir lernen verschiedene Begriffe der elementaren Kombinatorik kennen.

Dieses Kapitel dient nicht einer systematischen Einführung der natürlichen Zahlen, sondern der Sammlung einiger wichtiger Techniken und Begriffe, die mit diesen Zahlen zu tun haben.

Wir bezeichnen die Menge der natürlichen Zahlen mit

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

Die Literatur ist sich nicht ganz einig, ob $0 \in \mathbb{N}$, aber DIN 5473 meint, das sei so, und deshalb wollen wir uns darauf einigen.

„Zählen“ oder „Abzählen“ von Dingen ist eine der wichtigsten Anwendungen von natürlichen Zahlen. In der Mathematik benutzt man das, um verschiedene, aber analoge Dinge bezeichnungstechnisch durch einen *Index* zu unterscheiden:

„Seien $x_1, x_2, \dots, x_n$ reelle Zahlen.“

„Sind $A_1, A_2, A_3, \dots$ Mengen, so ist ihre Vereinigungsmenge definiert als

$$A := \bigcup_{i \in \mathbb{N}} A_i := \{x \mid \text{Es gibt ein } i \in \mathbb{N} \text{ mit } x \in A_i\}.$$

SUMMENZEICHEN

Für $m, n \in \mathbb{N}$ mit $m \leq n$ definiert man das Summenzeichen

$$\sum_{k=m}^n x_k := x_m + x_{m+1} + \dots + x_n.$$

Wichtige Eigenschaften:

$$\begin{aligned} \sum_{k=m}^n x_k + \sum_{k=m}^n y_k &= \sum_{k=m}^n (x_k + y_k), \\ \sum_{k=m}^n a x_k &= a \sum_{k=m}^n x_k, \\ \sum_{i=m}^n x_i \sum_{j=p}^q y_j &= \sum_{i=m}^n \sum_{j=p}^q x_i y_j = \sum_{j=p}^q \sum_{i=m}^n x_i y_j = \sum_{\substack{i=m, \dots, n \\ j=p, \dots, q}} x_i y_j, \\ \sum_{k=m}^n x_k &= \sum_{\mu=m}^n x_\mu = \sum_{i=m}^n x_i, \\ \sum_{k=m}^n x_k &= \sum_{k=m}^p x_k + \sum_{k=p+1}^n x_k, \text{ falls } m \leq p \leq n. \end{aligned}$$

Man setzt $\sum_{k=m}^n x_k = 0$, falls $m > n$.

Beispiel 20 (Momenten-Gleichgewicht).

Müller: Mechanik I, Abschnitt 4

Seien

$$F_{x1}, \dots, F_{xn}, F_{y1}, \dots, F_{yn}, x_1, \dots, x_n, y_1, \dots, y_n \in \mathbb{R}$$

reelle Zahlen, nämlich die Komponenten von n Kräften (F_{xi}, F_{yi}) , die an n -Punkten mit den Koordinaten (x_i, y_i) an einer Scheibe in der Ebene angreifen. Das *Gesamtmoment* bezüglich eines Punktes (a_x, a_y) ist

$$\sum_{i=1}^n (F_{yi}(x_i - a_x) - F_{xi}(y_i - a_y)). \quad (1)$$

Sind die Kräfte im Gleichgewicht d.h. gilt

$$\sum_{i=1}^n F_{xi} = 0 = \sum_{i=1}^n F_{yi},$$

so ist das Gesamtmoment unabhängig von gewählten Punkt (a_x, a_y) . Das wollen wir nachrechnen. Es ist

$$\begin{aligned} \sum_{i=1}^n (F_{yi}(x_i - a_x) - F_{xi}(y_i - a_y)) &= \sum_{i=1}^n F_{yi}(x_i - a_x) - \sum_{i=1}^n F_{xi}(y_i - a_y) \\ &= \sum_{i=1}^n F_{yi}x_i - \sum_{i=1}^n F_{yi}a_x - \sum_{i=1}^n F_{xi}y_i + \sum_{i=1}^n F_{xi}a_y \\ &= \sum_{i=1}^n F_{yi}x_i - a_x \underbrace{\sum_{i=1}^n F_{yi}}_{=0} - \sum_{i=1}^n F_{xi}y_i + a_y \underbrace{\sum_{i=1}^n F_{xi}}_{=0} \\ &= \sum_{i=1}^n F_{yi}x_i - \sum_{i=1}^n F_{xi}y_i. \end{aligned}$$

Offenbar ist das unabhängig von (a_x, a_y) .

Den Aspekt dieses Beispiels könnte man auch einfacher darstellen, ohne doppelte Indizes usw. Wir wollen aber deutlich machen, dass Sie lernen müssen, die Mathematik *im Rahmen Ihres Studienfachs* zu beherrschen. Und da sind zum Beispiel die Bezeichnungen nie darauf abgestellt, die Mathematik möglichst einfach zu gestalten. Die vorstehende Rechnung ist O-Ton des Mechanikskriptes, so müssen Sie mit Summen umgehen können.

Eine oft verwendete Regel ist die „Indexverschiebung“, die wir in einem Spezialfall notieren:

$$\sum_{k=m}^n x_k = \sum_{k=m-1}^{n-1} x_{k+1}.$$

Schreiben Sie beide Seiten aus, und die Gleichheit ist offensichtlich.

Beispiel 21 (Geometrische Summe). Wir wollen die Summe

$$\sum_{k=0}^n x^k = 1 + x + x^2 + \dots + x^n$$

berechnen, wobei das k nun kein Index sein soll, sondern x^k die k -te Potenz von x bezeichnet. Die Summe sei S . Ist $x = 1$, so ist offenbar $S = n + 1$. Andernfalls erhalten wir

$$xS = x + x^2 + \dots + x^{n+1},$$

und daher

$$(1 - x)S = S - xS = (1 + x + x^2 + \dots + x^n) - (x + x^2 + \dots + x^{n+1}) = 1 - x^{n+1}.$$

Insgesamt ergibt sich die sehr wichtige Formel

$$\sum_{k=0}^n x^k = \begin{cases} n + 1, & \text{falls } x = 1 \\ \frac{1-x^{n+1}}{1-x}, & \text{falls } x \neq 1 \end{cases}$$

□

Beispiel 22. Die geometrische Summe ist ein wichtiges Beispiel einer Summenbildung, das auch im Alltag seine Anwendung findet: Wenn man mit einer Bank einen Sparplan über n Jahre abschließt, bei dem zum Beginn eines jeden Jahres der feste Betrag B eingezahlt wird bei einer Verzinsung von p Prozent auf den jeweils am Ende eines Jahres insgesamt vorliegenden Betrag, kann man mit Hilfe der geometrischen Summe berechnen, wieviel Geld man nach Ablauf des Sparplanes erhalten wird.

Man kann die Situation so auffassen, dass der jeweils zu Jahresbeginn eingezahlte Betrag auf einem eigenen Konto geführt wird und zum Auszahlungszeitpunkt die Summe aller dieser n Kontostände ausgezahlt wird. Der zu Beginn des ersten Jahres eingezahlte Betrag wird am Ende des Jahres verzinst, d.h. mit dem Faktor $q := 1 + \frac{p}{100}$ multipliziert. Am Ende des zweiten Jahres wird dieser neue Betrag $B \cdot q$ wieder mit q multipliziert. Nach n Jahren ist der zu Beginn des ersten Jahres eingezahlte Betrag damit auf $B \cdot q^n$ angewachsen. Die zu Beginn der folgenden Jahre eingezahlten Beträge wachsen entsprechend zu $B \cdot q^{(n-1)}, \dots, B \cdot q$ an. Als Summe aller dieser Beträge wird daher bei Ablauf des Sparplanes nach n Jahren der folgende Betrag ausgezahlt:

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

PRODUKTZEICHEN

Das Produktzeichen definiert man wie das Summenzeichen:

$$\prod_{k=m}^n x_k := x_m \cdot x_{m+1} \cdot \dots \cdot x_{n-1} \cdot x_n.$$

Falls $n < m$, definiert man

$$\prod_{k=m}^n x_k := 1.$$

Die Rechenregeln sind ähnlich wie oben, wir gehen darauf nicht ein.

FAKULTÄT

Wir definieren n Fakultät durch

$$n! := 1 \cdot 2 \cdot \dots \cdot n = \prod_{k=1}^n k.$$

Man setzt $0! := 1$. Wir zeigen weiter unten, dass man $n!$ interpretieren kann als die Anzahl der Möglichkeiten, n Objekte linear anzuordnen (=Anzahl der Permutationen von n Elementen).

BINOMIALKOEFFIZIENTEN

Für $n, k \in \mathbb{N}$ und $k > 0$ definieren wir

$$\binom{n}{k} := \frac{n(n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}.$$

Das hat im Zähler wie im Nenner ein Produkt von k Faktoren. Man setzt

$$\binom{n}{0} := 1,$$

insbesondere also $\binom{0}{0} = 1$. Offenbar gilt $\binom{n}{n} = 1$ und $\binom{n}{k} = 0$, falls $k > n$. Für $k \leq n$ hat man

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

und wir rechnen nach, dass

$$\binom{n}{k-1} + \binom{n}{k} = \binom{n+1}{k} \quad \text{für } 1 \leq k \leq n. \quad (2)$$

Es ist nämlich

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

Die Formel (2) ist eine „Rekursionformel“: Wenn man alle $\binom{n}{k}$ für ein gewisses n schon hat, so kann man daraus sehr einfach die $\binom{n+1}{k}$ berechnen, nämlich gerade so, wie man es beim Pascalschen Dreieck tut:

$$\begin{array}{cccccccc} & & & & 1 & & & \\ & & & & & 1 & & \\ & & & 1 & & 2 & & 1 \\ & & 1 & & 3 & & 3 & & 1 \\ & 1 & & 4 & & 6 & & 4 & & 1 \\ 1 & & 5 & & 10 & & 10 & & 5 & & 1 \end{array}$$

INDUKTION

Eine Eigenschaft der natürlichen Zahlen von fundamentaler Wichtigkeit ist die folgende, nicht besonderes überraschende Feststellung: Wenn man bei 0 beginnend „immer eins weiter zählt“ erwischt man „schließlich“ jede natürliche Zahl. Wenn man also eine Aussage $A(n)$ hat, die für jede natürliche Zahl n Sinn macht, und wenn diese Aussage für 0 richtig ist und beim „Weiterzählen“ wahr bleibt, dann gilt sie für alle natürlichen Zahlen. Dieses ist

Satz 23 (Prinzip der vollständigen Induktion).

- Gilt $A(0)$ (Induktionsanfang, Verankerung),
- und gilt „ $A(n) \implies A(n+1)$ “ (Induktionsschluss),

so gilt $A(n)$ für alle $n \in \mathbb{N}$.

Ist vielleicht $A(0)$ falsch oder sinnlos, aber gilt $A(m)$ für ein gewisses $m > 0$, und gilt wieder die Implikation $A(n) \implies A(n+1)$, so folgt ebenso $A(n)$ für alle $n \geq m$.

Beispiel 24. Wir beweisen, dass die Anzahl der möglichen Reihenfolgen (= Anordnungen = Permutationen) von n Elementen $n!$ ist. Dabei wollen wir nur $n \geq 1$ betrachten. Die Aussage $A(n)$ heißt: „ n Elemente lassen sich auf genau $n!$ verschiedene Weisen (linear) anordnen.“

- Offenbar ist $A(1)$ richtig, weil sich ein einziges Element nur auf eine Weise anordnen lässt.
- Nun nehmen wir an,
 - $A(n)$ sei (für irgendein n) richtig (*Induktionsannahme, Induktionsvoraussetzung*).
 - Wir wollen zeigen, dass dann auch $A(n+1)$ richtig ist (*Induktionsschluss*).

Dazu lassen wir von den $n+1$ Elementen eines weg. Die verbleibenden n können wir dann nach unserer Induktionsannahme auf $n!$ verschiedene Weisen anordnen. In jede dieser Anordnungen kann das weggelassene Element aber an $n+1$ Stellen eingefügt werden (einschließlich am Anfang und am Ende). Also ist die gesuchte Anzahl möglicher Anordnungen von $n+1$ Elementen gerade $(n+1) \cdot n! = (n+1)!$.

□

Als weiteres Beispiel beweisen wir den

Satz 25 (Binomischer Satz). Es gilt für alle $n \in \mathbb{N}$ und $a, b \in \mathbb{R}$.

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

Zum Beispiel ist also

$$(a+b)^4 = a^4 + 4a^3b + 6a^2b^2 + 4ab^3 + b^4.$$

Wir beweisen den Satz durch vollständige Induktion:

- Der Fall $n=0$ ist wahr:

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

- $A(n) \implies A(n+1)$. Wir nehmen also an, dass

$$(a+b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k \quad (3)$$

für ein bestimmtes n gilt, und wollen zeigen, dass dann

$$(a+b)^{n+1} = \sum_{k=0}^{n+1} \binom{n+1}{k} a^{n+1-k} b^k. \quad (4)$$

Nun ist

$$\begin{aligned} (a+b)^{n+1} &= (a+b)(a+b)^n \\ &\stackrel{(3)}{=} (a+b) \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k \\ &= \sum_{k=0}^n \binom{n}{k} a^{n+1-k} b^k + \sum_{k=0}^n \binom{n}{k} a^{n-k} b^{k+1} \\ &\stackrel{(*)}{=} \sum_{k=0}^n \binom{n}{k} a^{n+1-k} b^k + \sum_{k=1}^{n+1} \binom{n}{k-1} a^{n-k+1} b^k \\ &= a^{n+1} + \sum_{k=1}^n \binom{n}{k} a^{n+1-k} b^k + \sum_{k=1}^n \binom{n}{k-1} a^{n-k+1} b^k + b^{n+1} \\ &= a^{n+1} + \sum_{k=1}^n \left(\binom{n}{k} + \binom{n}{k-1} \right) a^{n+1-k} b^k + b^{n+1} \\ &\stackrel{(2)}{=} a^{n+1} + \sum_{k=1}^n \binom{n+1}{k} a^{n+1-k} b^k + b^{n+1} \\ &= \sum_{k=0}^{n+1} \binom{n+1}{k} a^{n+1-k} b^k. \end{aligned}$$

Daraus folgt (4). Überlegen Sie sich, dass bei der mit (*) gekennzeichneten Gleichung die zweiten Summen auf beiden Seiten gleich sind, dass also die „Indexverschiebung“ den Wert der Summe nicht ändert.

Apropos „Beweisen“. Beweisen heißt erklären, warum etwas so funktioniert, wie man behauptet. Das scheint mir auch für Ingenieure ein wünschenswertes Unterfangen. Andererseits sind mathematische Beweise mitunter technisch so kompliziert, dass nur Könner die Erklärung verstehen. Solche Beweise werden wir in diesem Skriptum nicht bringen, sondern nur solche, die

- einen Sachverhalt auch für Sie verständlicher machen sollen oder
- der beispielhaften Einübung von wichtigen Rechentechniken dienen oder
- Ihnen neue Begriffe aus verschiedenen Blickwinkeln nahebringen sollen.

Der vorstehende Beweis ist von der zweiten Sorte: Wie geht man mit Summenzeichen und Binomialkoeffizienten um?

Beispiel 26 (Quantenzahlen).*Werkstoffe I, Abschnitt 1.5*

Durch vollständige Induktion beweist man (Sie?), dass die Summe der ersten n ungeraden Zahlen n^2 ist:

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

Diese Formel kommt etwa im Zusammenhang mit den Quantenzahlen eines Elektrons vor: Ein Elektron hat

- eine Hauptquantenzahl n ,
- eine Nebenquantenzahl l , welche die Werte $0, 1, \dots, n-1$ annehmen kann, und
- eine Magnetquantenzahl mit den möglichen Werten $-l, -(l-1), \dots, -1, 0, 1, \dots, l$, also mit $2l+1$ möglichen Werten.

Die Anzahl der möglichen Kombinationen von Neben- und Magnetquantenzahl (Anzahl der *Orbitals*) bei gegebenem n ist daher $\sum_{l=0}^{n-1} (2l+1) = n^2$.

Beispiel 27 (Kombinatorische Interpretation der Binomialkoeffizienten). Die binomische Formel liefert

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

Wenn man $(1+x)^n = (1+x) \cdot \dots \cdot (1+x)$ „von Hand“ ausmultipliziert, muss man alle n -fachen Produkte summieren, die aus jeder Klammer einen der beiden Summanden 1 oder x enthalten. Der Koeffizient $\binom{n}{k}$ von x^k gibt deshalb an, auf wieviel verschiedene Weisen man aus n Objekten (den Klammern) k auswählen kann (nämlich die, Klammern, aus denen man x nimmt).

$$\begin{aligned} \binom{n}{k} &= \text{Anzahl der Möglichkeiten, aus } n \text{ Elementen genau } k \text{ auszuwählen} \\ &= \text{Anzahl der } k\text{-elementigen Teilmengen einer } n\text{-elementigen Menge.} \end{aligned}$$

Die Anzahl *aller* Teilmengen einer n -elementigen Menge ist daher

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

□

2.2 Ganze, rationale und reelle Zahlen

- In diesem Abschnitt wiederholen wir vor allem wichtige Regeln für das Rechnen mit Ungleichungen.

GANZE ZAHLEN

In der Menge $\mathbb{N}$ der natürlichen Zahlen hat man die Operationen Addition und Multiplikation mit den üblichen Rechenregeln. Die Umkehroperationen Subtraktion und Division sind dagegen nur eingeschränkt möglich. In der Menge der ganzen Zahlen

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

kann man uneingeschränkt subtrahieren, aber noch nicht uneingeschränkt dividieren. Aus der Sicht der Analysis muss man das als Manko betrachten. Aber der Umstand, dass sich nicht jede ganze Zahl durch jede (von 0 verschiedene) andere ganze Zahl ohne Rest teilen lässt, hat interessante Folgen, denen man in der sogenannten *Zahlentheorie* nachgeht. Dazu gehört die Theorie der *Primzahlen*), die neuerdings in der Kodierungstheorie wichtige Anwendungen findet.

RATIONALE ZAHLEN

Die Menge der rationalen Zahlen ist die Menge aller „ausgekürzten“ Brüche ganzer Zahlen:

$$\mathbb{Q} = \left\{ \frac{a}{b} \mid a, b \in \mathbb{Z} \text{ teilerfremd, } b > 0 \right\},$$

Im Bereich $\mathbb{Q}$ der rationalen Zahlen kann man die Grundrechenarten Addition und Multiplikation mit ihren Umkehrungen Subtraktion und Division uneingeschränkt durchführen (nur Division durch 0 ist nicht erklärt), und es gelten die Ihnen aus der Schule vertrauten Rechenregeln (Assoziativ- und Kommutativgesetze, Distributivgesetz). Einen Zahlbereich, in dem eine Addition und eine Multiplikation mit diesen Regeln erklärt sind, nennt man einen *Körper*. Die Bereiche $\mathbb{N}$ und $\mathbb{Z}$ sind keine Körper, aber die reellen Zahlen oder die komplexen Zahlen liefern weitere Beispiele. In der Analysis sind wir in erster Linie an den reellen Zahlen interessiert. Bevor wir auf den Unterschied zwischen reellen und rationalen Zahlen eingehen, wollen wir aber auf das sehr wichtige Rechnen mit Ungleichungen eingehen, das man im Bereich der rationalen Zahlen ebenso gut beschreiben (und üben) kann, wie in dem der reellen Zahlen.

UNGLEICHUNGEN

Für die rationalen Zahlen hat man eine *Ordnungsrelation* $a < b$ „ a kleiner b “, so dass man für sie also Ungleichungen angeben kann. Das gilt ebenfalls für die reellen, nicht aber für die komplexen Zahlen.

Für $a < b$ schreibt man auch $b > a$ (b größer a). Und man schreibt $a \leq b$ für $a < b$ oder $a = b$.

Für die Ungleichungen gibt es die folgenden grundlegenden Rechenregeln (Axiome), aus denen sich alle anderen herleiten lassen:

$$\text{Es gilt genau einer der drei Fälle } x < y, x = y, y < x, \quad (5)$$

$$x < y \text{ und } y < z \implies x < z \quad (6)$$

$$x < y \text{ und } a \leq b \implies x + a < y + b, \quad (7)$$

$$x < y \text{ und } 0 < a \implies ax < ay. \quad (8)$$

Beispiel 28. Wir zeigen, wie die (natürlich vertraute) Tatsache dass $-x > 0$, falls $x < 0$, sich aus diesen Rechenregeln ableiten lässt. Die Zahl $-x$ ist offenbar $\neq 0$. Also ist $-x > 0$ oder $-x < 0$. Wäre $-x < 0$, so hätte man

$$x < 0 \text{ und } -x < 0,$$

also nach (7)

$$0 = x + (-x) < 0 + 0 = 0.$$

Aber $0 < 0$ ist ein Widerspruch! Daher war die Annahme $x < 0$ falsch, und es verbleibt als einzige Möglichkeit $x > 0$.

Das hier benutzte Beweisschema nennt man „Beweis durch Widerspruch“ oder „indirekten Beweis“. Man nimmt an, die gesuchte Aussage sei nicht wahr, und führt diese Annahme zum Widerspruch.

□

Beispiel 29. Ist $x > 0$, so folgt nach (8) auch

$$x^2 = x \cdot x > x \cdot 0 = 0.$$

Ist $x < 0$, so ist also $-x > 0$ und wie eben folgt

$$x^2 = (-x) \cdot (-x) > (-x) \cdot 0 = 0.$$

Damit gilt für alle $x \in \mathbb{Q}$ (oder $\in \mathbb{R}$)

$$x \neq 0 \implies x^2 > 0.$$

Das Quadrat einer rationalen (oder reellen) Zahl ist nie negativ.

□

Beispiel 30 (Arithmetisches und geometrisches Mittel). Für $0 < a$ und $0 < b$ gilt

$$0 \leq (a - b)^2 = a^2 + b^2 - 2ab.$$

Wir addieren dazu auf beiden Seiten $4ab$ und erhalten

$$4ab \leq a^2 + b^2 + 2ab = (a + b)^2.$$

Schließlich multiplizieren wir beide Seiten mit der positiven Zahl $\frac{1}{4}$ und erhalten

$$ab \leq \left(\frac{a+b}{2}\right)^2.$$

Beim Wurzelziehen bleiben Ungleichungen zwischen positiven Zahlen erhalten (warum?), und wir bekommen

$$\sqrt{ab} \leq \frac{a+b}{2}.$$

Man sagt, das *geometrische Mittel* von a und b ist höchstens so groß, wie das *arithmetische Mittel*.

Dazu ist allerdings noch anzumerken,

- dass $\sqrt{x^2} = x$ nur stimmt, wenn $x \geq 0$ ist, andernfalls ist $\sqrt{x^2} = |x|$, vgl. unten. Aber weil a und b positiv vorausgesetzt waren, ist ja $\frac{a+b}{2} > 0$.
- Außerdem ist die Operation des Quadratwurzelziehens im Bereich der rationalen Zahlen gar nicht uneingeschränkt möglich, so wenig wie die Division a/b im Bereich der ganzen Zahlen uneingeschränkt möglich ist. $\sqrt{4/49} = 2/7$ „geht“ innerhalb von $\mathbb{Q}$, aber $\sqrt{2}$ geht nicht: Es gibt keine rationale Zahl, deren Quadrat 2 ist. Das ist ein Grund dafür warum wir in der Analysis die reellen Zahlen brauchen und mit den rationalen nicht auskommen.

□

ABSOLUTBETRAG.

Mittels $<$ kann man den Absolutbetrag einer rationalen (oder reellen) Zahl x definieren:

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

Also gilt immer

$$-|x| \leq x \leq |x|$$

und

$$|x| = 0 \iff x = 0.$$

Offensichtlich ist $|xy| = |x||y|$. Weniger offensichtlich, aber sehr wichtig ist die sogenannte

Dreiecksungleichung:

$$|x + y| \leq |x| + |y|.$$

Zum Beweis betrachten wir zunächst den Fall, dass $x + y \geq 0$. Dann hat man

$$|x + y| = x + y \leq |x| + |y|.$$

Im anderen Falle ist $x + y < 0$ und

$$|x + y| = -(x + y) = (-x) + (-y) \leq |-x| + |-y| = |x| + |y|.$$

Die Dreiecksungleichung gilt auch für Summen mit mehr als zwei Summanden.

$$\left| \sum_{k=0}^n x_k \right| \leq \sum_{k=0}^n |x_k|.$$

Beispiel 31. Wendet man die Dreiecksungleichung auf x und $y - x$ statt x und y an, so ergibt sich $|y| \leq |x| + |y - x|$ oder $|y| - |x| \leq |y - x|$. Da sich die Rollen von x und y vertauschen lassen, folgt auch $|x| - |y| \leq |x - y| = |y - x|$. Zusammen ergibt sich

$$||y| - |x|| \leq |y - x|.$$

□

ZAHLDARSTELLUNGEN

Wir sind es gewohnt, Zahlen im Dezimalsystem darzustellen: Die Zahl

$$2 \cdot 10^3 + 7 \cdot 10 + 5 = \mathbf{2} \cdot 10^3 + \mathbf{0} \cdot 10^2 + \mathbf{7} \cdot 10^1 + \mathbf{5} \cdot 10^0 \quad (9)$$

schreiben wir durch Angabe der Ziffern als 2075. Dabei benötigt man zur Erfassung des Bereichs zwischen den Zehnerpotenzen gerade 10 Ziffern, nämlich $0, \dots, 9$. Auch Kombinationen von negativen Zehnerpotenzen lassen sich so erfassen, indem man die Stelle zwischen 10^0 und 10^{-1} durch ein Komma markiert:

$$3 \cdot 10^0 + 5 \cdot 10^{-1} + 6 \cdot 10^{-2} = 3,56.$$

Statt der Basis 10 lässt sich ebenso gut jede andere natürliche Zahl $m > 1$ verwenden, und die Fälle $m = 2, 8, 16$ spielen in der Computerarithmetik eine große Rolle. Man braucht für diese sogenannte m -adische Darstellung dann m Ziffern. Für $m \leq 10$ nimmt man einfach die Ziffern $0, \dots, m-1$, für $m = 16$ verwendet man $0, 1, \dots, 9, A, B, C, D, E, F$. Wenn man die verwendete Basis andeuten will, so schreibt man sie gern als unteren Index:

Beispiel 32. Im World Wide Web (WWW) werden neben reinen Texten auch viele Informationen mittels farbiger Bilder vermittelt und im Browser erscheinen farbige Texte vor farbigem Hintergrund. Vielleicht haben Sie sich schon einmal gefragt, wie die verschiedenen Farben der einzelnen Pixel Ihres Monitors dargestellt werden. Ein übliches Farbmodell gibt für jeden Pixel seinen **R**ot-, **G**rün- und **B**lauanteil an (RGB-Modell). Die Mischung dieser Anteile ergibt dann die tatsächlich erscheinende Farbe. Für die Darstellung wird jedoch nur eine beschränkte Anzahl von Farbwerten zur Verfügung gestellt. Der Rotanteil kann z.B. nur Werte von 0 (kein Rotanteil) bis 255 (kräftiges Rot) annehmen. Diese einzelnen Farbwerte werden im Hexadezimalsystem angegeben, wobei 00 keinem Rotanteil und FF dem kräftigen Rot entspricht. Sie können z.B. im Quelltext einer Webseite die Angabe `BGColor="#BADOE6"` finden, die die Hintergrundfarbe mischt aus $BA_{16} = 11 \cdot 16^1 + 10 \cdot 16^0 = 186_{10}$ Teilen Rot, $D0_{16} = 13 \cdot 16^1 + 0 \cdot 16^0 = 208_{10}$ Teilen Grün und $E6_{16} = 14 \cdot 16^1 + 6 \cdot 16^0 = 230_{10}$ Teilen Blau.

Die Kombination der möglichen $256 = 2^8$ Werte der drei Farbanteile ergibt $256^3 = 2^{3 \cdot 8} = 2^{24}$ verschiedene Farben. Da man zur Darstellung der Farbwerte jeden Pixels also 24 Bit Speicher zur Verfügung stellen muss, spricht man auch von 24 Bit Farbtiefe.

REELLE ZAHLEN

Zurück zum Dezimalsystem. Bei der Darstellung kommt man keineswegs immer wie in (9) mit endlich vielen Dezimalen aus. Es kann sich durchaus ein unendlicher Dezimalbruch ergeben, wie Sie das von $\frac{1}{3} = 0,33\dots$ kennen. Das führt dann also zu einer unendlichen Reihe

$$\pm \sum_{k=-m}^{\infty} c_k 10^{-k} \quad \text{mit } c_k \in \{0, \dots, 9\}, m \in \mathbb{Z}.$$

Die Frage, was die Summe von unendlich vielen Zahlen ist, müssen wir später genauer untersuchen (Konvergenz von Reihen im Abschnitt 10).

Man kann beweisen (nicht sehr schwer), dass für *rationales* $x \in \mathbb{Q}$ die Ziffernfolge $c_1, c_2, \dots$ schließlich immer periodisch wird, falls, wenn man abbrechende Folgen als „0-periodisch“ interpretiert:

$$0,53 = 0,530000\dots$$

Umgekehrt liefern periodische Dezimalbrüche immer rationale Zahlen. Das legt es nahe, die rationalen Zahlen zu erweitern, indem man auch die nicht-periodischen Dezimalbrüche zulässt. Man erhält dann einen größeren Zahlbereich, die *reellen Zahlen*:

$$\mathbb{R} = \{x \mid x \text{ ist unendlicher Dezimalbruch}\}.$$

Wir wollen allerdings Dezimalbrüche mit Periode 9 ausschließen, weil $0.\bar{9} = 1.0$, wie wir später noch genauer begründen wollen, vgl. Abschnitt 10.1.

Die reellen Zahlen $\mathbb{R}$ bilden wie die rationalen einen Körper: die Grundrechenarten funktionieren wie gewohnt. Auch das Rechnen mit Ungleichungen bleibt unverändert gültig. Der wesentliche Vorteil der reellen Zahlen vor den rationalen ist es, dass in ihnen *Grenzprozesse* erfolgreich sind, die in den rationalen Zahlen scheitern, weil man gewissermaßen in einem Loch zwischen den Zahlen landet. Dadurch werden in den reellen Zahlen Rechenoperationen möglich und sinnvoll, die in den rationalen Zahlen nicht möglich sind. So wissen wir, dass $y = x^2 \geq 0$. Aber ist auch jede positive Zahl y ein Quadrat? Kann man aus ihr die Wurzel

ziehen? Wir versuchen es, indem wir zunächst eine beliebige Zahl $x_0 > 0$ raten und testen, ob $x_0^2 = y$. Offenbar ist immer

$$y = x_0 \frac{y}{x_0}.$$

Gilt $x_0^2 = y$, so ist $\frac{y}{x_0} = x_0$ und beide Faktoren sind gleich. Haben wir aber x_0 , also den ersten Faktor, zu groß gewählt, so ist der zweite, also $\frac{y}{x_0}$ zu klein, und wir erwarten, dass der Mittelwert von beiden, also

$$x_1 := \frac{1}{2} \left(x_0 + \frac{y}{x_0} \right),$$

ein besserer Kandidat für die Wurzel aus y ist. Wir setzen das Verfahren fort. Durch die Vorschrift

$$x_{k+1} = \frac{1}{2} \left(x_k + \frac{y}{x_k} \right)$$

erhalten wir eine *rekursiv definierte* Folge $x_0, x_1, x_2, \dots$ von Zahlen, deren Quadrat gegen y konvergiert. Aber es kann eben sein, und zum Beispiel ist es für $y = 2$ der Fall, dass die x_k gegen einen nicht-periodischen Dezimalbruch konvergieren. In der Menge der rationalen Zahlen hat also die Gleichung $x^2 = 2$ keine Lösung. In den reellen Zahlen aber sehr wohl. Der Grenzwert $\sqrt{y} = \lim_{k \rightarrow \infty} x_k$ existiert für jedes $y > 0$. Ähnlich ist es mit vielen anderen Grenzprozessen, z.B. beim Differenzieren und Integrieren. Sie sehen aus, als wären sie konvergent, aber erst in den reellen Zahlen lässt sich das beweisen - weil es hier erst wahr ist. Die rationalen Zahlen sind für Grenzprozesse zu löcherig.

Beachten Sie, dass wir hier von Konvergenz reden, obwohl wir diesen Begriff noch nicht richtig eingeführt haben, das kommt im Abschnitt 3.1 Hier sollte nur der Grund für die Einführung der reellen Zahlen erläutert werden.

2.3 Komplexe Zahlen

- Wir lernen die Rechenregeln für die komplexen Zahlen und was wir uns darunter vorstellen sollen.
- Diese Zahlen sind ein ganz wichtiges Hilfsmittel der Analysis. Das wird deutlich bei der Untersuchung eines einfachen Schwingungsvorganges.

Komplexe Zahlen sind „Zahlen“ der Form $z = x + iy$ mit reellem x und y . Die Grundrechenarten für komplexe Zahlen sind die für reelle Zahlen, nur soll

$$i^2 = -1$$

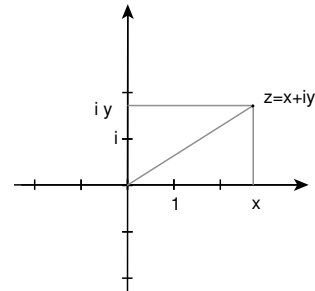
gelten¹. Also etwa

$$\frac{3-i}{5+2i} = \frac{(3-i)(5-2i)}{(5+2i)(5-2i)} = \frac{15-6i-5i+2i^2}{25-4i^2} = \frac{13}{29} + i\frac{-11}{29}. \quad (10)$$

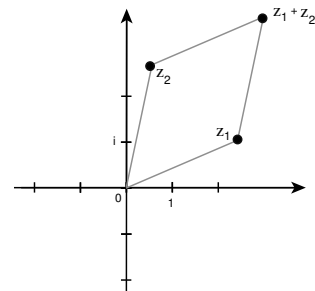
Der Trick mit der Erweiterung ist wichtig, weil man so das Ergebnis wieder in der Form $x + iy$ bekommt.

Die Zahlen $x + 0i$ sollen die alten reellen Zahlen sein, und die Zahlen $0 + iy$ nennt man *imaginäre* Zahlen. Damit sollte es klar sein, wie man mit komplexen Zahlen rechnet.

Die Gaußsche Zahlenebene. Man stellt sich die reellen Zahlen gern als Punkte auf der Zahlengeraden vor. Nach Carl Friedrich Gauß (1777-1855) kann man sich entsprechend die komplexen Zahlen als die Punkte einer Ebene vorstellen, die die reelle Zahlengerade enthält. Die „Zahl“ $x + iy$ ist der Punkt mit den Koordinaten (x, y) . Man muss dann nur erklären, wie man Punkte in der Ebene addiert und multipliziert (so wie man das früher für Punkte auf der Zahlengeraden tun musste).



Die Addition ist die aus der Schule bekannte(?) Vektoraddition. Auf die Multiplikation gehen wir weiter unten ein.



¹In der Elektrotechnik verwendet man anstelle von i häufig das Symbol j , um nicht mit der Standardbezeichnung für den Strom in Konflikt zu geraten.

Charakteristische Größen komplexer Zahlen. Für die komplexe Zahl $z = x + iy$ mit $x, y \in \mathbb{R}$ definieren wir:

$\operatorname{Re} z := x$	Realteil
$\operatorname{Im} z := y$	Imaginärteil, ist selbst reell!
$\bar{z} := x - iy$	Konjugierte
$ z := \sqrt{x^2 + y^2}$	(Absolut)betrag

Beispiel 33. Sei $z = 3 + 4i$. Dann ist

$$\operatorname{Re} z = 3, \operatorname{Im} z = 4, \bar{z} = 3 - 4i, |z| = 5.$$

□

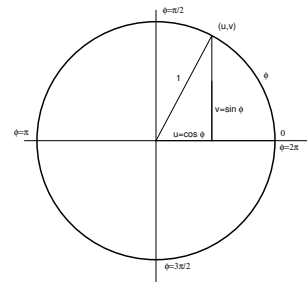
Rechenregeln für die Konjugation	Rechenregeln für den Betrag
$\bar{\bar{z}} = z$	$ z = \sqrt{z\bar{z}}$
$z\bar{z} = x^2 + y^2$	$ z_1 z_2 = z_1 z_2 $
$\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$	$\left \frac{z_1}{z_2} \right = \frac{ z_1 }{ z_2 }$
$\overline{\bar{z}_1 \bar{z}_2} = z_1 z_2$	$ z_1 + z_2 \leq z_1 + z_2 $ (Dreiecksungleichung)
$\overline{1/z} = 1/\bar{z}$	
$\operatorname{Re}(z) = \frac{1}{2}(z + \bar{z}), \quad \operatorname{Im}(z) = \frac{1}{2i}(z - \bar{z})$	

Komplexe Zahlen in Polarkoordinaten. Wir geben nun eine Darstellung der komplexen Zahlen, die besonders nützlich ist für die Multiplikation.

Die Zahl $z = x + iy \neq 0$ mit $x^2 + y^2 \neq 0$ lässt sich auch schreiben als

$$z = \sqrt{x^2 + y^2} \left(\underbrace{\frac{x}{\sqrt{x^2 + y^2}}}_u + i \underbrace{\frac{y}{\sqrt{x^2 + y^2}}}_v \right).$$

Wegen $u^2 + v^2 = 1$ ist dann (u, v) ein Punkt auf dem Einheitskreis. Daher gibt es eine – bis auf Vielfache von 2π eindeutig bestimmte – Zahl ϕ , so dass $u = \cos \phi$ und $v = \sin \phi$. Die formelmäßige Berechnung von ϕ ist nicht ganz einfach, sie hängt davon ab, wo $z = x + iy$ liegt.



Zum Beispiel kann man wählen:

$$\phi = \begin{cases} \arctan \frac{y}{x}, & \text{falls } x > 0, \\ \pi + \arctan \frac{y}{x}, & \text{falls } x < 0, \\ \pm \frac{\pi}{2}, & \text{falls } x = 0, \pm y > 0. \end{cases}$$

Aber auch $\phi = \arccos \frac{x}{\sqrt{x^2 + y^2}}$ wäre möglich (in welchem Bereich?)²

²Zu Ihrer Beruhigung: Mit den arc-Funktionen werden wir uns später noch genauer beschäftigen, vgl. Abschnitt 6.2.

Setzt man noch $r = |z|$, so erhält man

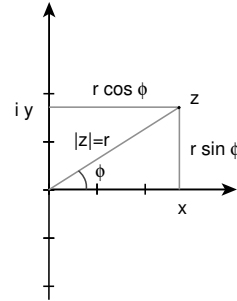
$$z = r(\cos \phi + i \sin \phi).$$

Die Zahl ϕ (im Bogenmaß gemessen!) heißt *das Argument* von z .

$r = |z|$ hatten wir schon als *Betrag* von z definiert. Man nennt diese Darstellung von z auch die Darstellung in *Polarkoordinaten* oder – etwas altertümelnd – die „Zeigerdarstellung“. Für $z = 0$ ist ϕ nicht bestimmt, aber natürlich kann man dann schreiben

$$z = 0(\cos \phi + i \sin \phi)$$

mit einem willkürlich gewählten ϕ .



Beispiel 34.

$$\begin{aligned} 1 &= 1(\cos 0 + i \sin 0), & -1 &= 1(\cos \pi + i \sin \pi), \\ i &= 1(\cos \frac{\pi}{2} + i \sin \frac{\pi}{2}), & 1+i &= \sqrt{2}(\cos \frac{\pi}{4} + i \sin \frac{\pi}{4}). \end{aligned}$$

□

Multiplikation in Polarkoordinaten. Für komplexe Zahlen $z_1 = r_1(\cos \phi_1 + i \sin \phi_1)$ und $z_2 = r_2(\cos \phi_2 + i \sin \phi_2)$ ist

$$\begin{aligned} z_1 z_2 &= r_1 r_2 ((\cos \phi_1 \cos \phi_2 - \sin \phi_1 \sin \phi_2) + i(\cos \phi_1 \sin \phi_2 + \sin \phi_1 \cos \phi_2)) \\ &= r_1 r_2 (\cos(\phi_1 + \phi_2) + i \sin(\phi_1 + \phi_2)). \end{aligned}$$

Bei der Multiplikation multiplizieren sich die Beträge
und addieren sich die Argumente.

Bei der Division dividieren sich die Beträge, und die Argumente subtrahieren sich.

Beispiel 35. Gesucht sind alle (reellen und komplexen) Lösungen der Gleichung

$$z^n = 1. \tag{11}$$

Setzen wir $z = r(\cos \phi + i \sin \phi)$, so erhalten wir nach der obigen Multiplikationsregel

$$z^n = r^n (\cos(n\phi) + i \sin(n\phi)) = 1(\cos 0 + i \sin 0).$$

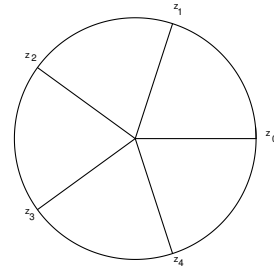
Daraus folgt für die nichtnegative reelle Zahl r , dass $r = 1$ und weiter

$$n\phi = 0 + 2k\pi \text{ für eine ganze Zahl } k \in \mathbb{Z}.$$

Es kommen also alle Werte $\phi = \frac{2k\pi}{n}$ in Frage. Weil aber $\frac{2(k+n)\pi}{n} = \frac{2k\pi}{n} + 2\pi$ als Argument dieselbe komplexe Zahl liefert, wie das Argument $\phi = \frac{2k\pi}{n}$, braucht man nur die Werte $k = 0, \dots, n-1$ zu betrachten, um alle verschiedenen Lösungen zu erhalten.

Die Gleichung (11) hat n verschiedene Lösungen, d.h. es gibt n verschiedene n -te Wurzeln aus 1, nämlich

$$z_k = \cos \frac{2k\pi}{n} + i \sin \frac{2k\pi}{n}, \quad k = 0, \dots, n-1.$$



□

Beispiel 36. Berechnen Sie $\sqrt{i}$. Genauer soll das heißen: Lösen Sie die Gleichung

$$z^2 = i.$$

Wie im vorigen Beispiel findet man unter Berücksichtigung von $i = \cos \frac{\pi}{2} + i \sin \frac{\pi}{2}$ die Werte

$$z_1 = \cos \frac{\pi}{4} + i \sin \frac{\pi}{4} = \frac{1}{2}\sqrt{2}(1+i), \quad z_2 = \cos\left(\frac{\pi}{4} + \pi\right) + i \sin\left(\frac{\pi}{4} + \pi\right) = -z_1.$$

Anders gesagt:

$$\sqrt{i} = \pm \frac{1+i}{\sqrt{2}}.$$

Wo liegen diese Wurzeln in der Gaußschen Zahlenebene?

□

WOZU BRAUCHT MAN KOMPLEXE ZAHLEN?

In diesem Abschnitt greifen wir ein wenig auf die Differentialrechnung vor. Zum Beispiel wissen Sie sicher, dass die Ableitung der Exponentialfunktion e^t wieder die Exponentialfunktion ist. Und mit der Kettenregel folgt daher

$$\frac{d}{dt} e^{\lambda t} = \lambda e^{\lambda t}.$$

Wachstums- oder Zerfallsvorgänge befolgen (wenigstens in gewissen Entwicklungsstadien) die sogenannte *Wachstumsgleichung*

$$\frac{dx}{dt} = cx \quad \text{oder} \quad \dot{x} = cx,$$

die besagt, dass die Wachstumsgeschwindigkeit proportional zur vorhandenen Menge ist. Das ist eine Gleichung für eine gesuchte Funktion $x(t)$, die deren Ableitung involviert, eine sogenannte *Differentialgleichung*. Und wir wissen schon, wie wir eine Lösung finden können: $x(t) = e^{ct}$ ist eine. Können Sie andere angeben, zum Beispiel eine mit dem „Anfangswert“ $x(0) = 5$?

Nun sei $x(t)$ die Auslenkung einer an einer Feder befestigten Masse aus der Ruhelage. Dann ist die Trägheitskraft gegeben durch

$$\text{Masse} \times \text{Beschleunigung} = m\ddot{x},$$

während die Rückstellkraft der Feder gerade $-kx(t)$ ist (k = Federkonstante). Das Newtonsche Gesetz liefert die *Schwingungsgleichung*

$$m\ddot{x} + kx = 0. \quad (12)$$

Vgl. [Müller: Mechanik II, Abschnitt 15]. Das ist wieder eine Differentialgleichung, und wenn wir sie

in der Form

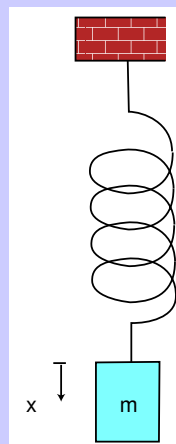
$$\ddot{x} = -\frac{k}{m}x =: cx$$

schreiben, haben wir sofort eine Idee, wie wir sie lösen können:

$$x(t) = e^{\sqrt{c}t} \text{ oder auch } x(t) = e^{-\sqrt{c}t}$$

sollte eine Lösung sein, denn dann ist $\ddot{x}(t) = (\sqrt{c})^2 e^{\sqrt{c}t} = cx(t)$. Aber dummerweise ist c in diesem Fall negativ! Wenn etwa $c = -1$ ist, erhalten wir $x(t) = e^{\pm it}$, und das sind anscheinend komplexe Funktionen mit:

$$x(0) = e^{\pm i0} = e^0 = 1, \quad \dot{x}(0) = \pm ie^0 = \pm i.$$



Aber macht das Sinn? Gibt es diese Funktionen überhaupt?

Bleiben wir der Einfachheit halber bei $c = -1$, also bei der Differentialgleichung

$$\ddot{x} = -x. \quad (13)$$

Wenn wir uns vergegenwärtigen, dass sie die Bewegung eines Federpendels beschreibt, fallen uns sofort ganz andere Lösungen ein, nämlich $x(t) = \sin t$ oder $x(t) = \cos t$. Die Rechenregeln der Differentiation zeigen, dass auch sogenannte Linearkombinationen

$$x(t) = a \cos t + b \sin t$$

für beliebige Konstanten a und b Lösungen sind, und zwar mit

$$x(0) = a \text{ und } \dot{x}(0) = b.$$

Aus physikalischen Gründen bestimmen aber Ort und Geschwindigkeit zur Zeit $t = 0$ die weitere Bewegung, d.h. die Lösung $x(t)$, eindeutig. Und wenn es so eine Funktion e^{it} wie oben gibt, muss gelten

$$e^{it} = \cos t + i \sin t. \quad (14)$$

Wir benutzen diese sogenannte *Eulersche Relation* nun einfach als Definition für e^{it} .

Definition 37 (Komplexe Exponentialfunktion, Eulerrelation). Für $x, y \in \mathbb{R}$ definieren wir

$$e^{x+iy} = e^x(\cos y + i \sin y).$$

Was ist zum Beispiel $e^{i\pi}$? Klar:

$$e^{i\pi} = \cos \pi + i \sin \pi = -1$$

und

$$e^{i\frac{\pi}{2}} = \cos \frac{\pi}{2} + i \sin \frac{\pi}{2} = i.$$

So haben wir das einfach definiert. Und die Funktionen

$$e^{\pm it} = \cos(\pm t) + i \sin(\pm t) = \cos t \pm i \sin t$$

sind wirklich (komplexe!) Lösungen von (13). Das gilt dann auch für alle Linearkombinationen

$$x(t) = ae^{it} + be^{-it}.$$

Zum Beispiel bekommt man die Lösungen

$$\frac{1}{2}e^{it} + \frac{1}{2}e^{-it} = \cos t \quad \text{und} \quad \frac{1}{2i}e^{it} - \frac{1}{2i}e^{-it} = \sin t, \quad (15)$$

aber nun drehen wir uns im Kreis.

Bisher sieht es aus, als hätten wir einfache Sachverhalte nur immer komplizierter gemacht. Der Nutzen dieser Überlegungen wird erst deutlich, wenn wir das Federpendel mit Reibung betrachten. Die Reibungskraft hängt typischerweise linear von der Geschwindigkeit ab, und die Bewegung wird beschrieben durch die Differentialgleichung

$$m\ddot{x} + a\dot{x} + kx = 0. \quad (16)$$

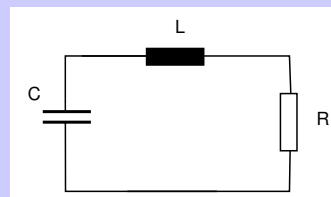
Wie soll man diese Differentialgleichung lösen? Das ist eine wichtige Frage, denn Differentialgleichungen von diesem Typ kommen nicht nur beim Federpendel, sondern an sehr vielen Stellen vor, vergleichen Sie die Einleitung zu Kapitel 6.

Beispiel 38.

In der Elektrotechnik wird der Strom I in einem RLC-Schwingkreis beschrieben durch die Differentialgleichung

$$L\ddot{I} + R\dot{I} + \frac{1}{C}I = 0,$$

also durch eine Differentialgleichung von ganz demselben Typ wie (16). Nur die Buchstaben sind andere. Weitere Beispiele im Anhang Abschnitt 11.2.



Weil wir bisher mit der Exponentialfunktion einigermaßen erfolgreich waren, machen wir für die Lösung (versuchsweise) den ANSATZ

$$x(t) = e^{\lambda t}.$$

Wir setzen das in (16) ein und erhalten wegen $\dot{x} = \lambda e^{\lambda t}$ und $\ddot{x} = \lambda^2 e^{\lambda t}$:

$$(m\lambda^2 + a\lambda + k)e^{\lambda t} = 0.$$

Wenn also

$$\lambda = \frac{1}{2m} \left(-a \pm \sqrt{a^2 - 4mk} \right)$$

eine Lösung der quadratischen Gleichung

$$m\lambda^2 + a\lambda + k = 0$$

ist, ist $x(t) = e^{\lambda t}$ eine Lösung der Differentialgleichung. Die quadratische Gleichung hat im allgemeinen zwei Lösungen λ_1 und λ_2 , und wieder sind alle Funktionen

$$x(t) = ae^{\lambda_1 t} + be^{\lambda_2 t}$$

ebenfalls Lösung. Und das sind alle, wie man zeigen kann!

- Wenn die Reibung a groß ist, sind beide Nullstellen reell und negativ, das Federpendel kriecht in die Nullstellung zurück (wie beim automatischen Türschließer).
- Wenn die Reibung klein ist, sind $\lambda_{1,2} = \alpha \pm i\omega$ komplexe Nullstellen mit negativem Realteil $\alpha = -\frac{a}{2m}$, und man findet Lösungen

$$e^{\lambda_{1,2}t} = e^{(\alpha \pm i\omega)t} = e^{\alpha t} e^{\pm i\omega t} = e^{\alpha t} (\cos \omega t \pm i \sin \omega t).$$

Daraus bekommt man wie in (15) reelle Lösungen

$$x(t) = e^{\alpha t} \cos \omega t \text{ und } x(t) = e^{\alpha t} \sin \omega t.$$

Wegen $\alpha < 0$ sind das gedämpfte Schwingungen.

Sie haben jetzt etwas sehr Wichtiges gelernt:

- Sie können mit dem Exponentialansatz Differentialgleichungen vom Typ(16) (sogenannte homogene lineare Differentialgleichungen mit konstanten Koeffizienten) lösen.
- Und Sie haben gelernt, dass komplexe Zahlen dabei unentbehrlich sind, weil die Nullstellen der auftretenden quadratischen Gleichung komplex sein können. Das Verfahren liefert dann erst einmal komplexe Lösungen mit der komplexen Exponentialfunktion. Daraus kann man reelle Lösungen gewinnen. Aber wenn das Komplexe erst einmal seine Schrecken verloren hat, werden Sie aus Bequemlichkeit ganz darauf verzichten und sich daran gewöhnen, im Komplexen zu arbeiten.

Bemerkungen. Wir sind nicht darauf eingegangen, was bei doppelten Nullstellen $\lambda_1 = \lambda_2$ passiert. Auch ist offensichtlich, dass der Exponentialansatz nicht klappt, wenn auf der rechten Seiten der Differentialgleichung nicht 0 sondern eine andere Konstante oder sogar eine Funktion von t steht, eine sogenannte Zwangskraft oder Inhomogenität.³ Noch schwieriger wird es, wenn die Koeffizienten m, a und k auch von der Zeit abhängen oder wenn die Differentialgleichung nicht linear ist, wenn sie zum Beispiel das Quadrat x^2 der gesuchten Funktion enthält. Diese Fälle werden in den Modulen *Differentialgleichungen für Ingenieure* bzw. *Integraltransformationen und partielle Differentialgleichungen für Ingenieure* behandelt.

³Vergleichen Sie dazu den Anhang 11.4, in dem wir einen systematischen Überblick über lineare Differentialgleichungen mit konstanten Koeffizienten geben.

3 Grenzwerte und Stetigkeit

Das Lösen von Gleichungen führt oft auf das Problem, die Nullstellen von Funktionen zu bestimmen. Bei quadratischen Polynomen ist das für die meisten von Ihnen kein Problem. Bei komplizierteren Funktionen helfen oft aber nur numerische Verfahren. Ein sehr naives, aber gar nicht so schlechtes Verfahren dafür ist das folgende

INTERVALLHALBIERUNGSVERFAHREN

Wir suchen eine Nullstelle der Funktion f und nehmen an, dass wir einen Punkt x_0 mit $f(x_0) < 0$ und einen Punkt y_0 mit $f(y_0) > 0$ kennen. Wir vermuten dazwischen eine Nullstelle und hoffen, dass der Mittelwert $\frac{x_0+y_0}{2}$ dichter an der Nullstelle liegt.

Wir berechnen $f(\frac{x_0+y_0}{2})$.

Wenn das 0 ist, sind wir fertig.

Wenn $f(\frac{x_0+y_0}{2}) < 0$ ist, setzen wir

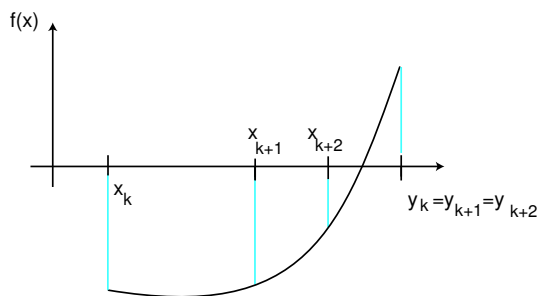
$$x_1 := \frac{x_0 + y_0}{2},$$

$$y_1 := y_0.$$

Und wenn $f(\frac{x_0+y_0}{2}) > 0$ ist, setzen wir

$$x_1 := x_0,$$

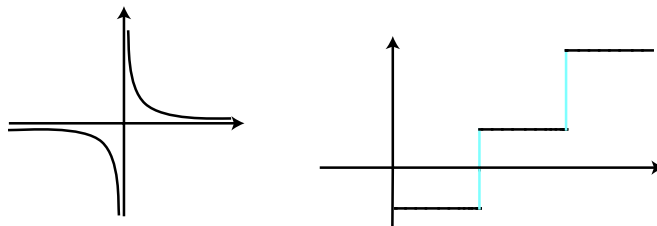
$$y_1 := \frac{x_0 + y_0}{2}.$$



Dann ist also $f(x_1) < 0$ und $f(y_1) > 0$ und die beiden Punkte x_1 und y_1 liegen doppelt so nah beisammen wie x_0 und y_0 . Wir setzen dieses Verfahren „unendlich“ fort und erwarten, dass jede der beiden Folgen $x_0, x_1, x_2, \dots$ und $y_0, y_1, y_2, \dots$ gegen „die Nullstelle x^* “ strebt. Was das bedeuten soll, müssen wir allerdings genauer erklären.

Was heißt das: Eine Folge strebt gegen einen Grenzwert x^* ?

Vielleicht gibt es auch andere Probleme. Betrachten Sie die beiden folgenden Funktionen:



Offenbar müssen wir sagen, wie eine Funktion beschaffen sein soll, damit unser Verfahren klappt. Irgendwie darf sie „keine Löcher“ haben. Das zu präzisieren ist gar nicht einfach.⁴ Entscheidend ist der Begriff der *Stetigkeit von Funktionen*. Um ihn zu erklären braucht man ebenfalls den Grenzwertbegriff.

Darum geht es in den nächsten Abschnitten.

⁴Die Mathematik hat dafür fast zweihundert Jahre gebraucht.

3.1 Zahlenfolgen und Konvergenz

- Was sind Folgen und wann sind sie konvergent?
- Wir lernen eine Menge Beispiele kennen.

Eine Folge reeller Zahlen $x_0, x_1, \dots$ ist mathematisch gesprochen eine Abbildungen der natürlichen Zahlen $\mathbb{N}$ in die reellen Zahlen: Jedem $n \in \mathbb{N}$ wird eine reelle Zahl x_n zugeordnet. Beliebte Notation: (x_n) oder $(x_n)_{n \in \mathbb{N}}$. Manchmal beginnt man die Nummerierung der Folge auch mit 1 statt mit 0, gelegentlich auch mit einer anderen Zahl.

Beispiel 39.

$$\left(\frac{1}{2^n}\right)_{n \in \mathbb{N}} : 1, \frac{1}{2}, \frac{1}{4}, \frac{1}{8}, \dots \quad (17)$$

$$(x_n = 1)_{n \in \mathbb{N}} : 1, 1, \dots \quad (18)$$

$$\left(\frac{1}{n}\right)_{n > 0} : 1, \frac{1}{2}, \frac{1}{3}, \dots \quad (19)$$

$$((-1)^{n+1}n^2)_{n \in \mathbb{N}} : 0, 1, -4, 9, -25, \dots \quad (20)$$

□

Beispiel 40 (Eulerfolge). Mit *Eulerfolge* wollen wir die Folge $(x_n)_{n \geq 1}$ mit den Gliedern

$$x_n := \left(1 + \frac{1}{n}\right)^n$$

bezeichnen. Ihre ersten Glieder sind

$$x_1 = 2,0000000$$

$$x_2 = 2,2500000$$

$$x_3 = 2,3703704$$

$$x_4 = 2,4414062$$

$$x_5 = 2,4883200$$

...

□

Häufig, zum Beispiel beim oben beschriebenen Halbierungsverfahren, kommen sogenannte *rekursive* Folgen vor: Man gibt einen (oder mehrere) Anfangswerte und eine Vorschrift, wie sich die Folgenglieder aus den vorangehenden Gliedern „entwickeln“. Dem folgenden Beispiel sind wir früher schon begegnet.

Beispiel 41 (Wurzelfolge).

$$x_0 := 1, \quad x_{n+1} := \frac{1}{2} \left(x_n + \frac{2}{x_n} \right) \quad (21)$$

liefert

$$x_1 = 1,5000000$$

$$x_2 = 1,4166667$$

$$x_3 = 1,4142157$$

$$x_4 = 1,4142136$$

$$x_5 = 1,4142136$$

...

□

Beispiel 42 (Fibonacci-Folge). Die Fibonacci-Folge ist gegeben durch

$$a_0 = a_1 = 1, \quad a_{n+2} = a_n + a_{n+1}. \quad (22)$$

Das ergibt 1, 1, 2, 3, 5, 8, 13, 21, ... Diese Folge hat zu tun mit der Vermehrung von unsterblichen Karnickeln. $\square$

Beispiel 43 (Collatz-Folge). Eine weitere rekursiv definierte Folge ist die Collatz-Folge: Man beginnt mit einer beliebigen *natürlichen* Zahl $c_0 > 0$ und definiert

$$c_{n+1} := \begin{cases} 3c_n + 1 & \text{falls } c_n \text{ ungerade und } \neq 1, \\ \frac{c_n}{2} & \text{falls } c_n \text{ gerade,} \\ 1 & \text{falls } c_n = 1. \end{cases}$$

Also etwa

$$9, 28, 14, 7, 22, 11, 34, 17, 52, 26, 13, 40, 20, 10, 5, 16, 8, 4, 2, 1, 1, 1, \dots$$

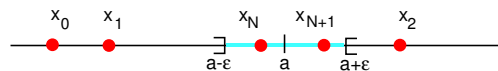
$\square$

Definition 44 (Konvergenz, Divergenz). Die Folge (x_n) heißt *konvergent gegen* $a \in \mathbb{R}$, wenn es zu jedem $\epsilon > 0$ eine natürliche Zahl N gibt, so dass

$$|x_n - a| < \epsilon \text{ für alle } n \geq N.$$

Man schreibt dann $x_n \rightarrow a$ oder $\lim_{n \rightarrow \infty} x_n = a$. Die Zahl a heißt der *Grenzwert* oder *Limes* der Folge. Man nennt die Folge (x_n) *konvergent*, wenn es ein a gibt, gegen das sie konvergiert. Andernfalls nennt man sie *divergent*. Sprachlich kann man das so formulieren:

Eine Folge ist konvergent gegen den Grenzwert a , wenn jede – noch so kleine! – Toleranz ϵ für die Abweichung vom Grenzwert a nur von *endlich vielen* Folgengliedern (nämlich höchstens denen mit $n < N$) verletzt wird. Vom Index N an liegen alle Folgenglieder innerhalb der Toleranzschränken.



Beispiel 45 (Ein Konvergenzbeweis: DAS fundamentale Beispiel). Die Folge $(\frac{1}{n})_{n>0}$ konvergiert gegen 0:

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

Beweis. Sei $\epsilon > 0$ (beliebig). Wir müssen zeigen: Es gibt eine natürliche Zahl N , so dass für alle $n \geq N$

$$|\frac{1}{n} - 0| < \epsilon. \quad (23)$$

Das bedeutet aber $\frac{1}{n} < \epsilon$ oder, wenn man die Reziproken nimmt, $n > \frac{1}{\epsilon}$. Also kann man irgendein $N > \frac{1}{\epsilon}$ wählen. Dann gilt für alle $n \geq N$ ebenfalls $n > \frac{1}{\epsilon}$ und daher (23). $\square$

Beispiel 46 (Divergenz: DAS fundamentale Beispiel). Die Folge $((-1)^n)_{n \in \mathbb{N}}$ ist divergent. Ihre Glieder sind nämlich abwechselnd $+1$ und -1 , d.h. aufeinander folgende Glieder haben den Abstand 2 voneinander. Wir beweisen die Divergenz nun indirekt: Wäre die Folge konvergent etwa gegen eine Zahl a , so gäbe es zu jedem $\epsilon > 0$ eine Stelle N , von der ab alle Folgenglieder näher als ϵ an a lägen. „Zu jedem $\epsilon > 0$ “, also auch zu $\epsilon = 1$. Wenn sich zwei

aufeinander folgende Glieder von a aber um weniger als 1 unterscheiden, dann unterscheiden sie sich voneinander um *weniger* als 2. Das ist ein Widerspruch: Sie unterscheiden sich nämlich um 2.

Wir schreiben den Beweis noch einmal mathematischer auf:

Beweis (indirekt). Annahme: Die Folge ist gegen ein $a \in \mathbb{R}$ konvergent. Dann gibt es insbesondere zu $\epsilon = 1$ ein N , so dass

$$|x_n - a| < 1 \quad \text{für alle } n \geq N.$$

Dann wären aber insbesondere $|x_N - a| < 1$ und $|x_{N+1} - a| < 1$. Daraus folgt mit der Dreiecksungleichung

$$|x_N - x_{N+1}| = |(x_N - a) - (x_{N+1} - a)| \leq |x_N - a| + |x_{N+1} - a| < 1 + 1 = 2.$$

Andrerseits ist $x_N - x_{N+1} = \pm(1 - (-1)) = \pm 2$, und wir erhalten einen Widerspruch. Also war die Annahme der Konvergenz falsch, und die Folge ist divergent. $\square$

Mit demselben oder einem ähnlichen Argument können Sie beweisen, dass

- die Folge $0, 1, 2, 3, 4, \dots$ divergent ist.
- allgemeiner eine *ganzzahlige* Folge dann und nur dann konvergent ist, wenn von einem gewissen Index N an alle Folgenglieder gleich sind.
- eine Folge nicht gegen *zwei verschiedene* Grenzwerte konvergieren kann.

Wichtige Bemerkungen. 1. Es ist hilfreich sich klarzumachen, dass aufgrund der Definition der „Anfang“ einer Folge keinen Einfluß auf die Frage ihrer Konvergenz hat, alles Wichtige passiert „bei den großen n “. Genauer gesagt ist die Folge $(x_n)_{n \geq 0}$ genau dann konvergent, wenn für ein N die Folge $(x_n)_{n \geq N}$ konvergent ist:

Ende gut, alles gut!

Als Konsequenz kann man Konvergenz nicht durch Berechnung von endlich vielen (wenn auch sehr vielen) Folgengliedern mit dem Rechner nachweisen. Wenn man Aussagen über die Konvergenz von Verfahren machen will, muss man etwas BEWEISEN, der Augenschein genügt nicht, vgl. Anhang Abschnitt 11.3. Das ist der Grund, warum in diesem Zusammenhang immer wieder Beweise auftreten. . .

2. Der Begriff der Konvergenz unendlicher Folgen ist nicht so einfach, und eine falsche sprachliche Formulierung erschwert sein Verständnis oder dokumentiert fehlende gedankliche Bewältigung. Hier einige (öfter anzutreffende) Beispiele falscher Behauptungen, mit denen Sie jederzeit Ihre Ignoranz beweisen können⁵:

Die Folge (x_n) konvergiert gegen a , wenn

... sie a immer näher kommt. FALSCH: $\lim_{n \rightarrow \infty} (1 + \frac{1}{n}) \neq 0$.

... sie a beliebig nah kommt. FALSCH: $\lim_{n \rightarrow \infty} ((-1)^n + \frac{1}{n}) \neq 1$.

... sie a beliebig nah kommt, es aber nie erreicht. Erst recht FALSCH.

⁵Ein langjähriger Personalchef eines großen Essener Konzerns pflegte bei Bewerbungsgesprächen von Ingenieuren nach der Definition der Konvergenz zu fragen . . .

Der Grenzwert der Folge $(\frac{1}{n})$ *geht* nicht gegen 0, er *ist* 0. Grenzwerte sind bereits angekommen...

Bemerkung: Komplexe Folgen. Bei der Definition von Zahlenfolgen und ihrer Konvergenz haben wir Folgen *reeller* Zahlen betrachtet. Ebenso gut kann man Folgen *komplexer* Zahlen betrachten und dieselben Definitionen wie oben benutzen. Zum Beispiel ist $(i^k)_{k \in \mathbb{N}}$ die Folge

$$1, i, -1, -i, 1, i, -1, -i, \dots$$

und die Folge $\left(\frac{1}{(2+3i)^k}\right)_{k \in \mathbb{N}}$ ist konvergent gegen 0, weil

$$\left| \frac{1}{(2+3i)^k} - 0 \right| = \frac{1}{\sqrt{4+9}^k} \rightarrow 0.$$

3.2 Konvergenzbeweise

- Wie erkennt man, ob eine Folge konvergent ist, und wie berechnet man den Grenzwert?

Es gibt kein einfaches Rezept, um die Konvergenz von Folgen zu beweisen. Ganz wesentlich ist allerdings, dass man mit Ungleichungen umgehen kann.

Beispiel 47. Für $x > 1$ gilt

$$\lim_{n \rightarrow \infty} \sqrt[n]{x} = 1.$$

Beweis. Sei $\epsilon > 0$. Wir müssen zeigen: Es gibt ein $N \in \mathbb{N}$ mit

$$|\sqrt[n]{x} - 1| < \epsilon \quad \text{für alle } n \geq N.$$

Weil $\sqrt[n]{x} > 1$, müssen wir nur zeigen, dass

$$y_n := \sqrt[n]{x} - 1 < \epsilon \quad \text{für alle } n \geq N. \quad (24)$$

Dazu verwenden wir einen Trick: Weil $1 + y_n = \sqrt[n]{x}$, ist

$$x = (1 + y_n)^n = 1 + ny_n + \binom{n}{2} y_n^2 + \dots + y_n^n \geq ny_n,$$

denn alle Summanden rechts sind nicht-negativ. Daraus folgt $y_n \leq \frac{x}{n}$ und daraus die Ungleichung (24) für alle $n \geq \frac{x}{\epsilon}$. $\square$

Beispiel 48 (Und nun Sie...). Der Trick im letzten Beweis ist wichtig. Um ihn zu lernen, beweisen Sie, dass

$$\lim_{n \rightarrow \infty} \sqrt[n]{n} = 1.$$

Benutzen Sie dabei nicht den zweiten, sondern den dritten Term der Binomialentwicklung.

Noch einfacher können Sie mit demselben Trick beweisen, dass für $a = 1 + b > 1$ gilt

$$\lim_{k \rightarrow \infty} \frac{k}{a^k} = 0.$$

$\square$

Beispiel 49 (Schwierige Fälle). Die Euler- und die Wurzelfolge aus den Beispielen 40 und 41 konvergieren gegen die Eulersche Zahl e bzw. gegen $\sqrt{2}$, aber das können wir noch nicht beweisen. $\square$

Beispiel 50 (Ein ungelöster Fall). Die Collatz-Folge ist konvergent gegen 1 für jedes $c_0 \in \mathbb{N} \setminus \{0\}$, für das man sie getestet hat. D.h. sie landet irgendwann bei 1 und bleibt dann natürlich dort. Aber bis heute ist unbekannt, ob das wirklich für jedes $c_0 \in \mathbb{N}$ der Fall ist. $\square$

Der Nachweis der Konvergenz oder Divergenz einer Folge ist oft eine anspruchsvolle Knobelei. Gelegentlich können einem dabei Rechenregeln für Grenzwerte helfen: Die Frage nach der Konvergenz zusammengesetzter Folgen lässt sich damit auf die nach der (vielleicht schon bekannten) Konvergenz der Komponenten reduzieren.

Satz 51 (Rechenregeln für konvergente Folgen). Seien (x_n) und (y_n) konvergente Folgen mit

$$\lim_{n \rightarrow \infty} x_n = a, \quad \lim_{n \rightarrow \infty} y_n = b,$$

und sei $c \in \mathbb{R}$. Dann gilt

$$\begin{aligned} \lim_{n \rightarrow \infty} (x_n + y_n) &= a + b \\ \lim_{n \rightarrow \infty} (x_n y_n) &= ab \\ \lim_{n \rightarrow \infty} c x_n &= ca \\ \lim_{n \rightarrow \infty} \frac{x_n}{y_n} &= \frac{a}{b}, \quad \text{falls } b \neq 0 \neq y_n \text{ für alle } n, \\ a \leq b, \quad &\text{falls } x_n \leq y_n \text{ für alle } n. \end{aligned}$$

Aber echte Ungleichungen $x_n < y_n$ übertragen sich nicht unbedingt auf die Grenzwerte!

Beispiel 52 (Ein Standard-Trick). Aus

$$\frac{1+n}{2n+3} = \frac{\frac{1}{n} + 1}{2 + \frac{3}{n}}$$

folgt mit den Rechenregeln und $\frac{1}{n} \rightarrow 0$, dass

$$\lim_{n \rightarrow \infty} \frac{1+n}{2n+3} = \frac{1}{2} \quad \text{und} \quad \lim_{n \rightarrow \infty} \left(\frac{1+n}{2n+3} \right)^2 = \frac{1}{4}.$$

□

Beispiel 53. Wir zeigen $\lim_{n \rightarrow \infty} \frac{x^n}{n!} = 0$. Wir können uns auf den Fall $x > 0$ beschränken. (Warum?) Wir wählen ein $k \in \mathbb{N}$ mit $k > 2x$. Dann gilt für alle $n \geq k$:

$$\frac{x^n}{n!} = \frac{x^k}{k!} \cdot \frac{x^{n-k}}{(k+1) \dots n} \leq \frac{x^k}{k!} \cdot \frac{x^{n-k}}{(2x)^{n-k}} = \frac{x^k}{k!} \cdot 2^k \left(\frac{1}{2} \right)^n \rightarrow 0 \quad \text{für } n \rightarrow \infty.$$

□

Definition 54 (Bestimmte Divergenz). Man sagt, die Folge (x_n) ist *bestimmt divergent gegen* $+\infty$, wenn es zu jedem (noch so großen) $\epsilon > 0$ ein $N \in \mathbb{N}$ gibt, so dass

$$x_n \geq \epsilon \text{ für alle } n \geq N.$$

Man schreibt dann auch

$$\lim_{n \rightarrow \infty} x_n = +\infty$$

und nennt $+\infty$ den *uneigentlichen Grenzwert* der Folge, obwohl $+\infty$ ja keine Zahl ist!

Der Begriff *bestimmte Divergenz* dient der Abgrenzung gegen die *unbestimmte Divergenz* zum Beispiel der Folge $((-1)^n)$.

Beachten Sie: Manchmal sagt man statt *bestimmt divergent gegen* $+\infty$ auch (*uneigentlich*) *konvergent gegen* $+\infty$. Wir bleiben aber bei der Sprachregelung

konvergent = konvergent gegen eine reelle Zahl.

Beispiel 55. Zeigen Sie wieder mit Hilfe der Binomialentwicklung, dass

$$\lim_{n \rightarrow \infty} x^n = +\infty, \text{ falls } x > 1,$$

indem Sie $x = 1 + y$ mit $y > 0$ benutzen. □

Ein schwieriges Problem ist es, die Konvergenz einer Folge zu beweisen, wenn man nicht weiß, was ihr Grenzwert ist. Relativ oft kann einem das folgende Kriterium helfen, dessen Beweis aber die Möglichkeiten dieser Vorlesung übersteigt.

Satz 56 (Monotoniekriterium). *Jede beschränkte monotone Folge reeller Zahlen ist konvergent.*

Dabei nennen wir eine Folge (x_n) *beschränkt*, wenn es Zahlen $a, b \in \mathbb{R}$ gibt, so dass

$$a \leq x_n \leq b \text{ für alle } n.$$

Wir nennen (x_n) *monoton wachsend*, wenn

$$x_n \leq x_{n+1} \text{ für alle } n.$$

Steht hier sogar $<$, so heißt die Folge *streng* monoton wachsend. Analog definiert man *monoton fallend*. Und eine Folge ist monoton, wenn sie monoton wachsend oder fallend ist. (Gibt es Folgen, die beides sind?)

Das Kriterium ist nur hinreichend, aber nicht notwendig: Zwar ist jede konvergente Folge beschränkt, aber nicht unbedingt monoton.

Beispiel 57 (Die Wurzelfolge). Die Wurzel-Folge (21) ist nicht monoton, wie man aus den berechneten Werten sieht. Allerdings stört scheint's nur das erste Glied. Für die Konvergenz ist das unerheblich, und wir ignorieren es deshalb.

Zunächst sind alle $x_n > 0$.

Weil das arithmetische Mittel größer gleich dem geometrischen ist, ist

$$x_{n+1} = \frac{x_n + \frac{2}{x_n}}{2} \geq \sqrt{x_n \frac{2}{x_n}} = \sqrt{2}.$$

und deshalb

$$x_{n+1} - x_n = \frac{x_n + \frac{2}{x_n}}{2} - x_n = \frac{1}{2x_n}(2 - x_n^2) \leq 0$$

Damit ist bewiesen, dass die Folge (21) nach Weglassen des ersten Gliedes monoton fallend und positiv, also beschränkt ist. Sie ist deshalb konvergent. Zur Bestimmung des Grenzwertes benutzen wir einen **Trick, der oft bei rekursiv definierten Folgen wirkt**, wenn man schon weiß, dass sie konvergent sind: Wenn $\lim_{n \rightarrow \infty} x_n = a$, dann ist auch der Grenzwert der in der Nummerierung um eins verschobenen Folge derselbe: $\lim_{n \rightarrow \infty} x_{n+1} = a$. Es folgt

$$\underbrace{x_{n+1}}_{\rightarrow a} = \frac{1}{2} \left(\underbrace{x_n}_{\rightarrow a} + \underbrace{\frac{2}{x_n}}_{\rightarrow 2/a} \right)$$

Es folgt $a = \frac{1}{2}(a + \frac{2}{a})$ und daraus $a^2 = 2$. Weil aber alle $x_n > 0$, ist $a = +\sqrt{2}$. □

Beispiel 58 (Die Eulerfolge). Nach dem binomischen Satz ist

$$\begin{aligned} x_n &= \left(1 + \frac{1}{n}\right)^n = 1 + \binom{n}{1} \frac{1}{n} + \sum_{k=2}^n \binom{n}{k} \frac{1}{n^k} \\ &= 1 + 1 + \sum_{k=2}^n \frac{n(n-1) \dots (n-k+1)}{k!} \frac{1}{n^k} \\ &= 1 + 1 + \sum_{k=2}^n \frac{1}{k!} \left(1 - \frac{1}{n}\right) \dots \left(1 - \frac{k-1}{n}\right). \end{aligned}$$

Beim Übergang von n zu $n+1$ werden alle Faktoren $(1 - \frac{j}{n})$ größer, und es kommt ein positiver Summand hinzu. Also ist die Folge monoton wachsend. Ebenso sieht man, dass

$$x_n \leq 1 + 1 + \sum_{k=2}^n \frac{1}{k!} \leq 1 + \sum_{k=1}^n \frac{1}{2^{k-1}} = 1 + \sum_{k=0}^{n-1} \left(\frac{1}{2}\right)^k = 1 + \frac{1 - (\frac{1}{2})^n}{1 - \frac{1}{2}} < 1 + 2 = 3.$$

Also ist die Eulerfolge auch beschränkt und konvergiert gegen eine Zahl $e \leq 3$, die man als Eulersche Zahl bezeichnet. Es ist

$$e = 2.71 \dots$$

□

3.3 Stetigkeit von Funktionen

- Den Konvergenzbegriff für Folgen benutzen wir, um einen “kontinuierlichen” Konvergenzbegriff für Funktionen zu erklären.
- Der Begriff der *Stetigkeit* von Funktionen ist wichtig für die erfolgreiche Suche nach Nullstellen und der Lösung von Gleichungen.
- Auch für Frage nach Extremwerten ist er von großer Bedeutung.

Sei $f : \mathbb{R} \supset D \rightarrow \mathbb{R}$ eine Funktion auf dem Definitionsbereich D und sei $x^* \in \mathbb{R}$.

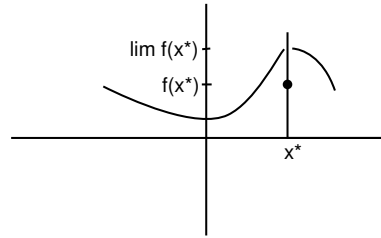
Definition 59. Wir sagen, dass f für x gegen x^* den Grenzwert y^* hat, wenn für jede Folge $(x_k)_{k \in \mathbb{N}}$ von Punkten in $D \setminus \{x^*\}$, die gegen x^* konvergiert, der Limes $\lim_{k \rightarrow \infty} f(x_k)$ existiert und gleich y^* ist.

Notation:

$$\lim_{x \rightarrow x^*} f(x) = y^*.$$

Dabei muss x^* selbst nicht unbedingt im Definitionsbereich D von f liegen; und die x_k sollen alle $\neq x^*$ sein.⁶

In der offensichtlichen Weise erweitern wir diese Definition auf den Fall, dass x^* und/oder y^* eine der „Nicht-Zahlen“ $\pm\infty$ ist.



In der Definition wird die kontinuierliche Bewegung von x gegen x^* ersetzt durch diskrete Folgen, die gegen x^* gehen. So wird der Grenzwertbegriff für Funktionen auf den für Folgen zurückgespielt. Allerdings genügt es nicht, *eine* Folge mit Grenzwert x^* zu betrachten, *alle* solche Folgen müssen zum selben Resultat führen.

Beispiel 60.

$$\lim_{x \rightarrow 0} x^2 = 0,$$

denn für jede Folge (x_n) mit $\lim_{n \rightarrow \infty} x_n = 0$ ist $\lim_{n \rightarrow \infty} x_n^2 = 0$.

Ebenso

$$\lim_{x \rightarrow 0} \frac{1+2x}{5+3x} = \frac{1}{5}, \quad \lim_{x \rightarrow +\infty} \frac{1+2x}{5+3x} = \frac{2}{3},$$

denn für jede Folge ...

□

Oft kommen auch *einseitige* Grenzwerte vor:

Definition 61. Man sagt y^* ist der *linksseitige Grenzwert* von f in x^* und schreibt

$$\lim_{x \rightarrow x^* -} f(x) = y^* \text{ oder } \lim_{x \nearrow x^*} f(x) = y^*,$$

wenn $\lim_{k \rightarrow \infty} f(x_k) = y^*$ für jede Folge (x_k) in $D \setminus \{x^*\}$, die gegen x^* konvergiert und bei der $x_k < x^*$ für alle k gilt.⁷

⁶Wir wollen stets voraussetzen, dass es wenigstens eine Folge in $D \setminus \{x^*\}$ gibt, die gegen x^* konvergiert, d.h. dass x^* kein isolierter Punkt des Definitionsbereiches ist.

Entsprechend definiert man den *rechtsseitigen Grenzwert*

$$\lim_{x \rightarrow x^*+} f(x) = \lim_{x \searrow x^*} f(x) = y^*.$$

Existieren $\lim_{x \searrow x^*} f(x)$ und $\lim_{x \nearrow x^*} f(x)$ und sind diese gleich, so existiert auch $\lim_{x \rightarrow x^*} f(x)$.

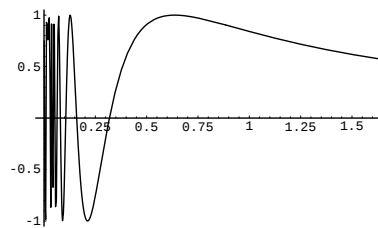
Beispiel 62.

$$\lim_{x \nearrow 0} \frac{x}{|x|} = -1, \quad \lim_{x \searrow 0} \frac{x}{|x|} = +1, \quad \lim_{x \searrow 0} \ln x = -\infty.$$

□

Beispiel 63.

$$\lim_{x \searrow 0} \sin \frac{1}{x} \quad \text{existiert nicht:}$$



Die Folge $\left(x_k = \frac{1}{k\pi + \frac{\pi}{2}}\right)_{k \in \mathbb{N}}$ geht von oben gegen null, aber die Folge mit den Gliedern

$$\sin \frac{1}{x_k} = \sin\left(k\pi + \frac{\pi}{2}\right) = (-1)^k$$

ist divergent.⁸

□

Definition 64 (Stetigkeit). Die Funktion f heißt *stetig* in x^* , wenn $x^* \in D$ und

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

Die Funktion heißt stetig (auf D), wenn sie in allen $x^* \in D$ stetig ist.

Beispiel 65 (Standard-Beispiele). Die Funktionen

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad x \mapsto x^2$$

oder

$$f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}, \quad x \mapsto \frac{1}{x}$$

sind stetig auf ihrem Definitionsbereich. (Im zweiten Fall reden wir nicht über 0, denn 0 liegt nicht im Definitionsbereich!)

Die Funktion f mit

$$f(x) := \begin{cases} +1 & \text{für } x \geq 0 \\ 0 & \text{für } x < 0 \end{cases}$$

⁷Wieder verlangt man überdies, dass es wenigstens eine solche Folge wirklich gibt.

⁸Das Bild des Graphen von $\sin \frac{1}{x}$ wurde mit dem Plot-Befehl von Mathematica erzeugt. Das Programm wertet die Funktion an bestimmten Stellen aus, die es selbst wählt. Die Stellen liegen enger zusammen, wenn die Funktionswerte heftig schwanken, aber das Programm erwischt im allgemeinen natürlich nicht automatisch die Maximal- und Minimalstellen. Deshalb sehen die Spitzen des Graphen so „angenibbelt“ aus. Um ein besseres Bild zu bekommen, muss man Mathematica genauere Anweisungen geben: Mathematische Software braucht menschliche Zuwendung und Fachkenntnis!

ist in 0 nicht stetig. Sie wird auch nicht stetig, wenn man ihr in 0 irgend einen anderen Wert verpasst. Gleiches gilt für die Funktion

$$f(x) := \begin{cases} \sin \frac{1}{x} & \text{für } x \neq 0 \\ 0 & \text{für } x = 0, \end{cases}$$

womit klar wird, dass Unstetigkeitsstellen von Funktionen nicht unbedingt „Sprungstellen“ sein müssen. $\square$

Rechenregeln für stetige Funktionen. Aus den Rechenregeln für Folgen ergeben sich unmittelbar einige Regeln für stetige Funktionen:

Summen, Differenzen und Produkte stetiger Funktionen sind stetig. Quotienten stetiger Funktionen sind überall dort stetig, wo der Nenner $\neq 0$ und damit die Quotientenfunktion definiert ist. Die Komposition von stetigen Funktionen ist stetig.

Beispiel 66. Reelle Polynome $a_n x^n + a_{n-1} x^{n-1} + \dots + a_0$ sind auf ganz $\mathbb{R}$ stetig. Rationale Funktionen (Quotienten von zwei Polynomen) sind stetig auf ihrem natürlichen Definitionsbereich, nämlich dort, wo der Nenner keine Nullstellen hat. Manchmal lassen sie sich allerdings auch in den Nullstellen des Nenners stetig ergänzen, zum Beispiel die Funktion $f(x) = \frac{x^2-9}{x-3}$. $\square$

Beispiel 67 (Wurzelfunktionen). Die Wurzelfunktionen

$$[0, +\infty[\rightarrow [0, +\infty[, \quad x \mapsto \sqrt[k]{x}$$

sind stetig. Auf den Beweis verzichten wir. $\square$

Komplexe Funktionen. Die vorstehenden Definitionen kann man unmittelbar auf komplexe Funktionen $f : \mathbb{C} \supset D \rightarrow \mathbb{C}$ übertragen, allerdings mit folgender Einschränkung: Im Komplexen machen Grenzwerte mit $\pm\infty$ und einseitige Grenzwerte keinen Sinn.

Zum Beispiel sind komplexe Polynome $a_n z^n + a_{n-1} z^{n-1} + \dots + a_0$ auf ganz $\mathbb{C}$ stetig. Rationale Funktionen (Quotienten von zwei Polynomen) sind überall dort stetig, wo der Nenner keine Nullstelle hat, vgl. oben.

BESTIMMUNG VON NULLSTELLEN.

Wir kommen nun zurück auf das Problem vom Beginn des Abschnitts 3, auf das Intervallhalbierungsverfahren zur Nullstellenbestimmung. Jetzt können wir dafür die hinreichenden Voraussetzungen formulieren.

Satz 68 (Zwischenwertsatz). Sei $f : I \rightarrow \mathbb{R}$ stetig auf einem Intervall $I \subset \mathbb{R}$. Seien $a, b \in I$ und

$$y^* \text{ eine Zahl zwischen } f(a) \text{ und } f(b).$$

Dann ist y^* ein Funktionswert von f :

$$\text{Es gibt (mindestens) ein } x^* \text{ zwischen } a \text{ und } b \text{ mit } f(x^*) = y^*.$$

Ein solches x^* läßt sich mit der Intervallhalbierungsmethode finden.

Beweis. Für $y^* = 0$ ist dies das Nullstellenproblem, und wir beweisen den Satz nur für diesen Fall. Der Fall mit beliebigem y^* ergibt sich dann sehr leicht, weil die y^* -Stellen von $f(x)$ gerade die Nullstellen von $\tilde{f}(x) = f(x) - y^*$ sind.

Wir nehmen an, dass $f(a) < 0 < f(b)$ und konstruieren mit der Intervallhalbierungsmethode zwei Folgen

$$\begin{aligned}x_0 &= a \leq x_1 \leq x_2 \leq \dots \leq b \\ y_0 &= b \geq y_1 \geq y_2 \geq \dots \geq a\end{aligned}$$

wie zu Beginn von Abschnitt 3. Diese Folgen sind monoton steigend bzw. fallend und beschränkt, also konvergent gegen ein x_* bzw. ein y_* . Wegen der Halbierung der Intervalle ist

$$y_k - x_k = \left(\frac{1}{2}\right)^k (b - a) \rightarrow 0.$$

Daraus folgt $x_* = y_*$. Weil nach unserer Konstruktion für alle $k \in \mathbb{N}$

$$f(x_k) \leq 0 \leq f(y_k),$$

folgt aus der Stetigkeit von f (hier wird sie gebraucht!), dass

$$f(x_*) = \lim_{k \rightarrow \infty} f(x_k) \leq 0 \leq \lim_{k \rightarrow \infty} f(y_k) = f(y_*),$$

also $f(x_*) = 0$. □

Beispiel 69. Ein reelles Polynom mit ungeradem Grad n

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

mit $a_n > 0$ geht für $x \rightarrow +\infty$ gegen $+\infty$ und für $x \rightarrow -\infty$ gegen $-\infty$. Insbesondere nimmt es positive und negative Werte an. Weil es stetig ist, hat es deshalb (irgendwo) eine Nullstelle. Mit dem Intervall-Halbierungsverfahren kann man für konkret gegebene Polynome eine solche ermitteln. Im Gegensatz dazu haben Polynome geraden Grades wie $x^2 + 1$ eventuell keine reelle Nullstelle. □

Beispiel 70. Die Funktion $\frac{1}{x}$ ist stetig auf ihrem ganzen Definitionsbereich, nämlich $\mathbb{R} \setminus \{0\}$. Sie nimmt dort positive und negative Werte an, hat aber keine Nullstelle! □

EXISTENZ VON EXTREMWERTEN.

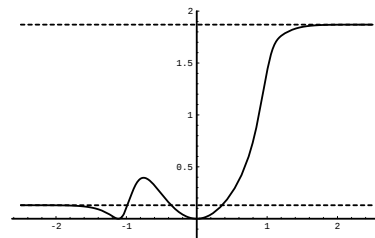
Ein anderes häufig auftretendes Problem neben der Bestimmung von Nullstellen ist die Ermittlung des Maximums (oder Minimums) einer reellwertigen Funktion. Zunächst müssen wir dazu allerdings die Begriffe klären: Das *Maximum* einer Funktion f ist der größte Wert, den diese annimmt, das *Minimum* entsprechend der kleinste. (Notation: $\max f, \min f$.)

Betrachten Sie folgenden Graphen einer Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$

Wir wollen verschiedene anschauliche Tatsachen daraus ableiten.⁹

Offenbar nimmt f gar keinen größten Wert an, sondern für $x \rightarrow +\infty$ wachsen die Werte monoton gegen 1.87. Also hat f kein Maximum!

Der Wert 1.87 ist aber etwas ähnliches wie ein Maximum, man nennt ihn *das Supremum* von f . Wird das Maximum angenommen, so ist das Supremum das Maximum.



⁹Mit der später eingeführten hyperbolischen Tangensfunktion kann man f schreiben als

$$f(x) = \tanh(x^2) + a \tanh(x^7), \quad a = 0.87,$$

und die gemachten Aussagen auch rechnerisch beweisen, aber im Augenblick verlassen wir uns auf unsere Augen.

Entsprechend definiert man das *Infimum* einer Funktion. Die oben dargestellte Funktion hat das Minimum 0, das also gleichzeitig das Infimum ist. Es wird an zwei Stellen angenommen, bei $x = 0$ und bei $x = -1.16$. Und dann ist da, etwa bei $x = -0.76$, noch ein *lokales Maximum*, das aber sicher nicht das (globale) Maximum ist (was es ja gar nicht gibt).

Wir fassen das in einer Definition zusammen:

Definition 71 (Maximum und Supremum). Sei $f : \mathbb{R} \supset D \rightarrow \mathbb{R}$ eine Funktion, $x^* \in D$ und $y^* \in \mathbb{R} \cup \{+\infty\}$.

- (i) f nimmt in x^* sein *Maximum* y^* an, wenn

$$f(x^*) = y^* \geq f(x) \text{ für alle } x \in D.$$

Notation: $y^* = \max_{x \in D} f(x) = \max f$.

Gilt die Gleichheit nur für $x = x^*$, so spricht man auch von einem *strengen* oder *eigentlichen* Maximum.

- (ii) y^* ist das *Supremum* von f , geschrieben $\sup f = \sup_{x \in D} f(x) = y^*$, wenn gilt

- (i) $f(x) \leq y^*$ für alle $x \in D$

- (ii) Es gibt eine (nicht unbedingt konvergente) Folge (x_k) in D mit

$$\lim_{k \rightarrow \infty} f(x_k) = y^*.$$

Man kann dann immer sogar eine *monotone* Folge (x_k) mit dieser Eigenschaft finden.

- (iii) f nimmt in x^* ein *lokales Maximum* an, wenn es ein $\epsilon > 0$ gibt, so dass

$$f(x^*) \geq f(x) \text{ für alle } x \in D \text{ mit } |x - x^*| < \epsilon,$$

d.h. für alle x in der Nähe von x^* .

Entsprechend definiert man die Begriffe *Minimum*, *Infimum* und *lokales Minimum*.

Hilfestellung zur Sortierung der Begriffe und Gedanken... Für die Bestimmung von Extremwerten differenzierbarer Funktionen gibt es die bekannte Bedingung $f'(x) = 0$. Diese Bedingung ist für innere Punkte des Definitionsbereichs eine notwendige Bedingung:

Nur (= höchstens) an diesen Stellen und in den Randpunkten des Definitionsbereichs kann f ein Extremum haben. Typischerweise sind das endlich viele Stellen, und man rechnet nach, an welcher davon $f(x)$ am größten (kleinsten) ist.

- Wenn man weiß, dass f sein Maximum irgendwo annimmt, dann hat man es damit gefunden.
- Wenn aber f sein Maximum gar nicht annimmt, dann auch nicht an der Nullstelle von f' mit größtem f -Wert. Die ganze Ableiterei und Nullstellensuche war umsonst. Schlimmer noch, sie liefert „Lösungen“ für ein Problem, das in Wirklichkeit keine hat.

Man ist also aus praktischen Gründen daran interessiert, einfache Bedingungen zu kennen, die einem garantieren, dass f sein Maximum irgendwo annimmt.

Satz 72 (Existenz vom Maximum und Minimum). Eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ auf einem kompakten Intervall nimmt ihr Supremum und ihr Infimum an, besitzt also ein Maximum und ein Minimum.

Die oben skizzierte Funktion f nimmt zwar nicht auf $\mathbb{R}$, sehr wohl aber z.B. auf $[-2, 2]$ ein Maximum an, nämlich an der Stelle $x = 2$.

Beweis. Wir zeigen den Satz nur für das Maximum. Der Beweis für das Minimum geht genauso. Sei $y^* := \sup f$. Wir müssen ein $x^* \in [a, b]$ finden, für das

$$f(x^*) = y^*.$$

Nach Definition des Supremums gibt es eine monotone Folge (x_k) in $[a, b]$ mit

$$\lim f(x_k) = y^*.$$

Die Folge (x_k) ist offenbar beschränkt:

$$a \leq x_k \leq b \quad \text{für alle } k.$$

Nach dem Monotoniekriterium ist sie also konvergent gegen ein x^* mit

$$a \leq x^* \leq b,$$

d.h. $x^* \in [a, b]$.

Und weil f stetig ist, folgt $f(x^*) = \lim_{k \rightarrow \infty} f(x_k) = y^*$. □

Kombiniert man die Begriffe $\sup$ und $\inf$ mit dem Zwischenwertsatz, so findet man:

Satz 73 (Surjektivitätskriterium). *Eine stetige Funktion $f : J \rightarrow \mathbb{R}$ auf einem Intervall J nimmt alle Werte (echt) zwischen ihrem Infimum und ihrem Supremum an:*

$$f(J) \supset] \inf_{x \in J} f(x), \sup_{x \in J} f(x) [.$$

Beispiel 74. Für ein Polynom $p(x)$ von ungeradem Grad (vgl. Abschnitt 4.1) ist $\inf_{x \in \mathbb{R}} p(x) = -\infty$ und $\sup_{x \in \mathbb{R}} p(x) = +\infty$. Warum?

Also ist $p(\mathbb{R}) = \mathbb{R}$. □

4 Elementare Funktionen I

- Für die Konstruktion von Funktionen aus einfacheren Funktionen sind *Linearkombinationen* das einfachste Hilfsmittel.

Sehr einfache Funktionen sind die *Potenzen* $x, x^2, \dots$ und erst recht die konstante Funktion $x^0 = 1$. Ein einfaches Verfahren, um aus gegebenen Funktionen neue zu machen, ist das folgende:

Definition 75 (Linearkombination). Sind $f_1(x), \dots, f_k(x) : D \rightarrow \mathbb{R}$ reellwertige Funktionen, und sind $a_1, \dots, a_k \in \mathbb{R}$, so heißt die Funktion

$$f(x) := a_1 f_1(x) + \dots + a_k f_k(x)$$

eine *Linearkombination* von $f_1(x), \dots, f_k(x)$ mit den *Koeffizienten* $a_1, \dots, a_k$. Dabei kann man überall die reellen Zahlen $\mathbb{R}$ auch durch die komplexen Zahlen $\mathbb{C}$ ersetzen.

Beispiel 76.

$$\sin(x + \frac{\pi}{3}) = \cos(\frac{\pi}{3}) \sin x + \sin(\frac{\pi}{3}) \cos x = \frac{1}{2} \sin x + \frac{1}{2} \sqrt{3} \cos x$$

ist eine Linearkombination von $\sin x$ und $\cos x$. □

Beispiel 77. Rechnen Sie nach: Sind $y_1(x)$ und $y_2(x)$ Lösungen der (linearen homogenen) Differentialgleichung

$$y'' + ay' + by = 0,$$

so ist auch jede Linearkombination der beiden eine Lösung. Wie kann man das verallgemeinern? Warum gilt das nicht für $y'' + ay' + by = x^2 - 3$? □

4.1 Polynome

Polynome sind

- Polynome sind Linearkombinationen von Potenzen.
- Wir lernen verschiedene Aspekte der Polynome und ihrer Anwendungen kennen.

Wir untersuchen die Linearkombinationen von Potenzen, die Polynome. Dabei ist es sehr empfehlenswert, komplexe Zahlen zugrunde zu legen, weil dann die Frage nach den Nullstellen viel einfacher wird.

Definition 78 (Polynom). Eine Linearkombination von endlich vielen Potenzen $z^0 = 1, z, \dots, z^n$ heißt ein *Polynom*:

$$f(z) := a_0 + a_1 z^1 + \dots + a_n z^n.$$

Die a_k heißen die *Koeffizienten* des Polynoms. Ist der höchste Koeffizient $a_n \neq 0$, so sagt man $f(z)$ hat den *Grad* n . Polynome vom Grad 0 sind also konstante Funktionen $\neq 0$. Der Nullfunktion gibt man den Grad $-\infty$.¹⁰

¹⁰Das hat den Vorteil, dass dann der Grad von $f(x)g(x)$ gleich der Summe der Grade von $f(x)$ und $g(x)$ ist, auch wenn eines der beiden 0 ist.

Beispiel 79. Die Funktion

$$\begin{aligned} f(z) &= 3 + 2(z-1) + 5(z-1)^3 \\ &= -4 + 17z - 15z^2 + 5z^3 \end{aligned}$$

ist ein Polynom. Überlegen Sie sich, dass man umgekehrt jedes Polynom $f(z)$ vom Grad n auch als Linearkombination

$$f(z) = a_0 + a_1(z-z_0)^1 + \dots + a_n(z-z_0)^n$$

mit einem beliebigen z_0 schreiben kann. □

Beispiel 80. Die Funktion

$$f(z) = 7(z-3)(z+i)(z-i) = 7z^3 - 21z^2 + 7z - 21$$

ist ein Polynom. In der ersten Form kann man die Nullstellen des Polynoms problemlos ablesen: Sie sind $z_1 = 3, z_2 = -i, z_3 = i$.

Überlegen Sie sich, dass man umgekehrt ...

Nein, das ist zu schwer, nämlich ein bedeutender Satz: □

Satz 81 (Fundamentalsatz der Algebra). *Ist f ein Polynom vom Grad $n \geq 1$, so gibt es eine bis auf die Reihenfolge eindeutige Linearfaktorzerlegung*

$$f(z) = a_n(z-z_1)(z-z_2)\dots(z-z_n).$$

In diesem Sinne hat jedes komplexe Polynom vom Grade n genau n Nullstellen (die aber nicht unbedingt voneinander verschieden sein müssen: „mehrfache“ Nullstellen muss man entsprechend ihrer Vielfachheit zählen).

Das Polynom $x^2 + 1$ zeigt, dass wir hier wirklich komplexe Zahlen brauchen!

Ein Polynom mit mindestens n Nullstellen ist daher entweder vom Grad $\geq n$, oder das Nullpolynom! Zwei Polynome $f(z), g(z)$ vom Grad $< n$, die an n Stellen denselben Wert haben, haben die gleichen Koeffizienten. Zum Beweis betrachtet man das Polynom $f(z) - g(z)$. Das hat dann auch Grad $< n$ und n Nullstellen, ist also das Nullpolynom. Daraus folgt

Satz 82 (Koeffizientenvergleich). *Sind*

$$f(z) = \sum a_k z^k, \quad g(z) = \sum b_k z^k$$

zwei Polynome mit $f(z) = g(z)$ für alle z , so haben $f(z)$ und $g(z)$ dieselben Koeffizienten: $a_k = b_k$ für alle k . Dabei genügt genau genommen die Gleichheit an $n+1$ Stellen, wenn die Grade von $f(z)$ und $g(z)$ höchstens n sind.

Beispiel 83 (Inhomogene lineare Differentialgleichungen). Wie findet man eine Funktion $y(x)$ mit

$$y'' - 3y' - 2y = x^2 + 1? \tag{25}$$

Gesucht ist also ein Lösung einer sogenannten *inhomogenen linearen Differentialgleichung mit konstanten Koeffizienten*. Durch scharfes Hinsehen kommt einem die Idee, $y(x)$ als Polynom vom Grad 2 anzusetzen:

$$y(x) = Ax^2 + Bx + C.$$

Dann ist $y'(x) = 2Ax + B$ und $y''(x) = 2A$. Einsetzen in die Gleichung liefert

$$2A - 3(2Ax + B) - 2(Ax^2 + Bx + C) = x^2 + 1$$

oder, nach Sortieren nach Potenzen,

$$-2Ax^2 - (6A + 2B)x + 2A - 3B - 2C = x^2 + 1.$$

Nach dem Satz ist also

$$-2A = 1, \quad -(6A + 2B) = 0 \text{ und } 2A - 3B - 2C = 1.$$

Das ist ein lineares Gleichungssystem, aus dem man findet:

$$A = -\frac{1}{2}, \quad B = \frac{3}{2}, \quad C = -\frac{13}{4}.$$

Also ist

$$y(x) = -\frac{1}{2}x^2 + \frac{3}{2}x - \frac{13}{4}$$

eine Lösung von (25). Machen Sie die Probe! □

Wie findet man Nullstellen? Ein Polynom in Linearfaktoren zu zerlegen ist gleichwertig damit, seine Nullstellen zu finden, d.h. eine Gleichung n -ten Grades zu lösen. Für quadratische Polynome haben Sie das in der Schule gelernt. Für Polynome von Grad 3 und 4 findet man Lösungen (durch Wurzelziehen) in Formelsammlungen, für Polynome von Grad ≥ 5 gibt es keine solchen Formeln mehr. Manchmal kann man eine Nullstelle raten, und kann dann durch Polynomdivision den Grad des Polynoms um 1 reduzieren. In speziellen Fällen helfen ad hoc-Überlegungen, wie im folgenden Beispiel:

Beispiel 84. Sei $p(z) = 3z^3 - 24i$. Die Nullstellen von p sind gerade die Lösungen der Gleichung

$$z^3 = 8i.$$

Wir erinnern uns an Beispiel 35 und schreiben das in Polarkoordinaten als

$$\begin{aligned} |z^3| &= |z|^3 = 8, \text{ also } |z| = 2, \\ 3 \arg z &= \arg(z^3) = \arg(8i) = \frac{\pi}{2} + 2k\pi, \end{aligned}$$

d.h.

$$\arg z_1 = \frac{\pi}{6}, \arg z_2 = \frac{5\pi}{6}, \arg z_3 = \frac{9\pi}{6} = \frac{3\pi}{2}.$$

Es folgt

$$\begin{aligned} z_1 &= 2\left(\cos \frac{\pi}{6} + i \sin \frac{\pi}{6}\right) = \sqrt{3} + i, \\ z_2 &= 2\left(\cos \frac{5\pi}{6} + i \sin \frac{5\pi}{6}\right) = -\sqrt{3} + i, \\ z_3 &= 2\left(\cos \frac{3\pi}{2} + i \sin \frac{3\pi}{2}\right) = -2i. \end{aligned}$$

Wir haben damit 3 Nullstellen des Polynoms gefunden, und nach dem Satz ist daher

$$3z^3 - 24i = 3(z - (\sqrt{3} + i))(z - (-\sqrt{3} + i))(z + 2i).$$

□

Bemerkung: Mathematik mit dem Rechner. Im allgemeinen aber findet man die Nullstellen von Polynomen höheren Grades nur numerisch, bei reellen Polynomen z.B. mit der Intervallhalbierungsmethode (vgl. Zwischenwertsatz). In der Praxis wendet man dazu natürlich geeignete mathematische Software an. In diesem Skript finden Sie in lockerer Folge Beispiele für den Einsatz eines Computeralgebra-Systems, nämlich **Mathematica**. Ähnlich ist das Maple-System, während z.B. Matlab eher auf numerische Probleme ausgerichtet ist. Es geht in dieser Vorlesung aber nicht darum, mit einem solchen System vertraut zu machen. Es geht vielmehr darum zu sehen, wie man derartige Systeme für die (Ingenieur-)Mathematik verwenden kann und Sie anzuregen, auf Ihrem eigenen Rechner weiter zu experimentieren. Die Syntax ist bei diesen Beispielen so einfach, dass Sie sie auch verstehen können, wenn Sie selbst bisher ein anderes oder gar kein Computer-Algebra-System benutzen.

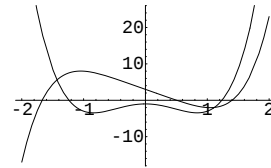
Beispiel 85. Mathematica löst numerisch eine Gleichung 6. Grades:

```
In[1]:= NSolve[z^5 - 6z + 3 == 0, z]
Out[1]={{z -> -1.67094}, {z -> -0.118104 - 1.58746 i},
{z -> -0.118104 + 1.58746 i}, {z -> 0.505501},
{z -> 1.40164}}
```

Reelle Polynome. Sind alle Koeffizienten reell, so nennt man das Polynom *reell*. Ein reelles Polynom gibt eine Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}$, die man durch ihren Graphen veranschaulichen kann. Das gibt die Bilder, die Sie aus der Schule kennen:

Beispiel 86.

Der Mathematica-Befehl
`Plot[{x^5 - 6x + 3, 5x^4 - 7x^2 - 1}, {x, -2, 2}]`
 liefert das nebenstehende Bild.
 Welche Kurve ist welches Polynom (und warum)?



□

Für ein Polynom mit reellen Koeffizienten, aber vielleicht komplexem Argument z , gilt nach den Rechenregeln für die Konjugation komplexer Zahlen:

$$\overline{f(z)} = \overline{\sum_{k=0}^n a_k z^k} = \sum_{k=0}^n \overline{a_k z^k} = \sum_{k=0}^n \overline{a_k} \overline{z^k} = \sum_{k=0}^n \overline{a_k} \overline{z}^k = f(\bar{z}).$$

Darum ist mit $z = a + ib$ auch $\bar{z} = a - ib$ eine Nullstelle. Wegen

$$(x - (a + ib))(x - (a - ib)) = (x - a)^2 + b^2,$$

erhält man deshalb für reelle Polynome eine Faktorisierung in lineare und quadratische Terme

$$f(x) = a_n(x - x_1) \cdot \dots \cdot (x - x_k)((x - a_{k+1})^2 + b_{k+1}^2) \cdot \dots \cdot ((x - a_{k+m})^2 + b_{k+m}^2).$$

Dabei ist $n = k + 2m$ der Grad, $x_1, \dots, x_k$ sind die reellen Nullstellen und

$$a_{k+1} + ib_{k+1}, a_{k+1} - ib_{k+1}, \dots, a_{k+m} + ib_{k+m}, a_{k+m} - ib_{k+m}$$

die komplexen.

Die komplexen Nullstellen eines reellen Polynoms treten in konjugiert-komplexen Paaren auf.

4.2 Rationale Funktionen

- Rationale Funktionen, also die Quotienten zweier Polynome, sind aus der Schule bekannt, vermutlich auch die Polynomdivision mit Rest.
- Weil rationale Funktionen oft vorkommen, ist es wichtig zu verstehen, dass sie sich aus besonders einfachen rationalen Funktionen aufbauen lassen. Dafür lernen wir die Methode der Partialbruchzerlegung.

Rationale Funktionen sind Quotienten von zwei Polynomen $f(z) = \frac{p(z)}{q(z)}$. Die Nullstellen des Nenners heißen auch die *Pole* der Funktion.

Beispiele sind Polynome ($q(z) = 1$)

$$a_n z^n + \dots + a_1 z + a_0,$$

oder die einfachen Ausdrücke

$$\frac{A}{z-a}, \quad \text{allgemeiner } \frac{A}{(z-a)^k}.$$

Wir wollen nun zeigen, dass sich jede **komplexe** rationale Funktion mit der sogenannten *Partialbruchzerlegung* additiv zerlegen lässt in Bausteine dieser einfachen Typen. Später sagen wir auch etwas zum reellen Fall, aber der ist technisch viel komplizierter.

Die Anwendungen. Das wird für Sie bei zwei Gelegenheiten wichtig:

- Bei der Integration von rationalen Funktionen und
- im Zusammenhang mit der sogenannten *Laplace transformation*, die in der Regelungstechnik, in der Netzwerktheorie, bei der Signalverarbeitung und in vielen anderen Anwendungsbereichen eine zentrale Rolle spielt, vgl. *Regelungstechnik I, Signale und Systeme, Einführung in die Netzwerktheorie* oder auch *Energie-, Impuls- und Stofftransport*.

In beiden Fällen möchte man rationale Funktionen tabellarisch erfassen und ist dann natürlich froh, wenn man sich auf die obigen einfachen Bausteine beschränken kann! Das Verfahren hat zwei Schritte, Polynomdivision und Partialbruchzerlegung.

POLYNOMDIVISION.

Durch Polynomdivision mit Rest ergibt sich für jede rationale Funktion eine Darstellung der Form

$$f(z) = g(z) + \frac{\tilde{p}(z)}{q(z)} \quad \tilde{p}(z), q(z), g(z) \text{ Polynome,}$$

wobei der Grad von $\tilde{p}(z)$ echt kleiner ist als der von $q(z)$. Wenn der Grad von $p(z)$ schon kleiner ist als der von $q(z)$, entfällt dieser Schritt natürlich.

Beispiel 87. Für $f(z) = \frac{z^4+3z+5}{z^2+1}$ ergibt sich:

$$\begin{array}{r} (z^4 + 3z + 5) : (z^2 + 1) = z^2 - 1 \quad \text{Rest } 3z + 6 \\ \underline{z^4 + z^2} \\ -z^2 + 3z \\ \underline{-z^2 - 1} \\ 3z + 6 \end{array}$$

Also ist $\frac{z^4+3z+5}{z^2+1} = z^2 - 1 + \frac{3z+6}{z^2+1}$. □

Beispiel 88. Für $f(z) = \frac{z^3 - z + 2}{5z^3 + 1}$ findet man ebenso

$$f(z) = \frac{1}{5} - \frac{z + \frac{24}{5}}{5z^3 + 1} = \frac{1}{5} - \frac{\frac{z}{5} + \frac{24}{25}}{z^3 + \frac{1}{5}}. \quad (26)$$

□

PARTIALBRUCHZERLEGUNG IM KOMPLEXEN.

Wir betrachten nun $f(z) = \frac{p(z)}{q(z)}$ mit

$$\text{Grad } p < \text{Grad } q$$

und nehmen an, dass wir die Nullstellen von q kennen:

$$q(z) = (z - z_1) \dots (z - z_n).^{11}$$

Die erfolgreiche Berechnung der Pole ist also die Voraussetzung für die Partialbruchzerlegung.

1. Fall: Einfache Pole. Wir wollen zunächst voraussetzen, dass $f(z)$ nur einfache Pole hat, d.h. dass $q(z)$ nur einfache Nullstellen hat. Die z_j sind also alle voneinander verschieden. Wir behaupten, dass es dann eindeutig bestimmte komplexe Zahlen A_j gibt, so dass

$$f(z) = \frac{p(z)}{q(z)} = \frac{A_1}{z - z_1} + \frac{A_2}{z - z_2} + \dots + \frac{A_n}{z - z_n} = \sum_{j=1}^n \frac{A_j}{z - z_j}. \quad (27)$$

Das nennen wir die *Partialbruchzerlegung* von f .

Zunächst nehmen wir an, dass die Existenz einer solchen Darstellung schon klar sei. Wie findet man dann die A_j ? (Diese Frage stellt sich natürlich immer, wenn man die Partialbruchzerlegung benutzen will.) Wir multiplizieren die Partialbruchzerlegung mit $q(z)$:

$$p(z) = A_1(z - z_2) \dots (z - z_n) + \dots + A_n(z - z_1) \dots (z - z_{n-1}).$$

Setzen wir $z = z_1$ ein,¹² so verschwinden auf der rechten Seite alle Summanden bis auf den ersten, und wir finden

$$A_1 = \frac{p(z_1)}{(z_1 - z_2) \dots (z_1 - z_n)}.$$

Dasselbe Ergebnis erhalten wir, wenn wir in $f(z) = \frac{p(z)}{q(z)}$ den Faktor $(z - z_1)$ im Nenner streichen und dann z_1 einsetzen. Für die andern A_j geht das entsprechend. Wir sehen also: Wenn sich $f(z)$ so darstellen lässt, sind die A_j zwangsläufig bestimmt und leicht zu berechnen:

Man streicht bei $\frac{p(z)}{q(z)}$ den Faktor $(z - z_j)$ im Nenner und setzt $z = z_j$ ein.
Das Resultat ist A_j .

¹¹Den Koeffizienten der höchsten Nennerpotenz können wir „in den Zähler schieben“ und ohne Einschränkung annehmen, dass $q(z) = z^n + \dots$. Vgl. (26).

¹²Eigentlich gelten die Partialbruchzerlegung und damit diese Gleichung nur für die $z \neq z_1, \dots, z_n$, so dass man z_1 gerade nicht einsetzen darf. Aber hier stehen zwei Polynome, die an unendlich vielen Stellen gleich sind. Also sind sie nach dem Koeffizientenvergleich (Satz 82) überall gleich, auch für $z = z_1$.

Weil man in der Praxis $(z - z_j)$ oft nicht streicht, sondern einfach zuhlt, heit diese Methode auch die *Zuhaltemethode*.

Eine Probe zeigt, dass mit den so berechneten A_j tatschlich (27) gilt. Das beweist die Existenz der Partialbruchzerlegung.

Beispiel 89. Wir betrachten

$$f(z) = \frac{z^2 + z + 4}{(z - 3)(z^2 + 1)} = \frac{z^2 + z + 4}{(z - 3)(z - i)(z + i)} = \frac{A}{z - 3} + \frac{B}{z - i} + \frac{C}{z + i}.$$

Nach der Zuhaltemethode ergibt sich:

$$\begin{aligned} A &= \frac{3^2 + 3 + 4}{3^2 + 1} = \frac{8}{5}, \\ B &= \frac{i^2 + i + 4}{(i - 3)(i + i)} = -\frac{3 + i}{2 + 6i} = -\frac{(3 + i)(2 - 6i)}{(2 + 6i)(2 - 6i)} = -\frac{3}{10} + \frac{2}{5}i, \\ C &= \frac{(-i)^2 - i + 4}{(-i - 3)(-i - i)} = -\frac{3 - i}{2 - 6i} = -\frac{3}{10} - \frac{2}{5}i. \end{aligned}$$

Die Partialbruchzerlegung ist also

$$\frac{z^2 + z + 4}{(z - 3)(z^2 + 1)} = \frac{\frac{8}{5}}{z - 3} + \frac{-\frac{3}{10} + \frac{2}{5}i}{z - i} + \frac{-\frac{3}{10} - \frac{2}{5}i}{z + i}.$$

□

Beispiel 90. Mit derselben Methode findet man

$$\frac{1}{z(z + 1)} = \frac{1}{z} - \frac{1}{z + 1}.$$

Damit erhlt man

$$\sum_{k=1}^n \frac{1}{k(k + 1)} = \sum_{k=1}^n \left(\frac{1}{k} - \frac{1}{k + 1} \right) = \left(1 - \frac{1}{2} \right) + \left(\frac{1}{2} - \frac{1}{3} \right) + \dots + \left(\frac{1}{n} - \frac{1}{n + 1} \right) = 1 - \frac{1}{n + 1} = \frac{n}{n + 1}.$$

Diese Formel knnte man auch mit vollstndiger Induktion beweisen. Knnen Sie? □

2. Fall: Mehrfache Pole. Bei einer k -fachen Nullstelle z_j enthlt die Faktorzerlegung von $q(z)$ einen Term $(z - z_j)^k$ (und keine hhere Potenz davon). In diesem Fall muss man in der Partialbruchzerlegung auer dem Term $\frac{A_j}{z - z_j}$ noch Terme mit hheren Nennerpotenzen zulassen, nmlich

$$\frac{A_{j1}}{z - z_j} + \frac{A_{j2}}{(z - z_j)^2} + \dots + \frac{A_{jk}}{(z - z_j)^k}.$$

Wir beweisen das nicht.

Die Zuhaltemethode liefert den Koeffizienten A_{jk} zur hchsten Potenz von $(z - z_j)$. Um die anderen Koeffizienten $A_{j1}, \dots, A_{jk-1}$ zu finden, kann man den Zerlegungsansatz mit dem Hauptnenner multiplizieren und einen Koeffizientenvergleich machen. Oft ist es aber einfacher vor oder nach der Multiplikation mit dem Hauptnenner spezielle Werte fr z einzusetzen.

Beispiel 91.

$$\frac{5z^2 - 4z + 7}{z^3 + z^2 - 5z + 3} = \frac{5z^2 - 4z + 7}{(z - 1)^2(z + 3)} = \frac{A}{z - 1} + \frac{B}{(z - 1)^2} + \frac{C}{z + 3}.$$

Die Zuhalttemethode liefert

$$B = \frac{5 - 4 + 7}{4} = 2, \quad C = \frac{45 + 12 + 7}{16} = 4.$$

Durch Multiplikation mit dem Hauptnenner finden wir

$$5z^2 - 4z + 7 = A(z - 1)(z + 3) + 2(z + 3) + 4(z + 1)^2,$$

und $z = 0$ liefert

$$7 = -3A + 6 + 4.$$

Also ist $A = 1$. □

PARTIALBRUCHZERLEGUNG IM REELLEN.

Wir betrachten jetzt den Fall reeller rationaler Funktionen $\frac{p(x)}{q(x)}$. Sind alle Pole reell, so sieht die Partialbruchzerlegung aus wie oben, nur dass die Koeffizienten A_j bzw. A_{jk} natürlich auch reell sind. Aber der Nenner kann ja auch komplexe Nullstellen haben, die dann in konjugiert-komplexen Paaren $a \pm ib$ auftreten. Wir behandeln hier nur den Fall *einfacher* komplexer Pole. Dann kann man die Linearfaktoren von $q(x)$ zusammenfassen als

$$(x - (a + ib))(x - (a - ib)) = (x - a)^2 + b^2.$$

Die entsprechenden komplexen Summanden kann man ebenfalls zusammenfassen

$$\frac{A}{x - (a + ib)} + \frac{B}{x - (a - ib)} = \frac{Cx + D}{(x - a)^2 + b^2}$$

und in dieser Form ansetzen.

Beispiel 92. Wir betrachten wieder die Funktion aus Beispiel 89. Aus

$$\frac{z^2 + z + 4}{(z - 3)(z^2 + 1)} = \frac{A}{z - 3} + \frac{Bz + C}{z^2 + 1}$$

folgt durch Multiplikation mit dem Nenner:

$$z^2 + z + 4 = A(z^2 + 1) + (Bz + C)(z - 3) = (A + B)z^2 + (C - 3B)z + A - 3C.$$

Also

$$\begin{array}{rcl} A & +B & = 1 \\ & -3B & +C = 1 \\ A & & -3C = 4. \end{array}$$

Mit dem Gaußalgorithmus findet man $A = \frac{8}{5}$, $B = -\frac{3}{5}$, $C = -\frac{4}{5}$. □

Rationale Funktionen und ihre Partialbruchzerlegung spielen eine wichtige Rolle in verschiedenen Anwendungsbereichen, zum Beispiel in Theorie linearer Systeme (LTI-Systeme), die durch Übertragungsfunktionen charakterisiert sind, etwa in der Regelungstechnik oder der Signalverarbeitung. Es lohnt sich daher, die Technik der Partialbruchzerlegung zu üben.

Wir beschließen diesen Abschnitt mit einer Zusammenfassung der Schritte zur Zerlegung einer rationalen Funktion in einfache Summanden:

1. Polynomdivision, falls $\text{Grad } p(z) \geq \text{Grad } q(z)$.
2. Partialbruchzerlegung:

- (a) Bestimmung der Nullstellen des Nenners.
- (b) Richtigen Typ für die Partialbruchzerlegung ansetzen:
 - Ein einfacher Pol z_j erfordert einen Term

$$\frac{A_j}{z - z_j}.$$

- Ein k -facher Pol z_j erfordert k Terme

$$\frac{A_{j1}}{z - z_j} + \frac{A_{j2}}{(z - z_j)^2} + \dots + \frac{A_{jk}}{(z - z_j)^k}.$$

- Bei reellen Koeffizienten und einem einfachen Pol $z_j = a_j + ib_j$ ist auch $a_j - ib_j$ ein Pol. Dann kann man die beiden zugehörigen komplexen Terme zu einem reellen zusammenfassen:

$$\frac{C_j x + D_j}{(x - a_j)^2 + b_j^2} \quad \text{statt} \quad \frac{A_j}{x - (a_j + ib_j)} + \frac{B_j}{x - (a_j - ib_j)}$$

- (c) Koeffizienten bestimmen (Zuhaltemethode oder Koeffizientenvergleich).

5 Differentiation

5.1 Die Ableitung

- Warum ist die Differentialrechnung für die Beschreibung von Naturvorgängen so wichtig?
- Wir wiederholen die aus der Schule bekannte Definition der Ableitung und interpretieren sie auf mehrere verschiedene Weisen.

Wir kommen nun zur Differentialrechnung. Ein Grund für deren große Bedeutung in der Naturwissenschaft und ihren Anwendungen in der Technik ist folgender: Unsere Modelle von der Natur sind weitgehend deterministisch: wir sind überzeugt, dass die zukünftigen Zustände eines Systems bestimmt sind durch die *Änderungen*, die dieses System erfährt, und dass diese Änderungen durch (äußere und innere) Einwirkungen verursacht werden (Kausalität). Dabei meinen wir mit *Änderung* die Änderung in jedem Augenblick, und die *Ableitung* einer Funktion (von der Zeit) beschreibt eben gerade diese momentane Änderung. Aber nicht nur zeitliche Entwicklungen beschreibt man auf solche Weise, auch bei räumlichen Entwicklungen bedient man sich gern der Methoden der Differentialrechnung, etwa wenn es darum geht zu erklären, welches elektrische Feld eine Ladung oder Ladungsverteilung erzeugt.

Das für die Anwendungen in der Technik wichtigste Werkzeug auf der Basis der Differentialrechnung sind die Differentialgleichungen: von Schwingungsphänomenen bis zur Gestalt von Zahnrädern, von der Wärmeleitung bis zur Strömungsmechanik, von der Beschreibung elektrischer Netzwerke bis zur Steuerung chemischer Prozesse bieten Differentialgleichungen die wichtigsten Modelle. Dafür legen wir hier mit der Differentialrechnung die Grundlagen.

Generalvoraussetzung. In diesem Abschnitt betrachten wir nur Funktionen, deren Definitionsbereich I aus einem oder mehreren Intervallen besteht, die jeweils nicht nur aus einem Punkt bestehen.¹³ Als Zielbereich der Funktion nehmen wir die reellen Zahlen, die folgende Definition macht aber auch Sinn, wenn die Werte von f komplex oder gar Vektoren sind. Dagegen ist die Differentialrechnung für Funktionen mit komplexem oder vektorielltem *Argument* eine ganz andere Theorie (Komplexe Funktionentheorie bzw. Theorie der partiellen Ableitungen).

Definition 93 (Differenzierbarkeit). $f : \mathbb{R} \supset I \rightarrow \mathbb{R}$ heißt in $x_0 \in I$ *differenzierbar*, wenn

$$f'(x_0) := \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0}$$

existiert. $f'(x_0)$ heißt dann *die Ableitung* von f in x_0 . Man nennt f differenzierbar auf I , wenn es in allen $x_0 \in I$ differenzierbar ist.

Bemerkungen. Man kann die Definition umformulieren, indem man x für x_0 und $\Delta x := x - x_0$ einführt:

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

Dann ist also f in x differenzierbar, wenn

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

existiert. Den Zähler rechts bezeichnet man auch gern mit Δf und schreibt

$$f'(x) = \lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x} =: \frac{df}{dx}(x).$$

¹³Der Definitionsbereich rationaler Funktionen erfüllt zum Beispiel diese Voraussetzung.

Wenn das Argument von f die Zeit bezeichnet, schreibt man dafür gern t statt x , obwohl es mathematisch gesehen auf die Bezeichnung gar nicht ankommt, und bezeichnet die Ableitung mit einem Punkt statt mit einem Strich:

$$\dot{f} := \frac{df}{dt}.$$

Beispiel 94. Sei $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$. Dann ist f differenzierbar:

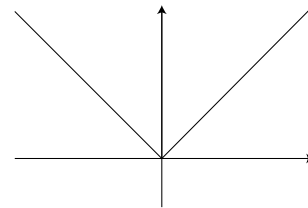
$$\lim_{\Delta x \rightarrow 0} \frac{(x + \Delta x)^2 - x^2}{\Delta x} = \lim_{\Delta x \rightarrow 0} \frac{2x\Delta x + (\Delta x)^2}{\Delta x} = \lim_{\Delta x \rightarrow 0} (2x + \Delta x) = 2x.$$

Also $f'(x) = 2x$. Ebenso zeigt man $g'(x) = 1$ für $g(x) := x$. □

Beispiel 95. Sei $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto |x|$. Dann ist f in 0 nicht differenzierbar, denn

$$\frac{f(0 + \Delta x) - f(0)}{\Delta x} = \frac{|\Delta x|}{\Delta x} = \begin{cases} +1 & \text{für } \Delta x > 0 \\ -1 & \text{für } \Delta x < 0 \end{cases}.$$

Der Grenzwert für $\Delta x \rightarrow 0$ existiert nicht. □



Die Funktion $|x|$ ist aber stetig. Stetige Funktionen sind also nicht unbedingt differenzierbar. Hingegen gilt:

Satz 96. Differenzierbare Funktionen sind stetig.

Genauer sind Funktionen dort, wo sie differenzierbar sind, auch stetig, wie der Beweis zeigt:

Beweis.

$$f(x) - f(x_0) = \frac{f(x) - f(x_0)}{x - x_0} (x - x_0) \rightarrow f'(x_0) \cdot 0 = 0 \quad \text{für } x \rightarrow x_0.$$

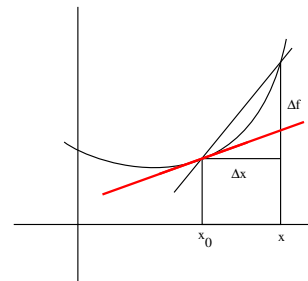
□

Geometrische (anschauliche) Interpretation: Die Tangente.

$\frac{f(x) - f(x_0)}{x - x_0}$ ist die Steigung der Sekante durch die Punkte $(x_0, f(x_0))$ und $(x, f(x))$. Der Grenzwert $f'(x_0)$ gibt die Steigung der Tangente an.

Die **Gleichung der Tangente** durch $(x_0, f(x_0))$ ist also gegeben durch

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



Differenzierbare Funktionen sind also solche, deren Graph eine Tangente besitzt. Vergleichen Sie insbesondere das obige (Gegen)beispiel der Funktion $|x|$.

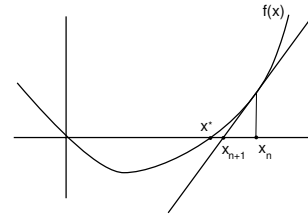
Beispiel 97 (Newtonverfahren). Beim Newtonverfahren zur Bestimmung der Nullstellen einer Funktion f beginnt man mit einem Wert x_1 möglichst nah an einer vermuteten

Nullstelle. Dann ermittelt man den Schnittpunkt x_2 der Tangente an den Graphen in diesem Punkt mit der x -Achse. Oft ist dies eine bessere Approximation für die Nullstelle. Iteration dieses Verfahrens liefert eine Folge, die unter gewissen, hier nicht diskutierten Voraussetzungen gegen eine Nullstelle x^* von f konvergiert. Die Formel für die Iteration erhält man durch Auflösen der Gleichung

$$f(x_n) + f'(x_n)(x_{n+1} - x_n) = 0.$$

Es ergibt sich

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



Wählt man etwa $f(x) = x^2 - a$, so findet man die bekannte Iteration für die Wurzel aus $a > 0$

$$x_{n+1} = x_n - \frac{x_n^2 - a}{2x_n} = \frac{1}{2} \left(x_n + \frac{a}{x_n} \right).$$

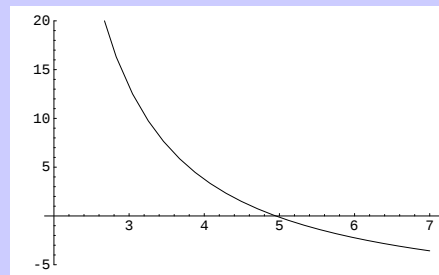
□

Beispiel 98 (Kapitalwert). In einer typischen Fragestellung der Investitionsrechnung ist eine (endliche oder unendliche) Zahlenfolge $a_0, a_1, a_2, \dots$ gegeben, wobei $a_0 < 0$ die Anschaffungsauszahlung bedeutet, während die positiven a_j für $j > 0$ die Einzahlungen zu äquidistanten Zeitpunkten bezeichnen, die mit x^{-j} abgezinst werden.

Der Kapitalwert ist dann die Funktion

$$K(x) = \sum_{j \geq 0} a_j x^{-j},$$

und man möchte wissen, bei welchem Zins x der Kapitalwert auf 0 fällt. Der Graph von K ist monoton fallend und konvex, wie in der Abbildung, und für die Berechnung der Nullstelle wird gern das Newtonverfahren benutzt.



Analytische Interpretation: Lineare Approximation. Die einfachsten Funktionen (nach den Konstanten) sind sicherlich die linearen Funktionen $a + bx$ mit konstanten a, b . Die Tangente ist der Graph einer solchen Funktion

$$f'(x_0)(x - x_0) + f(x_0) = \underbrace{(f(x_0) - f'(x_0)x_0)}_{=a} + \underbrace{f'(x_0)}_{=b} x,$$

und diese zeichnet sich dadurch aus, dass sie die Funktion in der Nähe des Punktes x_0 besonders gut approximiert: Für den Fehler

$$R(x) := f(x) - (f'(x_0)(x - x_0) + f(x_0)) = f(x) - f(x_0) - f'(x_0)(x - x_0)$$

gilt nämlich nicht nur $\lim_{x \rightarrow x_0} R(x) = 0$, sondern der Fehler geht für $x \rightarrow x_0$ schneller gegen null als $x - x_0$, also schneller als linear:

$$\lim_{x \rightarrow x_0} \frac{R(x)}{x - x_0} = 0.$$

Man kann daher die Differentiation auch auffassen als die Approximation von Funktionen durch lineare Funktionen mit „schnell verschwindendem“ Fehler:

$$f(x) = \underbrace{f(x_0) + f'(x_0)(x - x_0)}_{\text{linear}} + \underbrace{R(x)}_{\text{bei } x_0 \text{ „kleiner“ Fehler}}, \quad \lim_{x \rightarrow x_0} \frac{R(x)}{x - x_0} = 0.$$

Physikalische Interpretation: Die Geschwindigkeit. Bewegt sich ein Punkt auf einer Geraden mit konstanter Geschwindigkeit v , so legt er in der Zeit Δt die Strecke

$$\Delta s = v \Delta t$$

zurück. Umgekehrt kann man aus der Strecken- und Zeitdifferenz die Geschwindigkeit ermitteln:

$$v = \frac{\Delta s}{\Delta t}.$$

Bewegt sich der Punkt mit variabler Geschwindigkeit, und hat er zur Zeit t die Strecke $s(t)$ zurückgelegt, so kann man annehmen, dass wenigstens im *kleinen* Zeitabschnitt von t bis $t + \Delta t$ die Geschwindigkeit $v(t)$ annähernd konstant ist, also

$$v(t) \approx \frac{\Delta s}{\Delta t} := \frac{s(t + \Delta t) - s(t)}{\Delta t}.$$

Durch Grenzübergang $\Delta t \rightarrow 0$ erhält man die *Momentangeschwindigkeit* $v(t)$ zur Zeit t :

$$v(t) := \lim_{\Delta t \rightarrow 0} \frac{s(t + \Delta t) - s(t)}{\Delta t}.$$

Rechenregeln. Die Rechenregeln für die Differentiation sind ja aus der Schule gut bekannt: wenn f und g differenzierbar sind, gilt

$$\begin{aligned} \text{Summenregel:} \quad & (f + g)'(x) = f'(x) + g'(x) \\ \text{konstante Faktoren:} \quad & (cf)'(x) = cf'(x) \\ \text{Produktregel:} \quad & (fg)'(x) = f'(x)g(x) + f(x)g'(x) \\ \text{Quotientenregel:} \quad & \left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g^2(x)} \\ \text{Kettenregel:} \quad & (f \circ g)'(x) = (f' \circ g)(x)g'(x) = f'(g(x))g'(x) \end{aligned}$$

Wir verzichten hier auf die Beweise. Bis auf den zur Kettenregel sind sie alle sehr einfach.

Satz 99 (Ableitung der Umkehrfunktion). Sei $f : I \rightarrow J \subset \mathbb{R}$ differenzierbar und umkehrbar mit differenzierbarer Umkehrfunktion $f^{-1} : J \rightarrow I$. Aus

$$x = f(f^{-1}(x))$$

folgt dann durch Differentiation und Anwendung der Kettenregel:

$$1 = f'(f^{-1}(x)) \cdot (f^{-1})'(x),$$

also

$$(f^{-1})'(x) = \frac{1}{f'(f^{-1}(x))},$$

falls der Nenner $\neq 0$.

Man muss nicht unbedingt voraussetzen, dass auch f^{-1} differenzierbar ist. Wenn die Ableitung im Nenner $\neq 0$ ist, dann folgt die Differenzierbarkeit von f^{-1} von selbst.

Beispiel 100 (Wurzel). Die Funktion

$$]0, +\infty[\rightarrow \mathbb{R}, \quad x \mapsto x^2$$

ist differenzierbar und umkehrbar. Für ihre Inverse, nämlich die Wurzelfunktion

$$]0, +\infty[\rightarrow \mathbb{R}, \quad x \mapsto \sqrt{x}$$

gilt

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

□

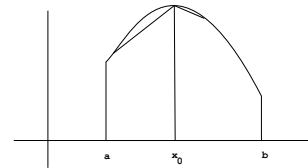
5.2 Extremwerte, Mittelwertsatz und Konsequenzen

- Mit dem Mittelwertsatz lernen wir den wichtigsten Satz der Differentialrechnung kennen.
- Der Satz, dass Funktionen mit verschwindender Ableitung konstant sind, gibt Anlass, über Ursache und Wirkung und über die logische Schlussrichtung nachzudenken.
- Weitere Anwendungen des Mittelwertsatzes.

Wenn die Funktion

$$f : \mathbb{R} \supset [a, b] \rightarrow \mathbb{R}$$

in einem inneren Punkt x_0 des Intervalls ihr Maximum annimmt, dann haben die Sekanten links davon eine Steigung ≥ 0 und die Sekanten rechts davon eine Steigung ≤ 0 .



Ist f in x_0 differenzierbar, so ist die Ableitung $f'(x_0)$ als Grenzwert der Sekantensteigungen einerseits ≥ 0 , andererseits aber ≤ 0 und deshalb $= 0$.

Satz 101 (Notwendiges Extremwertkriterium). *Nimmt eine differenzierbare Funktion ihr Maximum (oder Minimum) in einem inneren Punkt des Intervalls an, so ist ihre Ableitung dort 0.*

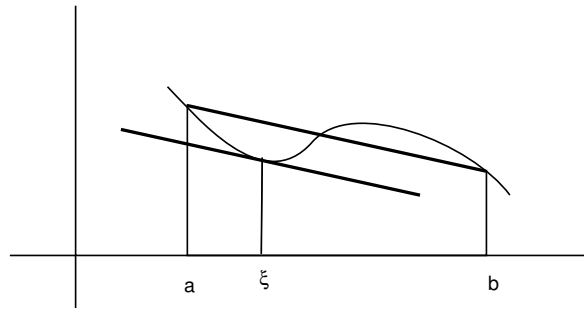
In Randpunkten muss das natürlich nicht so sein. Ist $f : \mathbb{R} \supset [a, b] \rightarrow \mathbb{R}$ differenzierbar, so sind die einzigen Stellen, die für Extremwerte von f in Frage kommen die Endpunkte a, b und die Punkte, wo die Ableitung $f'(x)$ verschwindet. Typischerweise sind das endlich viele Punkte, und man kann dann nachrechnen, wo f am größten bzw. kleinsten ist. Das ist dann wirklich das Maximum oder Minimum, weil f stetig auf einem kompakten Intervall ist.

Der folgende Satz ist von fundamentaler Wichtigkeit in der Differentialrechnung, obwohl er bei näherem Hinsehen ganz selbstverständlich erscheint.

Satz 102 (Mittelwertsatz). *Sei $f : I \rightarrow \mathbb{R}$ differenzierbar auf dem Intervall $I \subset \mathbb{R}$. Sind $a, b \in I$ beliebig mit $a \neq b$, so gibt es dazwischen ein ξ mit*

$$\frac{f(b) - f(a)}{b - a} = f'(\xi). \quad (28)$$

Anschaulich bedeutet das, dass es irgendwo eine Tangente gibt, die parallel zur Sekante durch $(a, f(a))$ und $(b, f(b))$ geht:



Wie beim Zwischenwertsatz belegt zum Beispiel die Funktion $f(x) = \frac{1}{x}$, dass I dabei ein *Intervall* sein muss.

Beweis. Die Sekante ist gegeben durch

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

und ein Blick auf die Abbildung zeigt, dass ein guter Kandidat für ξ der Punkt ist, wo die Differenz zwischen Funktion und Sekante maximal oder minimal wird. Wir betrachten deshalb Extremwerte von

$$g(x) := f(x) - \left(f(a) + \frac{f(b) - f(a)}{b - a}(x - a) \right).$$

Nach Konstruktion ist $g(a) = g(b) = 0$, aber das können Sie auch nachrechnen. Und dann rechnen Sie noch nach, dass aus $g'(\xi) = 0$ die Gleichung (28) folgt.

Die Funktion g ist differenzierbar und deshalb stetig. Wir betrachten sie auf dem kompakten Intervall zwischen a und b . Dort nimmt sie ihr Maximum und Minimum an. Sind die beiden gleich, so ist g konstant und $g'(\xi) = 0$ für alle ξ zwischen a und b . Fertig!

Sind Maximum und Minimum verschieden, so wird wegen $g(a) = g(b)$ eines der beiden in einem inneren Punkt ξ angenommen. Dort ist dann $g'(\xi) = 0$, und wir sind auch fertig. $\square$

Warum ist der Mittelwertsatz wichtig? Zunächst sagt der Mittelwertsatz nichts darüber aus, wo ξ eigentlich liegt. Vielleicht gibt es ja auch mehrere solche Stellen, vgl. Abbildung. Die wichtigen Anwendungen des Mittelwertsatzes „laufen andersherum“: Sie betreffen Situationen, wo man das Verhalten der Ableitung (überall!) sehr gut kennt, und daraus Schlüsse auf die Funktion zieht. *Der Mittelwertsatz ist die Brücke von Ableitungsinformationen zu Informationen über die Funktion selbst.* Hier zwei typische Beispiele:

Satz 103 (Schranksatz). Ist $|f'(x)| \leq M$ für alle $x \in [a, b]$, so ist

$$|f(b) - f(a)| \leq M(b - a).$$

Beweis.

$$|f(b) - f(a)| = |f'(\xi)|(b - a) \leq M(b - a).$$

$\square$

Beispiel 104. Wenn Sie 13 km von der Uni weg wohnen und in einer Viertelstunde zu Hause sind, sind Sie schneller als 50 km/h gefahren: $50 \frac{\text{km}}{\text{h}} \times \frac{1}{4} \text{ h} = 12.5 \text{ km}$. $\square$

Satz 105 (Monotoniekriterium). Sei $f : \mathbb{R} \supset I \rightarrow \mathbb{R}$ differenzierbar auf dem Intervall I und $f'(x) > 0$ für alle $x \in I$. Dann ist f streng monoton steigend:

$$a, b \in I \text{ und } a < b \implies f(a) < f(b).$$

Hat man nur $f'(x) \geq 0$, so folgt die schwache Monotonie: $f(a) \leq f(b)$.
Entsprechendes gilt für $f' < 0$ bzw. $f' \leq 0$.

Beweis. Seien $a, b \in I$ und $a < b$. Dann gibt es nach dem Mittelwertsatz ein ξ zwischen a und b mit

$$f(b) - f(a) = \underbrace{f'(\xi)}_{>0} \underbrace{(b-a)}_{>0} > 0.$$

Also ist $f(a) < f(b)$. □

Beispiel 106. Wenn Sie zeigen sollen (wollen?), dass der Tangens $\tan = \frac{\sin}{\cos}$ auf $]-\frac{\pi}{2}, +\frac{\pi}{2}[$ monoton wachsend ist, können Sie versuchen zu zeigen, dass

$$x < y \implies \frac{\sin x}{\cos x} < \frac{\sin y}{\cos y}.$$

Einfacher ist es aber, die Ableitung zu berechnen und das Monotoniekriterium anzuwenden:
Nach der Quotientenregel findet man

$$\tan' x = \frac{\cos^2 x + \sin^2 x}{\cos^2 x} = \frac{1}{\cos^2 x} > 0.$$

Fertig! □

Eine ganz wichtige Konsequenz des Mittelwertsatzes kombiniert die Monotonieaussagen:
Funktionen, die gleichzeitig monoton wachsen und fallen, sind konstant.

Satz 107 (Konstanzkriterium). Ist die Ableitung von $f : I \rightarrow \mathbb{R}$ auf dem Intervall I überall null, so ist f konstant.

Beachten Sie: Aus der Definition der Ableitung folgt trivial, dass konstante Funktionen die Ableitung 0 haben. Hier wird aber die *Umkehrung* behauptet und aus dem Mittelwertsatz bewiesen.

„Konstant“ heißt „Alle Sekanten haben Steigung 0“. Es ist aber eben nicht so klar, dass alle Sekanten die Steigung null haben, wenn die *Grenzwerte* der Sekantensteigungen alle null sind. Wenn der Grenzwert einer Folge null ist, können die Folgenglieder ja durchaus positiv sein!

Beispiel 108. Seien $f_1, f_2 : \mathbb{R} \rightarrow \mathbb{R}$ zwei Funktionen, die die Wachstumsgleichung

$$f' = af \tag{29}$$

mit derselben Konstanten $a > 0$ erfüllen. Wir wollen der Einfachheit halber auch noch voraussetzen, dass f_2 keine Nullstellen hat. Dann ist

$$\left(\frac{f_1}{f_2}\right)' = \frac{f_1' f_2 - f_1 f_2'}{f_2^2} = \frac{af_1 f_2 - af_1 f_2}{f_2^2} = 0.$$

Also ist $\frac{f_1}{f_2}$ konstant. Je zwei Lösungen der Wachstumsgleichung sind also gleich bis auf einen konstanten Faktor. □

5.3 Höhere Ableitungen

- Höhere Ableitungen kommen in der Physik oft vor, das Musterbeispiel ist die Beschleunigung, also die Ableitung der Geschwindigkeit, also die Ableitung der Ableitung der zurückgelegten Strecke.

Ist $f : \mathbb{R} \supset I \rightarrow \mathbb{R}$ auf ganz I differenzierbar, so ist die Ableitung wieder eine Funktion $f' : \mathbb{R} \supset I \rightarrow \mathbb{R}$. Ist diese auch differenzierbar, so nennt man ihre Ableitung $(f')' =: f''$ die 2. Ableitung von f und sagt f sei zweimal differenzierbar. Entsprechend definiert man induktiv die k -malige Differenzierbarkeit und k -te Ableitung für alle $k \in \mathbb{N}$. Notation, weil man nicht beliebig viele Striche machen möchte:

$$f^{(k)} = \frac{d^k f}{dx^k}.$$

Man setzt $f^{(0)} := f$. Nicht jede Funktion ist differenzierbar, nicht einmal jede stetige Funktion, wie die Funktion $|x|$ zeigt. Und eine Funktion, die einmal differenzierbar ist, muss nicht unbedingt zweimal differenzierbar sein, wie das folgende Beispiel zeigt.

Beispiel 109. Die Funktion

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

ist differenzierbar. Zeigen Sie das insbesondere im Nullpunkt! Die Ableitung ist $f'(x) = 2|x|$. Daher ist f in 0 nicht zweimal differenzierbar. $\square$

Beispiel 110. Ein typisches Beispiel für höhere Ableitungen ist die Beschleunigung als momentane Änderung der Geschwindigkeit, die selbst die momentane Änderung des Ortes ist:

$$\vec{b} = \dot{\vec{v}} = \ddot{\vec{x}}.$$

Hier benutzt man wieder Punkte zur Kennzeichnung der Ableitung nach einer Zeitvariablen.

Beispiel 111 (Gleichung für LRC-Glied). Sei $Q(t)$ die Ladung eines Kondensators in Abhängigkeit von der Zeit. Dann ist $I(t) = \dot{Q} = \frac{dQ}{dt}$ der zufließende Strom, und die Änderung der Stromstärke $\dot{I} = \frac{dI}{dt} = \ddot{Q} = \frac{d^2Q}{dt^2}$ ist proportional zur induktiven Spannung an einer Spule. Für den Spannungsabfall einer Reihen-LRC-Schaltung gilt daher

$$U(t) = L\dot{I} + RI + \frac{Q}{C} = L\ddot{Q} + R\dot{Q} + \frac{Q}{C}.$$

Beispiel 112. Wir berechnen die Ableitungen von

$$f(x) := (x - x_0)^n.$$

Es ist

$$\begin{aligned}
 f(x) &= (x - x_0)^n \\
 f'(x) &= n(x - x_0)^{n-1} \\
 f''(x) &= n(n-1)(x - x_0)^{n-2} \\
 &\vdots \\
 f^{(k)}(x) &= n \dots (n-k+1)(x - x_0)^{n-k} \\
 &\vdots \\
 f^{(n)}(x) &= n! \\
 f^{(n+1)}(x) &= 0
 \end{aligned}$$

An der Stelle $x = x_0$ ist deshalb nur die n -te Ableitung von Null verschieden, nämlich $= n!$. Die niedrigeren Ableitungen verschwinden, weil sie einen Faktor $x_0 - x_0 = 0$ enthalten, die höheren, weil „alles weg-differenziert ist“.

Ist daher $f(x) = a_0 + a_1(x - x_0) + \dots + a_n(x - x_0)^n$ ein Polynom vom Grade n , so ist für alle k

$$f^{(k)}(x_0) = k! a_k,$$

wenn man $a_k = 0$ setzt, falls $k > n$. Also gilt $a_k = \frac{f^{(k)}(x_0)}{k!}$ und

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k. \quad (30)$$

□

Beispiel 113. Schreibe $f(x) = x^3 + 2x^2 - x + 6$ in der Form $\sum_{k=0}^3 a_k(x-2)^k$. Dazu benutzen wir das letzte Beispiel und berechnen

$$f(2) = 20, f'(2) = 19, f''(2) = 16, f'''(2) = 6.$$

Es folgt

$$f(x) = 20 + \frac{19}{1!}(x-2) + \frac{16}{2!}(x-2)^2 + \frac{6}{3!}(x-2)^3 = 20 + 19(x-2) + 8(x-2)^2 + (x-2)^3.$$

□

5.4 Das Taylorpolynom

- Bessere als nur lineare Approximation liefert der Satz von Taylor.

Wir verallgemeinern Beispiel 112. Ist f eine n -mal differenzierbare Funktion, aber kein Polynom, so können wir die rechte Seite von (30) hinschreiben und erhalten ein Polynom, das an der Stelle x_0 denselben Wert wie f , dieselbe Ableitung, zweite Ableitung, ..., n -te Ableitung hat wie f selbst. Aber natürlich ist dieses Polynom nicht gleich f , wenn f kein Polynom war. Die Differenz zwischen beiden sei R_n :

$$f(x) = \underbrace{\sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k}_{\text{Polynom!}} + R_n(x).$$

Man nennt das Polynom das n -te Taylorpolynom von f im Entwicklungspunkt x_0 . Es schmiegt sich an der Stelle x_0 also so gut an f an, dass alle Ableitungen bis zur n -ten Ordnung übereinstimmen.

Für $n = 0$ hat man ganz bescheiden

$$f(x) = f(x_0) + R_0(x),$$

und wir wissen aus dem Mittelwertsatz, dass

$$R_0(x) = f'(\xi)(x - x_0)$$

für ein ξ zwischen x und x_0 , falls f differenzierbar ist.

Für $n = 1$ hat man

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + R_1(x),$$

und aus der Definition der Differenzierbarkeit wissen wir, dass

$$\lim_{x \rightarrow x_0} \frac{R_1(x)}{x - x_0} = 0.$$

Der Satz von Taylor sagt nun etwas über den Fehler R_n bei der Approximation durch Taylorpolynome höherer Ordnung n .

Satz 114 (Taylorformel). Sei $f : I \rightarrow \mathbb{R}$ n -mal differenzierbar im Intervall I , und sei $x_0 \in I$.

Dann gilt

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(x_0)}{k!} (x - x_0)^k + R_n(x),$$

mit

$$\lim_{x \rightarrow x_0} \frac{R_n(x)}{(x - x_0)^n} = 0. \quad (31)$$

Ist f sogar $(n + 1)$ -mal differenzierbar, so kann man das Restglied auch schreiben als

$$R_n(x) = \frac{f^{(n+1)}(\xi)}{(n + 1)!} (x - x_0)^{n+1} \quad (32)$$

mit einem ξ zwischen x und x_0 . Das Restglied in dieser Darstellung nennt man auch das Lagrangesche Restglied. Es sieht also fast so aus, wie der nächste Term der Summe, nur das Argument der Ableitung ist ξ und nicht x_0 .

Der Beweis benutzt den verallgemeinerten Mittelwertsatz 231, soll hier aber nicht gegeben werden.

Der Satz zeigt, dass der Fehler $R_n(x)$ für $x \rightarrow x_0$ sehr schnell gegen 0 geht, d.h. dass sich das Taylorpolynom wirklich sehr gut an die Funktion anschmiegt. Natürlich um so besser, je größer das n ist.

Eine wichtige Anwendung der Taylorformel, die heute allerdings in der Software der Computer versteckt ist, ist die Berechnung von Funktionswerten:

Das Taylorpolynom gibt die Möglichkeit, eine (mehrfach differenzierbare) Funktion in der Umgebung eines Entwicklungspunktes durch ein Polynom zu approximieren und den Approximationsfehler qualitativ mit (31) oder quantitativ mit (32) abzuschätzen.

Beispiel 115 (Schwingungen). In der Elektrotechnik wie in der Mechanik, Akustik, Hydrodynamik usw. wird die Grundform eines Schwingungsvorganges $f(t)$ charakterisiert durch die Bedingung

$$\ddot{f}(t) = -af(t). \quad (33)$$

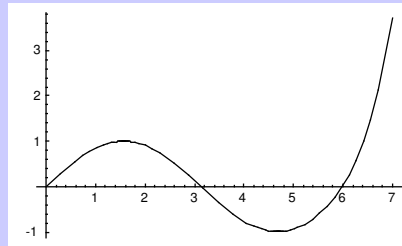
Wenn wir der Einfachheit halber annehmen, dass $a = 1$ und $f(0) = 0$ ist und sich das System zur Zeit $t = 0$ mit der Geschwindigkeit $\dot{f}(0) = 1$ bewegt, wie sieht dann der Schwingungsvorgang aus? Wir schreiben das Taylorpolynom von $f(t)$ im Punkt $t = 0$ auf. Die dafür benötigten Ableitungen bekommt man durch Differentiation von (33) sehr leicht:

$$\begin{aligned} f(0) &= 0, & \dot{f}(0) &= 1 \\ \ddot{f}(0) &= -f(0) = 0, & \overset{(3)}{f}^{(3)}(0) &= -\dot{f}(0) = -1 \\ &\dots \end{aligned}$$

Alle geraden Ableitungen in 0 sind 0, die ungeraden abwechselnd 1 und -1. Es ergibt sich mit $n = 2m + 1$

$$f(t) = t - \frac{t^3}{3!} + \frac{t^5}{5!} - \dots \pm \frac{t^{2m+1}}{(2m+1)!} + R_n(t).$$

Wir haben im Augenblick nicht genug Information bereit, um das Restglied abzuschätzen. Wenn wir das Polynom für $m = 6$ mit Mathematika „plotten“, stellen wir fest, dass über ein großes Intervall tatsächlich das Bild einer Sinusschwingung entsteht. Aber wie bei jedem Polynom verabschiedet sich der Graph schließlich und geht gegen $\pm\infty$.



Beispiel 116 (Diskretisierung). Aus der Taylorformel folgt, wenn f'' stetig und h „klein“ ist,

$$\begin{aligned} f(x+h) &\approx f(x) + f'(x)h + \frac{1}{2}f''(x)h^2, \\ f(x-h) &\approx f(x) - f'(x)h + \frac{1}{2}f''(x)h^2. \end{aligned}$$

Addition der beiden Gleichungen liefert

$$f''(x) \approx \frac{f(x+h) + f(x-h) - 2f(x)}{h^2}. \quad (34)$$

Und natürlich hat man

$$f'(x) \approx \frac{f(x+h) - f(x-h)}{2h}. \quad (35)$$

Viele physikalische Gesetze sind durch Differentialgleichungen gegeben, und viele dieser Differentialgleichungen lassen sich nicht explizit lösen, so dass man auf numerische Lösungsverfahren angewiesen ist. Ein wichtiges Hilfsmittel dabei ist die Diskretisierung des Problems. Die gesuchte Funktion $y(x)$ wird dabei durch eine (diskrete) Zahlenfolge y_k ersetzt, die die Funktionswerte an den Stellen

$$x_k = x_0 + kh, k = 0, 1, 2, \dots$$

approximieren soll. Dabei ist h die sogenannte *Diskretisierungskonstante*, und man hofft, für sehr kleines h eine gute Approximation $y_k \approx y(x_k)$ zu bekommen.

Um eine Differentialgleichung 2. Ordnung zu diskretisieren, also in eine sogenannte *Differenzengleichung* für die y_k umzuschreiben, benutzt man die Approximationen (34) und (35). Konkret wird aus

$$y'' + 6y' + 9y = 0, \quad y(0) = 0, \quad y'(0) = 0.5$$

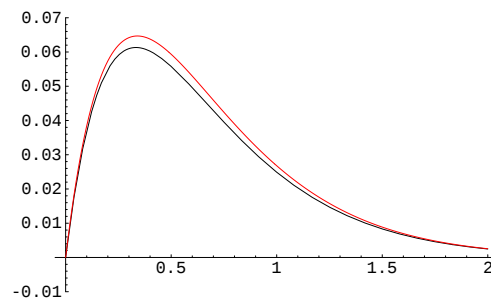
mit $x_0 = 0$ und $x_k = kh$ nach Einsetzen

$$\frac{y_{k+1} + y_{k-1} - 2y_k}{h^2} + 6 \frac{y_{k+1} - y_k}{h} + 9y_k = 0, \quad y_0 = 0, \quad y_1 = 0.5h.$$

Auflösen nach y_{k+1} liefert die Rekursion

$$y_{k+1} = \frac{(2 + 6h - 9h^2)y_k - y_{k-1}}{1 + 6h}, \quad y_0 = 0, \quad y_1 = 0.5h.$$

In diesem Fall kann man die Differentialgleichung allerdings auch elementar lösen und die Lösung mit der Diskretisierung vergleichen. Die folgende Abbildung zeigt die exakte Lösung (untere Kurve) überlagert mit der Diskretisierung für $h = 0.01$.



□

Beispiel 117 (Arbeitspunkt einer Verstärkers). Wir schließen die Reihe der Beispiele mit einer weiteren Anwendung der Taylorformel auf ein elektrotechnisches Problem, nämlich die Wahl des Arbeitspunktes für einen Verstärker, der eine Eingangsspannung u auf die Ausgangsspannung $f(u)$ verstärkt. Man möchte häufig gern, dass die Verstärkung im Arbeitsbereich linear ist.

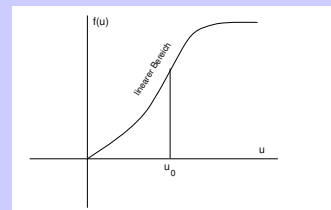
Nach der Taylorformel ist aber

$$f(u) \approx f(u_0) + f'(u_0)(u - u_0) + \frac{1}{2}f''(u_0)(u - u_0)^2$$

mit einem Fehler $\frac{f'''(\xi)}{6}(u - u_0)^3$ von dritter Ordnung in $\Delta u = u - u_0$, den man für kleines $|\Delta u|$ vernachlässigen kann.

Man erreicht also ein Optimum an Linearität, wenn man den Arbeitspunkt u_0 so legt, dass $f''(u_0) = 0$, d.h. in einen Wendepunkt. Dann ist

$$f(u) \approx f(u_0) + f'(u_0)(u - u_0).$$



Beispiel 118 (Relativistische Korrektur). In den meisten „Alltagsanwendungen“ der Taylorapproximation benutzt man die Approximation 1. Ordnung (Das ist eigentlich nur die Definition der Differenzierbarkeit und der Name „Taylor“ in diesem Zusammenhang fast eine Beleidigung des Herrn Taylor.) oder 2. Ordnung, also zum Beispiel (mit $x_0 = 0$)

$$\frac{1}{\sqrt{1-x^2}} \approx 1 + \frac{1}{2}x^2$$

oder

$$\frac{1}{\sqrt{1-x^2}} \approx 1 + \frac{1}{2}x^2 + \frac{3}{8}x^4.$$

(Betrachten Sie zunächst das Taylorpolynom von $(1-x)^{-1/2}$ und setzen Sie dann x^2 statt x ein.)

In der relativistischen Mechanik begegnet man diesem Faktor sehr häufig, wobei $x = v/c$ die mit der Lichtgeschwindigkeit normierte Geschwindigkeit bedeutet. Das m in Einsteins berühmter Formel $E = mc^2$ ist zum Beispiel die relative Masse, die sich aus der „Ruhmasse“ m_0 durch $m = \frac{m_0}{\sqrt{1-(v/c)^2}}$ ergibt. Die relativistische Korrektur ist nach obigem von der Größenordnung $\approx 1 + \frac{1}{2}(v/c)^2$, bei 150 km/h also

$$\approx 1 + \frac{1}{2} \left(\frac{150}{108 \cdot 10^7} \right)^2 = 1 + 0.96 \cdot 10^{-14}.$$

6 Elementare Funktionen II

An diesem Skriptum ist vielfach kritisiert worden, dass es die einfachsten Sachverhalte auf möglichst komplizierte Weise erklärt. Wenn Sie dieses Kapitel überfliegen, finden Sie das bestätigt: Den Sinus kennen Sie aus der Schule und Sie wissen natürlich, dass $\sin(x + 2\pi) = \sin x$ oder $\sin'' x = -\sin x$ ist. Und dann tut das Skriptum so, als wüssten Sie nichts und *definiert* den Sinus so umständlich wie im Satz 119. Komplizierter geht's nicht. Oder vielleicht doch: Da wird die Taylorformel bemüht, um die Sinuskurve zu berechnen, die doch jeder Schüler in seiner Formelsammlung finden oder sogar malen kann. Und $\pi = 3.14\dots$ wird *definiert* als das Doppelte der ersten positive Nullstelle von Cosinus, wozu man meint, erstmal beweisen zu müssen, dass es die überhaupt gibt ...

Solche Kritik geht völlig am entscheidenden Punkt vorbei. Einmal benutzen wir zur Einführung des Sinus nicht Dreiecksverhältnisse wie in der Schule, sondern die Situation, in der er Ihnen in den Ingenieurwissenschaften vor allem begegnet, nämlich Schwingungsvorgänge. Das ist ein ehrenwertes Motiv, aber angesichts der dann folgenden „Formalitäten“ vielleicht doch nicht zu rechtfertigen. Es geht auch um viel mehr: Wenn Sie sich mit der Technik von Mobilfunkantennen beschäftigen oder mit der Schadstoffausbreitung in der Luft, mit der Belastung von Radreifen eines ICE oder mit der Wärmedämmung eines Reaktors, mit unerwünschten Fahrwerksschwingungen oder mit der Steuerung von Regelkreisen, mit der Diffusion in Halbleitern oder mit Skineffekten in der Hochfrequenztechnik, mit der Erdbebensicherheit von Gebäuden oder mit was für Ingenieurproblemen auch immer: Sie haben es mit **Differentialgleichungen** zu tun. Und in aller Regel mit Differentialgleichungen, deren Lösungen Sie in keiner noch so guten Formelsammlung finden, einfach deshalb, weil sich die meisten Differentialgleichungen gar nicht elementar lösen lassen.¹⁴ Dann ist es aber von fundamentaler Wichtigkeit, Wege zu finden, um trotzdem das Verhalten der Lösungen in den Griff zu bekommen. Die einfachsten Beispiele dafür stellen wir hier vor: Trockenübungen gewissermaßen, weil wir eben die Lösungen der elementaren Schwingungs- oder Wachstumsgleichungen schon kennen. „Echte“ Beispiele sind leider *erheblich* viel aufwendiger, und deshalb scheint es vernünftig, hier erst einmal die simpelsten Fälle zu betrachten. Sie sind, wie Sie sehen werden, kompliziert genug.

Nebenbei demonstrieren wir hier auch die Stärken der bisher entwickelten Verfahren und Begriffe (Differentiation, Zwischenwertsatz) an praktischen Beispielen, nämlich gerade an Schwingungs- oder Wachstumsgleichungen.

¹⁴Vergleichen Sie dazu auch die Einleitung von *Regelungstechnik I*.

6.1 Trigonometrische Funktionen I

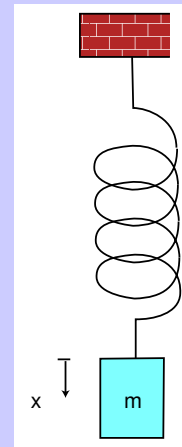
- Den Sinus kann man über Dreiecksverhältnisse oder über seine Potenzreihe erklären.
- Wir untersuchen stattdessen ein einfachstes Schwingungsproblem der Mechanik, das als Lösung den Sinus und Cosinus liefert.
- Wir sehen, dass diese Funktionen durch eine Differentialgleichung mit gewissen Anfangsbedingungen eindeutig festgelegt sind - mit allen ihren Eigenschaften.

Vergleichen Sie zum folgenden auch die Mechanik-Skripte, zum Beispiel *Müller: Mechanik II, Abschnitt 15*.

Wir betrachten wie früher die Federschwingung der Masse m mit Federkonstante k . Mit $x(t)$ bezeichnen wir die Auslenkung aus der Gleichgewichtslage zur Zeit t und mit einem Punkt die Ableitung nach t . Wie im Abschnitt über die komplexen Zahlen erklärt, führt das auf die Schwingungsgleichung

$$m\ddot{x} + kx = 0. \quad (36)$$

Aus physikalischen Gründen ist uns klar, dass diese Bewegungsgleichung eine Lösung hat. Genauer hat sie sehr viele Lösungen: Erst indem wir zu einem Zeitpunkt, etwa $t = 0$, den Ort $x(0)$ und die Geschwindigkeit $\dot{x}(0)$ festlegen, spezifizieren wir eine eindeutig bestimmte Bewegung.

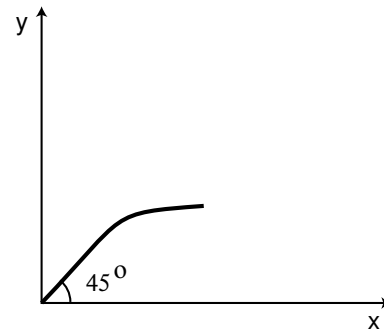


Ziel: Wenn wir unterstellen, dass das Newtonsche Gesetz in Form der Gleichung (36) die Schwingung eines Federpendels beschreibt, wie kann diese Schwingung dann aussehen? Welche Eigenschaften haben die Lösungen der Differentialgleichung (36)?

Zur Vereinfachung wählen wir die Konstanten m und k gleich, so dass wir sie herauskürzen können, ignorieren die physikalischen Dimensionen und bezeichnen die gesuchte Funktion nicht mit $x(t)$, sondern mit $y(x)$, d.h. wir betrachten die Differentialgleichung

$$y'' + y = 0.$$

Wir legen uns auch auf konkrete Anfangsbedingungen $y(0) = 0$ und $y'(0) = 1$ fest. Die Lösung „beginnt“ also ungefähr so, wie nebenstehend skizziert.



Satz 119 (und Definition). *Es gibt eine eindeutig bestimmte zweimal differenzierbare Funktion*

$$y : \mathbb{R} \rightarrow \mathbb{R}$$

mit folgenden Eigenschaften:

$$y'' + y = 0 \tag{37}$$

und

$$y(0) = 0, \quad y'(0) = 1. \tag{38}$$

Diese Funktion nennen wir Sinus und ihre Ableitung Cosinus. Wir bezeichnen sie mit $\sin$ und $\cos$.

Beispiel 120. Zeigen Sie, dass dann für beliebige $a, b \in \mathbb{R}$ die Funktion

$$x(t) := a \cos \sqrt{\frac{k}{m}} t + b \sin \sqrt{\frac{k}{m}} t,$$

die Schwingungsgleichung (37) löst und bestimmen Sie die Anfangswerte $y(0)$ und $\dot{y}(0)$ zur Zeit $t = 0$. $\square$

Die einleitende Motivation ist Physik, aber kein mathematischer Beweis für Satz 119. Wir müssen zweierlei beweisen:

- dass es wirklich solche Funktionen y gibt. (Existenzbeweis)
- dass es nicht mehrere, sondern höchstens *eine* solche Funktion y gibt. (Eindeutigkeitsbeweis)

Muss das sein? Lesen Sie noch einmal die Einleitung zu Kapitel 6. Den Existenzbeweis vertagen wir einstweilen, nämlich bis zum Anfang des folgenden Abschnittes. Wir nehmen also an, dass es Lösungen gibt. Aber die Einzigkeit der Lösung wollen wir beweisen, nicht zuletzt deshalb, weil dabei schon wesentliche Eigenschaften der Lösung zutage treten.

Eindeutigkeitsbeweis zum Satz 119. Sei zunächst y eine Lösung. Dann gilt

$$(y(x)^2 + y'(x)^2)' = 2y(x)y'(x) + 2y'(x)y''(x) \stackrel{(37)}{=} 2y(x)y'(x) - 2y'(x)y(x) = 0.$$

Also ist $y(x)^2 + y'(x)^2 = \text{const}$ und aus (38) folgt

$$y(x)^2 + y'(x)^2 = 1.$$

Damit haben wir aus der Differentialgleichung die Beziehung

$$\sin^2 x + \cos^2 x = 1 \tag{39}$$

hergeleitet.

Sei nun $\tilde{y} : \mathbb{R} \rightarrow \mathbb{R}$ eine weitere Lösung der Gleichung (37):

$$\tilde{y}'' + \tilde{y} = 0. \tag{40}$$

Wir definieren eine Funktion

$$f(x) := y'(x)\tilde{y}(x) - y(x)\tilde{y}'(x). \tag{41}$$

Dann ist

$$\begin{aligned} f'(x) &= y''(x)\tilde{y}(x) + y'(x)\tilde{y}'(x) - y'(x)\tilde{y}'(x) - y(x)\tilde{y}''(x) \\ &\stackrel{(37),(40)}{=} -y(x)\tilde{y}(x) + y(x)\tilde{y}(x) = 0. \end{aligned}$$

Also ist $f(x) = f(0)$ konstant, und Einsetzen der Anfangswerte aus den Gleichungen (38) liefert $f(0) = \tilde{y}(0)$.

Ebenso rechnet man nach, dass

$$g(x) := y(x)\tilde{y}(x) + y'(x)\tilde{y}'(x) \quad (42)$$

konstant ist und zwar $g(x) = \tilde{y}'(0)$.

Damit haben wir ein lineares Gleichungssystem für $\tilde{y}(x)$ und $\tilde{y}'(x)$:

$$\begin{aligned} y'(x)\tilde{y}(x) - y(x)\tilde{y}'(x) &= \tilde{y}(0), \\ y(x)\tilde{y}(x) + y'(x)\tilde{y}'(x) &= \tilde{y}'(0). \end{aligned}$$

Auflösen unter Benutzung von (39) liefert

$$\tilde{y}(x) = \tilde{y}'(0)y(x) + \tilde{y}(0)y'(x). \quad (43)$$

Erfüllt $\tilde{y}$ dieselben Anfangsbedingungen $\tilde{y}(0) = 0$ und $\tilde{y}'(0) = 1$ wie y , so folgt $\tilde{y} = y$. Damit ist gezeigt, dass es höchstens eine Lösung gibt. $\square$

Weitere Eigenschaften der Lösung $y(x)$. Aus der Differentialgleichung folgt

$$(y')' = y'' = -y,$$

d.h.

$$\cos' x = -\sin x.$$

Nun wenden wir die Gleichung (43) noch zweimal auf verschiedene Funktionen $\tilde{y}$ an.

1. Die Funktion $\tilde{y}(x) := y(-x)$ erfüllt (37) mit den Anfangsbedingungen

$$y(0) = 0, y'(0) = -1.$$

Daher folgt aus (43)

$$\sin(-x) = -\sin x,$$

d.h. $\sin$ ist eine sogenannte *ungerade Funktion*. Durch Differenzieren ergibt sich daraus

$$\cos(-x) = \cos x,$$

d.h. $\cos$ ist eine sogenannte *gerade Funktion*.

2. Mit einem festen x_2 definiert man

$$\tilde{y}(x) := y(x + x_2).$$

Auch diese Funktion erfüllt (37) mit den Anfangsbedingungen

$$\tilde{y}(0) = \sin x_2, \quad \tilde{y}'(0) = \cos x_2.$$

Aus (43) folgt

$$\sin(x + x_2) = \sin x \cos x_2 + \cos x \sin x_2 \quad (44)$$

und durch Differenzieren

$$\cos(x + x_2) = \cos x \cos x_2 - \sin x \sin x_2. \quad (45)$$

Mit $x = x_1$ ergeben sich die bekannten(?) Additionstheoreme.

Damit haben wir schon eine Menge Informationen über die Lösung der Schwingungsgleichung, also über den Sinus und die Ableitung Cosinus gesammelt. Wir halten noch fest, dass wegen (39) ihre Werte zwischen -1 und +1 liegen, und dass aus der im Satz postulierten Differenzierbarkeit zusammen mit der Differentialgleichung (37) folgt, dass der Sinus beliebig oft differenzierbar und insbesondere stetig ist.

Periodizität. Was uns noch fehlt, ist die fundamentale Schwingungseigenschaft: die Periodizität. Sinus und Cosinus sind 2π -periodisch. Dazu müssen wir aber zunächst wissen, was π eigentlich ist. Dafür zeigen wir, dass der Cosinus $y'(x)$ eine kleinste positive Nullstelle bei einem Wert x_0 hat. Dann werden wir definieren $\pi = 2x_0$ und zeigen, dass Sinus und Cosinus 2π -periodisch sind.

Satz 121 (und Definition). Die Funktion $\cos : \mathbb{R} \rightarrow \mathbb{R}$ hat eine kleinste positive Nullstelle, die wir mit $\frac{\pi}{2}$ bezeichnen.

Beweis. Zunächst zeigen wir, dass der Cosinus überhaupt eine Nullstelle hat. Weil

$$\cos 0 = y'(0) = 1 > 0$$

genügt es nach dem Zwischenwertsatz zu zeigen, dass der Cosinus in $]0, +\infty[$ auch negative Werte annimmt. Andernfalls wäre er überall positiv, und wir wollen zeigen, dass diese Annahme zu einem Widerspruch führt.

Die Ableitungen des Cosinus kennen wir schon

$$\cos 0 = 1, \cos' 0 = -\sin 0 = 0, \cos'' 0 = -\cos 0 = -1, \cos''' 0 = \sin 0 = 0.$$

Nach dem Satz von Taylor ist also

$$\cos x = 1 + \frac{0}{1!}x - \frac{x^2}{2!} + \frac{0}{3!}x^3 + \frac{\cos^{(4)} \xi}{4!}x^4$$

Weil $|\cos^{(4)} \xi| = |\cos \xi| \leq 1$ für alle ξ , ist

$$\cos 2 \leq 1 - \frac{2^2}{2} + \frac{2^4}{4!} = -\frac{1}{3} < 0.$$

Damit hat der Cosinus in $]0, +3[$ jedenfalls Nullstellen. Weil er stetig und $\cos 0 = 1$ ist, kann man leicht zeigen, dass es unter diesen eine kleinste gibt. Wir führen das aber nicht aus. $\square$

Nun geht alles ganz einfach. Auf $[0, \frac{\pi}{2}[$ ist der Cosinus positiv und der Sinus deshalb monoton wachsend. Insbesondere ist $\sin \frac{\pi}{2} > 0$ und aus (39) folgt

$$\sin \frac{\pi}{2} = 1$$

Die Additionstheoreme (44), (45) liefern dann sofort für alle $x \in \mathbb{R}$

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

Also hat man auch die 2π -Periodizität hergeleitet.

Schlußbemerkung. Die Einführung der trigonometrischen Funktionen über ihre Differentialgleichung folgt dem für Mathematiker geschriebenen Buch *Analysis* von *Serge Lang*. Wir haben eingangs erläutert, warum sie sich für Ingenieure ebenfalls besonders empfiehlt.

Natürlich erscheinen die obigen Überlegungen sehr trickreich und an der speziellen Differentialgleichung zu hängen. Das ist teilweise wahr, aber Sie sollen ja lernen zu fragen: „Welche Informationen habe ich in der gegebenen speziellen Situation und wie kann ich sie benutzen?“ Teilweise stehen dahinter aber auch allgemeine Prinzipien aus der Theorie der Differentialgleichungen bzw. der Physik. So schreibt sich (39) für Lösungen der Schwingungsgleichung (36) als

$$kx^2 + m\dot{x}^2 = \text{const},$$

und das ist der Energieerhaltungssatz. Erhaltungsgrößen, der Mathematiker sagt auch *Integrale* liefern Informationen zum besseren Verständnis der Lösungen, auch in Fällen, wo man die letzteren nicht explizit angeben kann.

Wir stellen die in diesem Abschnitt gefundenen Eigenschaften von Sinus und Cosinus noch einmal zusammen.

Satz 122 (Eigenschaften von Sinus und Cosinus). *Die Funktionen $\sin : \mathbb{R} \rightarrow \mathbb{R}$ und $\cos : \mathbb{R} \rightarrow \mathbb{R}$ sind 2π -periodisch, und es gilt für alle $x, x_1, x_2 \in \mathbb{R}$:*

$$\begin{aligned}\sin' x &= \cos x, & \cos' x &= -\sin x \\ \sin^2 x + \cos^2 x &= 1 \\ \sin(-x) &= -\sin x, & \cos(-x) &= \cos x \\ \sin\left(x + \frac{\pi}{2}\right) &= \cos x, & \cos\left(x + \frac{\pi}{2}\right) &= -\sin x \\ \sin(x_1 + x_2) &= \sin x_1 \cos x_2 + \cos x_1 \sin x_2 \\ \cos(x_1 + x_2) &= \cos x_1 \cos x_2 - \sin x_1 \sin x_2.\end{aligned}$$

6.2 Trigonometrische Funktionen II

- Wir tragen weitere Eigenschaften von Sinus und Cosinus zusammen und widmen uns kurz dem Tangens und den Umkehrfunktionen der trigonometrischen Funktionen.

SPEZIELLES ZU SINUS UND COSINUS

Die Graphen und der Existenzbeweis. Für die numerische Berechnung von Funktionswerten benutzen wir die Taylorapproximation, vergleichen Sie Beispiel 115. Weil Sinus und Cosinus sich beim Differenzieren bis aufs Vorzeichen gegenseitig reproduzieren und in 0 den Wert 0 bzw. 1 haben, ist das Taylorpolynom für $\sin x$ im Nullpunkt leicht aufzuschreiben. Die Terme mit geraden x -Potenzen fallen weg, und man erhält

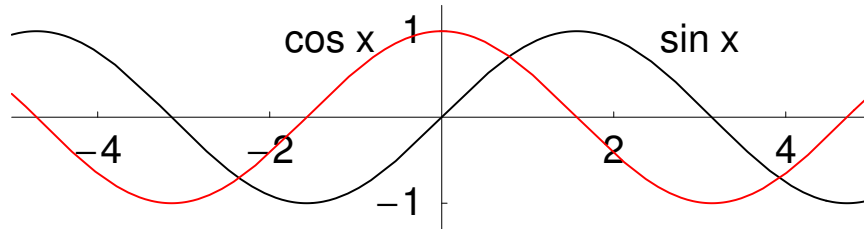
$$\sin x = \sum_{k=0}^m \frac{(-1)^k}{(2k+1)!} x^{2k+1} + R_{2m+2}(x).$$

Das Restglied ist erst einmal natürlich R_{2m+1} , aber weil der Koeffizient von x^{2m+2} im Taylorpolynom 0 ist, gilt die Formel auch mit dem Restglied R_{2m+2} , und das liefert in der Regel schärfere Abschätzungen:

$$|R_{2m+2}(x)| = \left| \frac{\sin^{(2m+3)} \xi}{(2m+3)!} x^{2m+3} \right| \leq \frac{|x|^{2m+3}}{(2m+3)!}.$$

Für $m = 6$, also das Taylorpolynom der Ordnung 13, und $|x| < 4$ ist der Fehler $< 0,02$. Eine genauere numerische Berechnung der Funktionswerte beschränkt sich auf kleinere x -Werte und benutzt die Periodizitätseigenschaften.

Mit einer entsprechenden Rechnung für den Cosinus ergeben sich folgende Graphen und damit auch eine approximative Berechnung der Zahl $\pi \approx 3.14$.



Wir haben nun so viele Eigenschaften vom Sinus hergeleitet, dass wir sehr genau wissen, wie er aussehen muss – wenn es ihn denn gibt. Weil nämlich nach Beispiel 53 das Restglied für $m \rightarrow \infty$ gegen 0 geht, muss gelten:

$$\sin x = \lim_{m \rightarrow \infty} \sum_{k=0}^m \frac{(-1)^k}{(2k+1)!} x^{2k+1}.$$

Wir überlassen den Mathematikern die Mühe zu zeigen, dass dieser Limes tatsächlich für jedes $x \in \mathbb{R}$ existiert und eine differenzierbare Funktion definiert, die die Bedingungen von Satz 119 erfüllt.

Spezielle Werte. Unabhängig von der numerischen Berechnung lassen sich spezielle Werte von $\sin$ und $\cos$ auch exakt bestimmen.

Aus (45) und (39) folgt

$$\cos(2x) = 2 \cos^2 x - 1$$

und daraus

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

also

$$\cos\frac{\pi}{4} = \frac{1}{\sqrt{2}}, \quad \sin\frac{\pi}{4} = \sqrt{1 - \frac{1}{2}} = \frac{1}{\sqrt{2}}.$$

Aus (45) und (44) erhält man ferner

$$\cos(3x) = \cos x(1 - 4\sin^2 x)$$

und mit $x = \pi/6$

$$0 = \cos\frac{\pi}{6}(1 - 4\sin^2\frac{\pi}{6}).$$

Daher ist

$$\sin\frac{\pi}{6} = \frac{1}{2}, \quad \cos\frac{\pi}{6} = \sqrt{1 - 1/4} = \frac{1}{2}\sqrt{3}.$$

Mit $\cos(\frac{\pi}{2} - x) = \sin x$ erhalten wir schließlich folgende spezielle Werte für die Funktionen Sinus und Cosinus:

x	0	$\pi/6$	$\pi/4$	$\pi/3$	$\pi/2$
	0°	30°	45°	60°	90°
$\sin x$	$\sqrt{0}/2$	$\sqrt{1}/2$	$\sqrt{2}/2$	$\sqrt{3}/2$	$\sqrt{4}/2$
$\cos x$	$\sqrt{4}/2$	$\sqrt{3}/2$	$\sqrt{2}/2$	$\sqrt{1}/2$	$\sqrt{0}/2$

Wir stellen noch einmal wichtige Identitäten in einem Satz zusammen. Teilweise sind sie oben bewiesen, die anderen seien dem Leser zur Übung empfohlen.

Satz 123 (Trigonometrische Identitäten).

$$\cos^2 x + \sin^2 x = 1$$

$$\cos 2x = \cos^2 x - \sin^2 x = 2\cos^2 x - 1$$

$$\sin 2x = 2\cos x \sin x$$

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

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

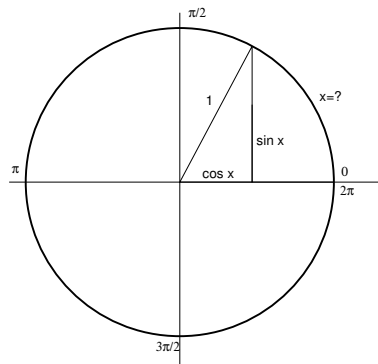
$$\cos x - \cos y = -2\sin\frac{x+y}{2}\sin\frac{x-y}{2}$$

Beziehungen zu Kreis und Dreieck. Die folgende Überlegung erklärt, warum Sinus und Cosinus auch als *Kreisfunktionen* oder *trigonometrische* (=Dreiecks-)Funktionen bezeichnet werden.

Wir lassen x von 0 an wachsen und betrachten den Punkt mit den Koordinaten $(\cos x, \sin x)$ in der Ebene. Nach (39) ist $\cos^2 x + \sin^2 x = 1$, d.h. dieser Punkt erfüllt die Kreisgleichung

$$x_1^2 + x_2^2 = r^2$$

mit $r = 1$, er liegt auf dem sogenannten *Einheitskreis*. Und wegen der Stetigkeit der Funktionen wird von 0 bis 2π jeder Punkt auf dem Kreis erfaßt.



Allerdings bleibt es offen, ob x dabei wirklich die Länge des Winkelsegments auf dem Kreisbogen ist. Das stimmt, aber wir können es erst im Rahmen der Integralrechnung beweisen.

Beispiel 124 (Harmonische Schwingungen). Eine wichtige Anwendung der vorstehenden Überlegung ist die folgende: Eine Funktion

$$f(t) = a \cos \omega t + b \sin \omega t \quad (46)$$

ist die Überlagerung zweier Schwingungen mit Frequenz ω . Ist $a^2 + b^2 > 0$, so kann man das schreiben als

$$f(t) = \sqrt{a^2 + b^2} \left(\frac{a}{\sqrt{a^2 + b^2}} \cos \omega t + \frac{b}{\sqrt{a^2 + b^2}} \sin \omega t \right),$$

und weil

$$\left(\frac{a}{\sqrt{a^2 + b^2}} \right)^2 + \left(\frac{b}{\sqrt{a^2 + b^2}} \right)^2 = 1,$$

gibt es ein ϕ mit

$$\cos \phi = \frac{a}{\sqrt{a^2 + b^2}}, \quad \sin \phi = \frac{b}{\sqrt{a^2 + b^2}}.$$

Also ist

$$\begin{aligned} f(t) &= \sqrt{a^2 + b^2} (\cos \phi \cos \omega t + \sin \phi \sin \omega t) \\ &= \sqrt{a^2 + b^2} \cos(\omega t - \phi). \end{aligned}$$

Jede Schwingung der Form (46) läßt sich als reine Cosinusschwingung mit einer Phasenverschiebung schreiben.

TANGENS UND COTANGENS

Zunächst definieren wir zwei weitere trigonometrische Funktionen: Wir definieren

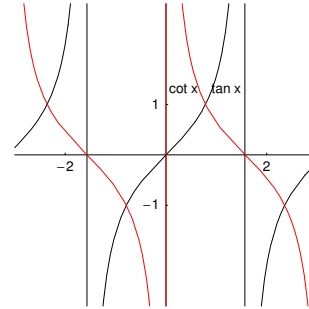
$$\tan x := \frac{\sin x}{\cos x}, \quad \cot x := \frac{\cos x}{\sin x}.$$

Diese Funktionen sind definiert für $x \neq (k + \frac{1}{2})\pi, k \in \mathbb{Z}$ bzw. für $x \neq k\pi, k \in \mathbb{Z}$. Sie sind in ihrem Definitionsbereich stetig und π -periodisch, weil Zähler und Nenner gerade ihr Vorzeichen wechseln, wenn man x durch $x + \pi$ ersetzt.

Aus den Rechenregeln für Sinus und Cosinus folgen zum Beispiel leicht die Formeln

$$\tan(x+y) = \frac{\tan x + \tan y}{1 - \tan x \tan y}$$

$$1 + \tan^2 x = \frac{1}{\cos^2 x}$$



Beispiel 125 (Schraube).

Müller: *Mechanik II*, Abschnitt 8.3.2

Beim Eindrehen einer Schraube vom Radius r und Schraubenwinkel α und Reibungskoeffizient μ gegen eine eingeprägte Kraft F braucht man nach der Mechanik ein Drehmoment von mindestens

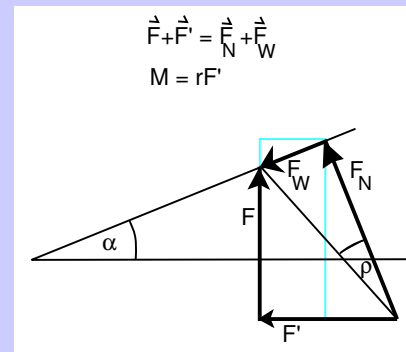
$$M = r \frac{\sin \alpha + \mu \cos \alpha}{\cos \alpha - \mu \sin \alpha} F. \quad (47)$$

um die Haftreibung zu überwinden.

Der Reibungskoeffizient μ ist definiert als Verhältnis der Haftreibungskraft F_W zur verursachenden Normalkraft F_N . Er läßt sich durch den Reibungswinkel ρ mit

$$\mu = \tan \rho = \frac{F_W}{F_N}$$

veranschaulichen.



Dann kann man (47) mit $\cos \alpha$ kürzen und erhält

$$M = r \frac{\tan \alpha + \tan \rho}{1 - \tan \rho \tan \alpha} F = r \tan(\alpha + \rho) F. \quad (48)$$

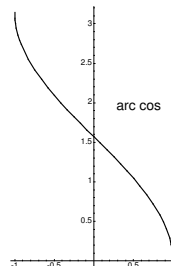
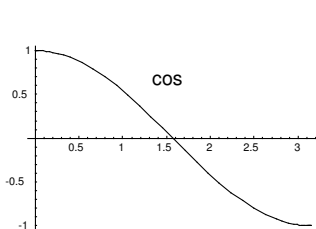
Die Ableitungen von Tangens und Cotangens ergeben sich mit der Quotientenregel:

$$\tan' x = \frac{1}{\cos^2 x} = 1 + \tan^2 x, \quad \cot' x = \frac{-1}{\sin^2 x} = -(1 + \cot^2 x).$$

ARCUSFUNKTIONEN

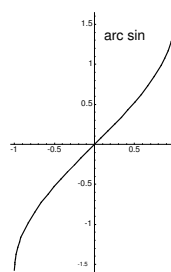
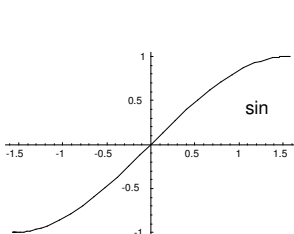
Keine der Funktionen $\cos x, \sin x, \tan x, \cot x$ ist injektiv, sie sind vielmehr „im Gegenteil“ alle periodisch. Aber wir können sie auf Teilintervalle einschränken, wo sie injektiv sind, und dann gibt es zu den so eingeschränkten Funktionen eine Umkehrfunktion. Diese Funktionen nennt man *Arcus-Funktionen* (=Bogenfunktionen), weil sie zu einem gegebenen Wert (z.B. $y = \cos x$) die zugehörige *Bogenlänge* x liefern.

Konkreter: Auf dem Intervall $[0, \pi]$ ist der Cosinus streng monoton fallend und besitzt eine Umkehrfunktion $\arccos : [-1, 1] \rightarrow [0, \pi]$, Arcuscosinus genannt:

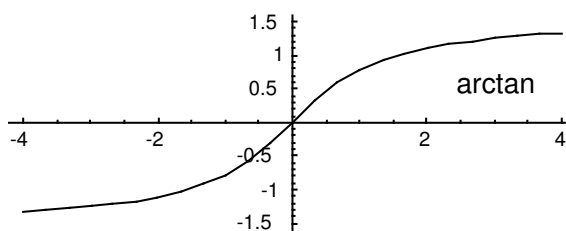
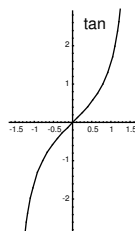


Die Auswahl des Intervalls $[0, \pi]$ ist willkürlich, man kann z.B. auch das Intervall $[\pi, 2\pi]$ oder allgemein $[k\pi, (k+1)\pi]$ nehmen und erhält Umkehrfunktionen, deren Werte jeweils in diesem Intervall liegen. Die anfangs definierte Funktion nennt man auch den *Hauptwert* des Arcuscosinus.

Entsprechend definiert man den Hauptwert des arcsin : $[-1, 1] \rightarrow [-\pi/2, \pi/2]$ (Arcussinus)



und des arctan : $\mathbb{R} \rightarrow]-\pi/2, \pi/2[$ (Arcustangens).



Mit dem Satz 99 erhält man die Differenzierbarkeit der Arcusfunktionen im Inneren ihres Definitionsbereiches und die folgenden Formeln:

Satz 126 (Ableitungen der Arcusfunktionen).

$$\arcsin' x = \frac{1}{\sqrt{1-x^2}}, \quad \arccos' x = \frac{1}{-\sqrt{1-x^2}}, \quad \arctan' x = \frac{1}{1+x^2}.$$

Zum Beispiel ist für $-1 < x < 1$

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

Dabei haben wir benutzt, dass die Werte $\arcsin x$ im Intervall $]-\frac{\pi}{2}, \frac{\pi}{2}[$ liegen, auf dem der cos positiv ist. Darum ist $\cos(\arcsin x) = +\sqrt{1-\sin^2(\arcsin x)}$.

6.3 Exponentialfunktion und Logarithmus

- Wie die trigonometrischen Funktionen, entwickeln wir auch die Exponentialfunktion aus ihrer charakteristischen Differentialgleichung, die für Wachstums- und Zerfallsprozesse zuständig ist.
- Ihre Umkehrfunktion, der Logarithmus war früher von eminenter Wichtigkeit als Rechenhilfe zum Beispiel in der Navigation und Astronomie.
- Heute ist sie vor allem wichtig, weil ihre so einfache Ableitung oft vorkommt.

Vorbemerkung. Wir haben schon wiederholt von der Exponentialfunktion Gebrauch gemacht, insbesondere von ihrer Eigenschaft $\frac{d}{dx}e^x = e^x$. Wir haben sogar im Kapitel über den Nutzen komplexer Zahlen die komplexe Exponentialfunktion e^{x+iy} betrachtet. Jetzt wollen wir eine exakte Definition für diese Funktionen geben.

Dieses scheinbare Chaos rührt daher, dass ein mathematisch sauber strukturierter Aufbau der Analysis viel zu theoretisch wäre, und die für Sie interessanten Dinge alle erst am Schluß kämen, zum Beispiel in der Vorlesung über Differentialgleichungen im 3. Semester. Deshalb greifen wir *in den Beispielen(!)* dieses Skriptums immer wieder auf Schulwissen zurück, das wir dann bei Gelegenheit auch systematisch-theoretisch einbauen. Dieser „Einbau“ soll nicht nur ein Tribut an die mathematische Logik sein, sondern insbesondere der Wiederholung und Festigung von aus der Schule (oft nur schemenhaft) bekannten Sachverhalten dienen.

DIE REELLE EXPONENTIALFUNKTION

Der Zerfall radioaktiver Elemente erfolgt so, dass die Zahl der in der Zeiteinheit zerfallenden Atome proportional ist zur Zahl der vorhandenen Atome des Elementes. Dieses Verhalten ist typisch für eine Vielzahl von Vorgängen: Wachstums-, Sättigungs-, Entladungs-, Verzinsungsvorgänge usw. folgen wenigstens in gewissen Bereichen einem solchen Gesetz: $\dot{y} = cy$.

In völliger Analogie zu unserem Vorgehen bei den trigonometrischen Funktionen benutzen wir auch hier einen Existenz- und Eindeutigkeitssatz für solche Vorgänge zur Definition der Exponentialfunktion.

Satz 127 (und Definition). *Es gibt genau eine differenzierbare Funktion $y : \mathbb{R} \rightarrow \mathbb{R}$ mit*

$$y'(x) = y(x) \quad \text{für alle } x \in \mathbb{R} \quad (49)$$

und

$$y(0) = 1.$$

Diese Funktion heißt die Exponentialfunktion

$$\exp : \mathbb{R} \rightarrow \mathbb{R}.$$

Der Wert

$$e := \exp 1$$

heißt die Eulersche Zahl. Wenn wir demnächst die allgemeinen Potenzen definiert haben, ergibt sich

$$\exp x = e^x.$$

Den Beweis der Einzigkeit für die exp-Funktion haben wir schon im Beispiel 108 gegeben. Ganz ähnlich wie (und noch leichter als) im Fall der trigonometrischen Funktionen bekommen Sie eine Fülle von Eigenschaften für die Exponentialfunktion:

Satz 128 (Exponentialfunktion). Es gilt

$$\begin{aligned}\exp' x &= \exp x, & \exp 0 &= 1, \\ \exp x &> 0 \text{ für alle } x \in \mathbb{R}, \\ \exp(x_1 + x_2) &= \exp x_1 \exp x_2 \text{ für alle } x_1, x_2 \in \mathbb{R}, & (\text{Additionstheorem}) \\ \exp(-x) &= \frac{1}{\exp x}, \\ \exp x &\text{ ist streng monoton wachsend,} \\ \lim_{x \rightarrow -\infty} \exp x &= 0, & \lim_{x \rightarrow +\infty} \exp x &= +\infty.\end{aligned}$$

Beweis als Übung.

Beispiel 129 (Abschätzung für $e = \exp 1$). Wegen $\exp^{(k)} x = \exp x$ liefert die Taylorapproximation in $x_0 = 0$

$$\exp x = 1 + x + \frac{x^2}{2} + \frac{x^3}{6} \exp \xi.$$

Wir wenden das an auf $x = -1$. Dann liegt ξ zwischen -1 und 0 , also ist $|\exp \xi| < 1$ und wegen $\exp x > 0$ ist $0 < \exp \xi < 1$. Wir erhalten

$$\exp(-1) = \frac{1}{2} + R, \quad -\frac{1}{6} < R < 0.$$

Also ist $\frac{1}{3} < \exp(-1) < \frac{1}{2}$. Für die Eulersche Zahl $e := \exp 1$ erhalten wir

$$2 < \exp 1 < 3.$$

□

Beispiel 130. Für beliebiges $n \in \mathbb{N}$ und positives x gilt

$$\exp x = 1 + \frac{x}{1!} + \dots + \frac{x^n}{n!} + \frac{\exp \xi}{(n+1)!} x^{n+1} \geq 1 + \frac{x}{1!} + \dots + \frac{x^n}{n!}.$$

Schließen Sie daraus, dass die Exponentialfunktion „schneller wächst als jede Potenz von x “:

$$\lim_{x \rightarrow \infty} \frac{\exp x}{x^n} = \infty.$$

□

Der Graph und der Existenzbeweis.

Wie im Beispiel 129 ergibt sich das Taylorpolynom

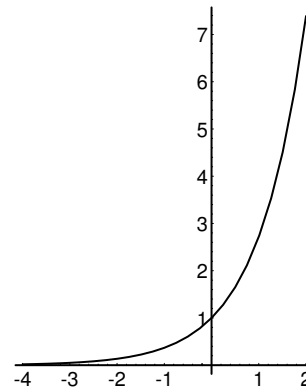
$$\exp x = \sum_{k=0}^n \frac{1}{k!} x^k + R_n(x).$$

Für $a < x < b$ ist

$$|R_n(x)| \leq \exp(b) \frac{|x|^{n+1}}{(n+1)!}.$$

Nach Beispiel 53 geht die rechte Seite für $n \rightarrow \infty$ gegen Null. Ist $m \in \mathbb{N}$ mit $m > b$, so ist

$$\exp b < \exp m = e^m < 3^m.$$



Damit können wir die Funktionswerte mit kontrolliertem Fehler berechnen. Wir finden den skizzierten Graphen.

Die Eulersche Zahl ergibt sich genauer als

$$e = 2.71828 \dots$$

Wie beim Sinus können wir nun auch durch

$$y(x) := \lim_{n \rightarrow \infty} \sum_{k=0}^n \frac{1}{k!} x^k$$

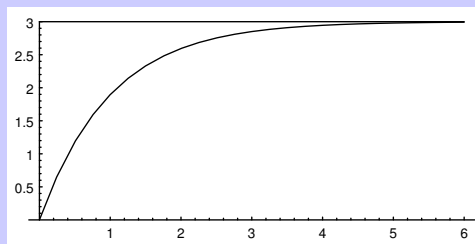
eine Funktion definieren, von der die Mathematiker zeigen können, dass sie differenzierbar ist und die in Satz 127 geforderten Eigenschaften erfüllt.

Beispiel 131 (Sättigung und Entladung). Der Entladungsstrom eines Kondensators ist proportional zur noch vorhandenen Ladung, so dass die Ladung der Gleichung

$$Q(t) = Q_0 \exp(-at)$$

genügt. Ebenso ist die Polarisierung eines Dielektrikums durch ein elektrisches Feld gegeben durch

$$P(t) = P_s(1 - e^{-\frac{t}{\tau}}).$$



Beispiel 132 (Boltzmann-Maxwellsche Energieverteilung).

Werkstoffe II, Abschnitt 1.12

Wir untersuchen den Funktionsverlauf von

$$f(x) = \sqrt{x}e^{-x}, \quad x \geq 0.$$

Diese Funktion tritt in der Form

$$p(W) = A \sqrt{\frac{W}{kT}} e^{-\frac{W}{kT}}, \quad W \geq 0$$

als Boltzmann-Maxwellsche Energieverteilung in der kinetischen Gastheorie auf. Sie gibt die Wahrscheinlichkeitsverteilung für die Geschwindigkeit der Gasmoleküle in Abhängigkeit von der Energie W an.

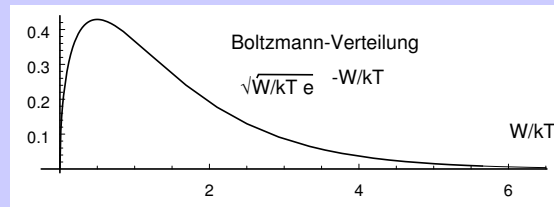
Klar ist, dass

$$\begin{aligned} f(x) &\geq 0, \\ f(x) &= 0 \iff x = 0, \\ \lim_{x \rightarrow \infty} f(x) &= 0. \end{aligned}$$

Es ist

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

Das ist positiv in $]0, \frac{1}{2}[$, hat eine Nullstelle für $x = \frac{1}{2}$ und ist negativ für $x > \frac{1}{2}$. Also ist f auf $[0, \frac{1}{2}]$ streng monoton steigend, hat bei $x = \frac{1}{2}$ ein Maximum und fällt dann monoton gegen 0 ab.



DER LOGARITHMUS

Definition 133 (Logarithmus). Die Exponentialfunktion $\exp : \mathbb{R} \rightarrow]0, \infty[$ ist streng monoton wachsend, also injektiv. Weil sie stetig ist und

$$\lim_{x \rightarrow -\infty} \exp x = 0 \text{ und } \lim_{x \rightarrow +\infty} \exp x = +\infty$$

gilt, ist nach dem Zwischenwertsatz

$$\exp(\mathbb{R}) =]0, \infty[.$$

Die Umkehrabbildung

$$\ln :]0, \infty[\rightarrow \mathbb{R}$$

heißt *der (natürliche) Logarithmus*.

Nach Satz 99 ist der Logarithmus ebenfalls differenzierbar mit der Ableitung

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

Aus den Eigenschaften der Exponentialfunktion folgen dann Eigenschaften des Logarithmus, zum Beispiel

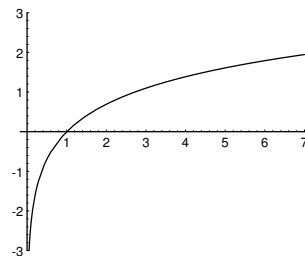
$$\lim_{x \rightarrow +\infty} \ln x = +\infty, \quad \lim_{x \searrow 0} \ln x = -\infty$$

und die fundamentale Gleichung

$$\ln(xy) = \ln x + \ln y.$$

Die letztere Gleichung (Reduktion der Multiplikation auf die Addition) hat den Logarithmen vor der Erfindung elektronischer Rechenmaschinen eine prominente Bedeutung bei allen schwierigeren praktischen Rechenaufgaben (z.B. in der Astronomie oder Navigation) verschafft.

Den Graphen der Umkehrfunktion erhält man durch Spiegelung des Exponentialgraphen an der Winkelhalbierenden.



Beispiel 134. Berechnen Sie die Ableitungen der Funktion $f(x) = \ln(1+x)$ in $x_0 = 0$ und zeigen Sie dadurch die Taylorformel

$$\ln(1+x) = \frac{x}{1} - \frac{x^2}{2} + \frac{x^3}{3} + \dots + (-1)^{n-1} \frac{x^n}{n} + R_n(x)$$

Schreibt man diese Formel für $-x$ statt x und subtrahiert, so ergibt sich

$$\ln \frac{1+x}{1-x} = 2 \left(\frac{x}{1} + \frac{x^3}{3} + \dots + \frac{x^{2m+1}}{2m+1} \right) + \text{Restterm}.$$

□

DIE KOMPLEXE EXPONENTIALFUNKTION

Wir erinnern daran, dass wir im Abschnitt über die komplexen Zahlen im Vorgriff schon zu folgender Definition „gedrängt“ worden sind:

Definition 135 (Komplexe Exponentialfunktion, Eulerrelation). Für reelles x sei

$$\exp(ix) := e^{ix} := \cos x + i \sin x.$$

Für reelles x und y sei

$$e^{x+iy} := e^x e^{iy} := e^x (\cos y + i \sin y).$$

Damit ist also die Exponentialfunktion erweitert worden zu einer komplexen Funktion

$$\exp : \mathbb{C} \rightarrow \mathbb{C}.$$

Sie erfüllt ebenfalls das Additionstheorem mit beliebigen komplexen Argumenten, aber sie ist zum Beispiel nicht mehr injektiv, weil

$$e^{2\pi i} = \cos 2\pi + i \sin 2\pi = 1 = e^0.$$

Deshalb gibt es keine Umkehrfunktion. Man kann sich behelfen, indem man wie bei den Arcusfunktionen Umkehrfunktionen der eingeschränkten exp-Funktion definiert. Oder man akzeptiert, dass die Umkehrfunktion „mehrwertig“ und damit eigentlich keine Funktion im Sinne unserer Konventionen ist.

Definition 136 (Komplexer Logarithmus.). Für $z = |z|e^{i\phi} \neq 0$ definiert man

$$\log z := \ln |z| + i\phi, \quad \phi = \arg z.$$

Das ist mehrdeutig, ϕ ist nur bis auf ganzzahlige Vielfache von 2π definiert. Um Eindeutigkeit zu erreichen, fordert man etwas willkürlich $-\pi < \phi \leq \pi$ oder $0 \leq \phi < 2\pi$.

6.4 Allgemeine Potenz und Hyperbelfunktionen

- Es ist leicht zu erklären, was 3^{16} ist: 3 als Faktor 16 mal mit sich multipliziert. Aber was ist e^π ?
- Wir lernen, wie man mit Logarithmus und Exponentialfunktion auch Potenzen mit nicht ganzzahligem Exponenten definieren kann, so dass die üblichen Rechenregeln bestehen bleiben.
- Weil $\cos^2 + \sin^2 = 1$ ist, liegen die Punkte $(\cos x, \sin x)$ auf dem Einheitskreis und der Zusammenhang von trigonometrischen Funktionen und Kreis ist oft wichtig. Wir lernen die *Hyperbelfunktionen* kennen, die eine analoge Beziehung zur Hyperbel $x^2 - y^2 = 1$ haben.

DIE ALLGEMEINE POTENZ

Wir definieren a^x für reelles $a > 0$ und $x \in \mathbb{R}$ durch

$$a^x = \exp(x \ln a).$$

Rechenregeln für die allgemeine Potenz. Diese ergeben sich sofort aus denen für $\exp$ und $\ln$.

$$\begin{aligned} a^{x+y} &= a^x a^y, \\ a^{-x} &= \frac{1}{a^x}, \\ (a^x)^y &= a^{xy}, \\ a^0 &= 1. \end{aligned}$$

Aus der ersten Identität folgt durch vollständige Induktion für natürliches n :

$$a^n = \underbrace{a \cdot \dots \cdot a}_n.$$

Das heißt, die Definition von a^x stimmt für natürliches n mit der alten Definition überein.

Weiter ergibt sich

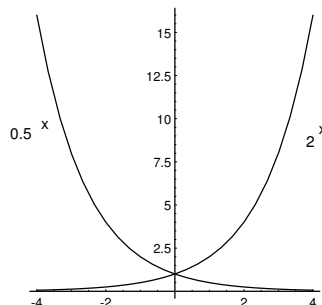
$$(a^{\frac{1}{n}})^n = \exp(n \ln(a^{\frac{1}{n}})) = \exp(n \ln(\exp(\frac{1}{n} \ln a))) = \exp(\ln a) = a.$$

Also ist $a^{\frac{1}{n}} = \sqrt[n]{a}$.

Aus $e = \exp 1$ folgt weiter

$$e^x = \exp(x \ln(\exp 1)) = \exp x.$$

Man erhält



Nach der Kettenregel sind die Potenzfunktionen differenzierbar, und es gilt (nachrechnen!)

$$\frac{dx^b}{dx} = b x^{b-1}, \quad \frac{da^x}{dx} = a^x \ln a.$$

Beispiel 137 (Adiabaten).

Thermodynamik I, Abschnitt 4.13

Die Zustände eines thermodynamischen Systems werden häufig durch die Punkte einer (V, P) -Ebene beschrieben. Kurven in dieser Ebene repräsentieren thermodynamische Prozesse, und die sogenannten *adiabatischen Prozesse* sind durch eine Gleichung

$$PV^\kappa = \text{const.}$$

mit einem *Adiabatenexponenten* $\kappa \in]1, 2[$ charakterisiert.

Beispiel 138. Für $x > -1$ betrachten wir

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

Für die ersten Ableitungen bei $x = 0$ finden wir

$$f'(x) = -\frac{1}{2}(1+x)^{-\frac{3}{2}}, \quad f''(x) = +\frac{3}{4}(1+x)^{-\frac{5}{2}}, \quad f'''(x) = -\frac{15}{8}(1+x)^{-\frac{7}{2}}.$$

Die Taylorsche Formel liefert

$$(1+x)^{-\frac{1}{2}} \approx 1 - \frac{1}{2}x + \frac{3}{8}x^2$$

mit einem Fehler

$$R_2(x) = -\frac{5}{16} \frac{1}{(1+\xi)^{\frac{7}{2}}} x^3.$$

Betrachten wir z.B. nur Werte $0 < x < \frac{1}{10}$, so ist $\xi > 0$ und daher der Fehler höchstens

$$|R_2(x)| < \frac{5}{16} \frac{1}{1^{\frac{7}{2}}} 10^{-3} < 0,0004.$$

Ähnlich findet man für kleines $|x|$

$$\sqrt{1+x} \approx 1 + \frac{1}{2}x - \frac{1}{8}x^2.$$

□

Beispiel 139 (Dipol).

Theoretische Elektrotechnik, Abschnitt 7.5 oder Energie-, Impuls- und Stofftransport, Abschnitt 2.2

Die aus dem obigen Beispiel sich ergebenden Formeln

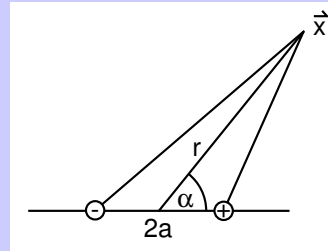
$$\boxed{\sqrt{1+x} \approx 1 + \frac{1}{2}x, \quad \frac{1}{\sqrt{1+x}} \approx 1 - \frac{1}{2}x \text{ für kleines } |x|}$$

findet man häufig in Anwendungen, zum Beispiel bei der approximativen Berechnung des Potentials eines Dipols (Theoretische Elektrotechnik) oder der Isothermen zweier Rohre entgegengesetzter Temperatur (Energie-, Impuls- und Stofftransport).

Dieses ist bis auf eine Konstante gegeben durch

$$V = \frac{1}{\sqrt{r^2 + a^2 - 2ar \cos \alpha}} - \frac{1}{\sqrt{r^2 + a^2 + 2ar \cos \alpha}}.$$

Dabei ist $2a$ die Länge des Dipols, r der Abstand des Aufpunktes vom Zentrum des Dipols und α der Winkel zwischen der Dipolrichtung und dem Aufpunkt.



Man findet für $r \gg a > 0$

$$\begin{aligned} V &= (r^2 + a^2 - 2ar \cos \alpha)^{-1/2} - (r^2 + a^2 + 2ar \cos \alpha)^{-1/2} \\ &= (r^2)^{-1/2} \left(\left(1 + \frac{a^2 - 2ar \cos \alpha}{r^2}\right)^{-1/2} - \left(1 + \frac{a^2 + 2ar \cos \alpha}{r^2}\right)^{-1/2} \right) \\ &\approx \frac{1}{r} \left(1 - \frac{a^2 - 2ar \cos \alpha}{2r^2} - 1 + \frac{a^2 + 2ar \cos \alpha}{2r^2} \right) \\ &= \frac{2ar \cos \alpha}{r^3} = 2 \frac{\vec{a} \cdot \vec{x}}{|\vec{x}|^3}, \end{aligned}$$

wenn man den Ortsvektor des Aufpunkts mit $\vec{x}$ und die Ortsvektoren der Dipolladungen mit $\pm \vec{a}$ bezeichnet.

DIE HYPERBELFUNKTIONEN

Die trigonometrischen Funktionen können wir schreiben als

$$\cos x = \frac{e^{ix} + e^{-ix}}{2}, \quad \sin x = \frac{e^{ix} - e^{-ix}}{2i}.$$

Eine wesentliche Eigenschaft war $\cos^2 x + \sin^2 x = 1$. Dies bedeutete, dass der Punkt $(\cos x, \sin x)$ auf dem Einheitskreis liegt.

Definition 140 (Hyperbelfunktionen). Die Hyperbelfunktionen sind gegeben durch

$$\cosh x = \frac{e^x + e^{-x}}{2}, \quad \sinh x = \frac{e^x - e^{-x}}{2}.$$

Sie heißen *Cosinus hyperbolicus* und *Sinus hyperbolicus*.

Für sie gilt

$$\cosh^2 x - \sinh^2 x = 1,$$

denn

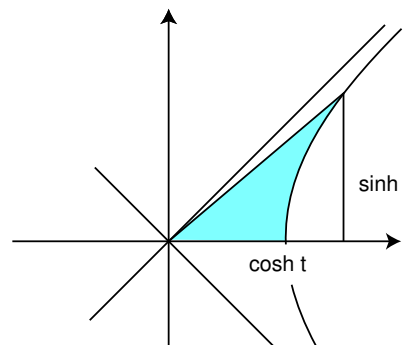
$$\cosh^2 x - \sinh^2 x = \frac{1}{4} ((e^x + e^{-x})^2 - (e^x - e^{-x})^2)$$

Das bedeutet aber, dass der Punkt $(\cosh x, \sinh x)$ auf der Einheitshyperbel

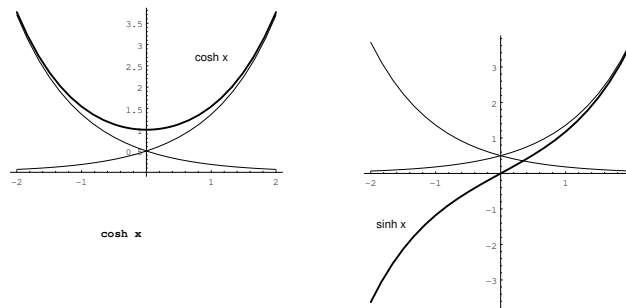
$$\{(x, y) \mid x^2 - y^2 = 1\} \quad (50)$$

liegt, daher der Name.

Der Wert x bedeutet in diesem Fall aber nicht etwa die Länge der Hyperbel bis zum Punkt $(\cosh x, \sinh x)$, sondern den doppelten Flächeninhalt der getönten Dreiecks.



Hier sind die Graphen der beiden Funktionen mit den Graphen von $\frac{1}{2}e^{\pm x}$ zum Vergleich:



Aus der Ableitungsformel für $\exp$ folgt sofort

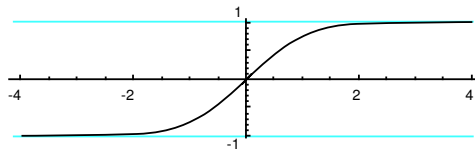
$$\cosh' x = \sinh x, \quad \sinh' x = \cosh x.$$

Auch für die Hyperbelfunktionen gibt es eine Menge Identitäten (wie Additionstheoreme u.a.), auf die wir hier aber nicht eingehen wollen.

Definition 141. Man definiert die hyperbolischen Cotangens- und Tangensfunktionen durch

$$\coth x = \frac{\cosh x}{\sinh x}, \quad \tanh x = \frac{\sinh x}{\cosh x}.$$

Der Tangens hyperbolicus bildet die reelle Achse streng monoton auf das Intervall $] -1, +1[$ ab.



Beispiel 142.

Energie-, Impuls- und Stofftransport, Abschnitt 2.3

Die Differentialgleichung

$$y'' = a^2 y$$

hat die allgemeine Lösung

$$y(x) = A e^{ax} + B e^{-ax}, \quad A, B \in \mathbb{R}.$$

Wegen

$$e^x = \sinh x + \cosh x, \quad e^{-x} = \sinh x - \cosh x$$

kann man die allgemeine Lösung aber auch schreiben als

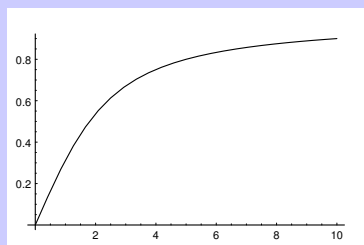
$$y(x) = C \cosh ax + D \sinh ax, \quad \text{wobei } C = A - B, D = A + B.$$

Beispiel 143 (Langevin-Funktion).*Werkstoffe I, Abschnitt 2.7*

Bei der Beschreibung von paramagnetischem Verhalten wie bei der Beschreibung der Polarisation von Dielektrika spielt die Langevin-Funktion

$$L(x) = \coth x - \frac{1}{x}$$

eine Rolle. Ihr Graph beschreibt einen Sättigungsvorgang.



Der Grenzwert $\lim_{x \rightarrow \infty} L(x) = 1$ ergibt sich leicht aus $L(x) = \frac{e^x + e^{-x}}{e^x - e^{-x}} - \frac{1}{x}$. Das Verhalten bei 0 ist hingegen schwieriger zu analysieren. Wir benutzen die Taylorentwicklungen im Nullpunkt

$$\begin{aligned} \cosh x &= \cosh 0 + \frac{\sinh 0}{1!}x + R_1(x) = 1 + R_1(x), \\ \sinh x &= \sinh 0 + \frac{\cosh 0}{1!}x + \frac{\sinh 0}{2!}x^2 + R_2(x) = x + R_2(x). \end{aligned}$$

Damit finden wir

$$\begin{aligned} L(x) &= \frac{\cosh x}{\sinh x} - \frac{1}{x} = \frac{x \cosh x - \sinh x}{x \sinh x} = \frac{x + xR_1(x) - x - R_2(x)}{x^2 + xR_2(x)} \\ &= \frac{\frac{R_1(x)}{x} - \frac{R_2(x)}{x^2}}{1 + \frac{R_2(x)}{x^2}} \rightarrow 0 \quad \text{für } x \rightarrow 0. \end{aligned}$$

Dieess ist ein weiteres Beispiel, wie man durch die Taylorapproximation komplizierte Funktionen vereinfachen kann.

7 Weitere Anwendungen der Differentiation

- Wir lernen Anwendungen der Differentialrechnung auf Extremalprobleme, Fehlerrechnung und verschiedene physikalische Probleme kennen.

EXTREMWERTE

Eine Funktion $f : I \rightarrow \mathbb{R}$ nimmt in $x_0 \in I$ ihr (*globales*) *Maximum* an, wenn

$$f(x) \leq f(x_0) \quad \text{für alle } x \in I. \quad (51)$$

Sie nimmt in x_0 ein *lokales Maximum* an, wenn

$$f(x) \leq f(x_0) \quad \text{für alle } x \in I \text{ nah bei } x_0, \quad (52)$$

das heißt, wenn es ein $\epsilon > 0$ gibt, so dass (52) für alle $x \in I$ mit $|x - x_0| < \epsilon$ gilt. Wir sprechen von einem strengen oder strikten (lokalen) Maximum, wenn in (51) bzw. (52) das Gleichheitszeichen nur für $x = x_0$ gilt. Analoge Begriffsbildungen definiert man für Minima.

Wir wissen über Extremwerte (=Maxima oder Minima) schon das folgende:

- Ist $f : [a, b] \rightarrow \mathbb{R}$ auf dem kompakten Intervall $[a, b]$ stetig, so nimmt f sein Maximum und Minimum (wirklich!) an.
- Ist $f : I \rightarrow \mathbb{R}$ differenzierbar, so sind die einzigen Kandidaten für Extremstellen die Endpunkte von I und die Punkte x , wo $f'(x) = 0$.

Wir beweisen nun den:

Satz 144 (Lokale Extremwerte). Sei $f : I \rightarrow \mathbb{R}$ eine n -mal differenzierbare Funktion und sei x_0 ein innerer Punkt von I . Es gelte

$$f'(x_0) = \dots = f^{(n-1)}(x_0) = 0, \quad f^{(n)}(x_0) \neq 0.$$

Dann gilt:

- (i) Ist n ungerade, so hat f in x_0 kein lokales Extremum, also erst recht kein globales.
- (ii) Ist n gerade, so hat f in x_0 ein lokales Extremum und zwar ein Maximum, wenn $f^{(n)}(x_0) < 0$, und ein Minimum, wenn $f^{(n)}(x_0) > 0$.

Für $n = 1$ sagt der Satz noch einmal, dass innere Punkte mit $f'(x) \neq 0$ nicht als Extremstellen infrage kommen.

Für $n = 2$ findet man, dass f in x_0 ein lokales Maximum hat, wenn $f'(x_0) = 0$ und $f''(x_0) < 0$ ist. Das ist vermutlich aus der Schule bekannt.

Beweis von Satz 144. Zum Beweis benutzen wir die Taylorformel:

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x - x_0) + \dots + \frac{f^{(n-1)}(x_0)}{(n-1)!}(x - x_0)^{n-1} + \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n + R_n(x) \\ &= f(x_0) + \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n + R_n(x). \end{aligned}$$

Weil $\lim_{x \rightarrow x_0} \frac{R_n(x)}{(x-x_0)^n} = 0$, bestimmt für x nah bei x_0 der Term $\frac{f^{(n)}(x_0)}{n!}(x-x_0)^n$ das Vorzeichen von $f(x) - f(x_0)$. Genauer ist

$$\left| \frac{R_n(x)}{(x-x_0)^n} \right| < \frac{1}{2} \left| \frac{f^{(n)}(x_0)}{n!} \right|$$

für x hinreichend nah bei x_0 . Für diese x haben daher

$$\frac{f^{(n)}(x_0)}{n!}(x-x_0)^n \text{ und } \frac{f^{(n)}(x_0)}{n!}(x-x_0)^n + R_n(x)$$

dasselbe Vorzeichen.

Bei ungeradem n wechselt $\frac{f^{(n)}(x_0)}{n!}(x-x_0)^n$ sein Vorzeichen in x_0 und es liegt kein Extremum vor.

Bei geradem n ist $(x-x_0)^n \geq 0$ für alle x , und deshalb bestimmt die n -te Ableitung $f^{(n)}(x_0)$ das Vorzeichen von $f(x) - f(x_0)$.

Daraus folgt die Behauptung. $\square$

Dieser Beweis zeigt noch einmal sehr gut, wie man die Taylorformel benutzen kann, um Eigenschaften von (einfachen) Polynomen auf (komplizierte) Funktionen zu übertragen.

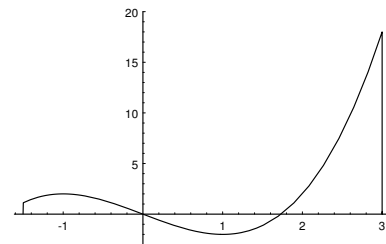
Beispiel 145. Sei $I = [-\frac{3}{2}, 3]$ und $f(x) := x^3 - 3x$. Es ist

$$f'(x) = 3x^2 - 3, \quad f''(x) = 6x.$$

Also sind die möglichen Kandidaten für Extrema die Endpunkte des Intervalls $-\frac{3}{2}$ und 3 und die Nullstellen -1 und 1 der Ableitung. Wegen $f''(-1) < 0$ und $f''(1) > 0$ hat f in -1 ein lokales Maximum und in 1 ein lokales Minimum. Die Funktionswerte in diesen Punkten sind

$$f(-\frac{3}{2}) = \frac{9}{8}, \quad f(-1) = 2, \quad f(1) = -2, \quad f(3) = 18.$$

Das globale Maximum im Intervall $[-\frac{3}{2}, 3]$ liegt bei $x = 3$, das globale Minimum bei 1. An der Stelle $x = -\frac{3}{2}$ liegt ein lokales Minimum. (Warum?)



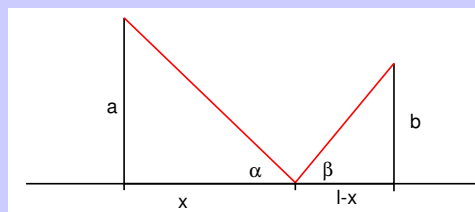
$\square$

Beispiel 146 (Reflexion nach Fermat). Die Lichtreflexion an einem Spiegel folgt dem Fermatschen Prinzip des kürzesten Laufweges. Wir betrachten zwei Punkte der Höhen a und b über dem Spiegel, deren Fußpunkte einen Abstand l haben.

Für welches x ist die Funktion

$$L(x) = \sqrt{a^2 + x^2} + \sqrt{b^2 + (l-x)^2}$$

minimal?



Für sehr große positive oder negative x ist der Laufweg offenbar sehr groß. Deshalb gibt es dazwischen ein Minimum. Wir lösen $L'(x) = 0$. Aus

$$\frac{x}{\sqrt{a^2 + x^2}} - \frac{l - x}{\sqrt{b^2 + (l - x)^2}} = 0 \quad (53)$$

folgt, dass x und $l - x$ gleiches Vorzeichen haben, also beide positiv sind, und dass

$$\frac{x^2}{a^2 + x^2} = \frac{(l - x)^2}{b^2 + (l - x)^2}.$$

Dann ist aber

$$x^2 b^2 = a^2 (l - x)^2$$

oder

$$bx = \pm a(l - x).$$

Da a , b , x und $l - x$ positiv sind, muss hier $+$ stehen:

$$x = \frac{al}{a + b}.$$

Es folgt

$$\cot \alpha = \frac{x}{a} = \frac{l}{a + b} = \cot \beta,$$

und wir erhalten das Spiegelungsgesetz „Einfallswinkel=Ausfallswinkel“.

FEHLERABSCHÄTZUNG

Ein wichtiger Punkt bei experimentellen Versuchen ist eine Diskussion der Messfehler. Ungenauigkeiten der Messung bewirken Ungenauigkeiten der daraus abgeleiteten Ergebnisse. Wie ist das Verhältnis von Ursache und Wirkung? Informationen liefert die Taylorapproximation. Es ist

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2}f''(\xi)(x - x_0)^2.$$

Wir ändern die Bezeichnungen und schreiben x statt x_0 und $x + \Delta x$ statt x :

$$f(x + \Delta x) = f(x) + f'(x)\Delta x + \frac{1}{2}f''(\xi)(\Delta x)^2.$$

Weiter schreiben wir $\Delta f := f(x + \Delta x) - f(x)$ für den Fehler in f :

$$\Delta f = f'(x)\Delta x + \frac{1}{2}f''(\xi)(\Delta x)^2.$$

Am häufigsten finden Sie eine qualitative *Fehlerapproximation*

$$\Delta f \approx f'(x) \Delta x$$

oder, für den absoluten Fehler,

$$|\Delta f| \approx |f'(x)| |\Delta x|.$$

Eine verlässliche maximale Schranke für den Fehler liefert

$$|\Delta f| = |f'(x) \Delta x + \frac{1}{2}f''(\xi)(\Delta x)^2| \leq |f'(x) \Delta x| + \frac{\max |f''(\xi)|}{2} (\Delta x)^2,$$

wenn man die zweite Ableitung abschätzen kann. Wir betrachten ein konkretes Beispiel:

Beispiel 147 (Wheatstone-Brücke). Bei der Wheatstone-Brücke zur Vergleichsmessung von Widerständen greift man auf einem Widerstandsdraht der Länge l eine Strecke x so ab, dass ein Potentiometer Nullstellung anzeigt. Dann ist das Verhältnis von unbekanntem Widerstand zu dem bekannten Vergleichswiderstand

$$\frac{R}{R_0} = \frac{x}{l-x}.$$

Also ist

$$R(x) = \frac{R_0 x}{l-x}.$$

Ein Messfehler von Δx am Abgleich produziert also einen Fehler von

$$\Delta R \approx R'(x) \Delta x = \frac{R_0 l}{(l-x)^2} \Delta x.$$

Ist konkret die Länge der Messstrecke $l = 10 \text{ cm}$ und erfolgt der Abgleich gegen den Vergleichswiderstand von $300 \, \Omega$ bei $x = 2 \pm 0.1 \text{ cm}$, so erhält man den gesuchten Widerstand $R = \frac{300 \cdot 2}{8} \, \Omega = 75 \, \Omega$ mit einer Genauigkeit von $\Delta R = \frac{300 \cdot 10}{64} 0,1 \, \Omega = 4,7 \, \Omega$.

8 Integration

8.1 Das bestimmte Integral

- Wir lernen das Integral als Grenzwert von Summen kennen.

VON DER SUMME ZUM INTEGRAL

Ein Punkt, der sich mit der Geschwindigkeit v bewegt, legt in einem Zeitintervall der Länge Δt die Strecke

$$\Delta s = v \Delta t$$

zurück. Ist die Geschwindigkeit hingegen zeitabhängig,

$$v = v(t),$$

so ist ein Standardverfahren zur Ermittlung der zurückgelegten Strecke das folgende. Man teilt das Zeitintervall, sagen wir $[a, b]$, in kleine Teilintervalle

$$a = t_0 \leq t_1 \leq \dots \leq t_n = b$$

und nimmt an, dass auf jedem dieser Teilintervalle $[t_{j-1}, t_j]$ die Geschwindigkeit annähernd konstant v_j ist. Setzen wir $\Delta t_j := t_j - t_{j-1}$, so ist die in diesem Teilintervall zurückgelegte Strecke annähernd $v_j \Delta t_j$ und die Gesamtstrecke ist

$$s \approx \sum_{j=1}^n v_j \Delta t_j.$$

Dieselbe Methode benutzt man zum Beispiel zur Bestimmung des Widerstandes oder der Masse eines (homogenen) Drahtes wechselnden Querschnittes und bei vielen anderen Gelegenheiten. Sie ist unter anderem *das* typische Argumentationsmuster beim Übergang von der Mechanik einzelner Massenpunkte zur Kontinuumsmechanik, etwa bei der Berechnung des Schwerpunktes eines Körpers, vgl. *Müller: Mechanik II, Abschnitt 14.1.1* und den Anhang Abschnitt 11.7.

Die Methode liefert zunächst allerdings nur eine Approximation, weil v ja auf $[t_{j-1}, t_j]$ im allgemeinen nicht wirklich konstant sein wird. Aber wenn man die Teilintervalle „genügend klein“ macht, erwartet man eine „genügend gute“ Approximation. „Approximation“ von was eigentlich? In unserem Beispiel von der zurückgelegten Gesamtstrecke, die ja eine physikalische Realität ist. Mathematisch erwartet man diese Gesamtstrecke als Grenzwert der Summen bei immer feinerer Intervallzerlegung, als ein *Integral*.

Die Präzisierung dieser vagen Vorstellungen muss klären, unter welchen Voraussetzungen tatsächlich ein Grenzwert existiert, dessen Eigenschaften die physikalische Realität vernünftig wiedergeben. Der dafür benötigte begriffliche Aufwand hängt sehr stark von den betrachteten Funktionen ab. Wir wählen hier eine Klasse von Funktionen, die alle in der Ingenieurspraxis auftretenden Funktionen umfaßt, aber gleichzeitig erlaubt, fast ohne jeden mathematischen Formalismus auszukommen.¹⁵

Wir bezeichnen die Funktionen allgemein wieder mit $f(x)$ statt $v(t)$. Sei also $f : [a, b] \rightarrow \mathbb{R}$ eine Funktion. Wir unterteilen $[a, b]$ in n Teile gleicher Länge:

$$a = x_0 < x_1 < \dots < x_n = b.$$

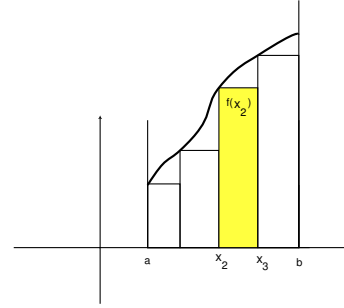
Also ist

$$x_j := a + j \frac{b-a}{n} \quad \text{und} \quad \Delta x_j = \Delta x = \frac{b-a}{n}.$$

¹⁵Vergleiche Anhang Abschnitt 11.6.

Indem wir f auf $[x_{j-1}, x_j]$ durch die Konstante $f(x_{j-1})$ approximieren, bilden wir

$$F_n := \sum_{j=1}^n f(x_{j-1}) \Delta x_j = \frac{b-a}{n} \sum_{j=1}^n f\left(a + (j-1) \frac{b-a}{n}\right).$$



Solche Summen F_n nennt man auch *Riemannsche Summen* der Funktion f .

Satz 148 (und Integraldefinition). Ist $f : [a, b] \rightarrow \mathbb{R}$

- *monoton oder*
- *stetig,*

so existiert der Grenzwert

$$\lim_{n \rightarrow \infty} F_n =: \int_a^b f(x) dx.$$

Wir nennen ihn das Integral von f über $[a, b]$.

Wir verzichten auf einen Beweis.

Beispiel 149 (Exakte Berechnung der Fallstrecke).

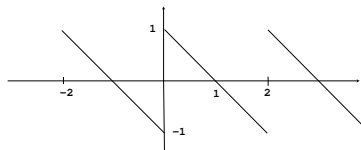
Müller: *Mechanik II*, Abschnitt 12.1.2

Befindet sich ein Massenpunkt im freien Fall, so ist die Geschwindigkeit proportional zur Fallzeit, also $v(t) = gt$. Wir berechnen die Fallstrecke für das Zeitintervall $[0, T]$. Wir erhalten

$$\begin{aligned} \int_0^T gt \, dt &= \lim_{n \rightarrow \infty} \frac{T}{n} \sum_{j=1}^n \frac{g(j-1)T}{n} \\ &= \lim_{n \rightarrow \infty} \frac{gT^2}{n^2} \sum_{j=1}^n (j-1) = \lim_{n \rightarrow \infty} \frac{gT^2}{n^2} \frac{(n-1)n}{2} \\ &= \lim_{n \rightarrow \infty} \frac{gT^2}{2} \frac{(1 - \frac{1}{n})}{1} = \frac{gT^2}{2}. \end{aligned}$$

□

Es kommt nicht so selten vor, dass man in den Ingenieurwissenschaften Funktionen integrieren muss, die weder monoton noch stetig sind, etwa die Steuerspannung des Kathodenstrahls einer Fernsehröhre:



Wir müssen daher unsere Integraldefinition noch etwas verallgemeinern:

Definition 150. (i) Wir nennen $f : [a, b] \rightarrow \mathbb{R}$ *stückweise stetig* bzw. *stückweise monoton*, wenn wir das Intervall $[a, b]$ in endlich viele Teilintervalle zerlegen können:

$$a = x_0 \leq x_1 \leq \dots \leq x_n = b,$$

so dass $f :]x_{j-1}, x_j[\rightarrow \mathbb{R}$ für jedes i die Einschränkung einer stetigen bzw. monotonen Funktion

$$f_j : [x_{j-1}, x_j] \rightarrow \mathbb{R}$$

ist. Die Werte von f an den Teilpunkten selbst sind belanglos.

(ii) Ist f wie in (i), so definieren wir

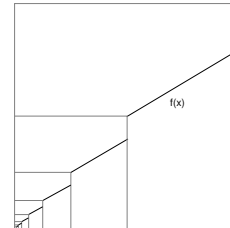
$$\int_a^b f(x) dx := \int_{x_0}^{x_1} f_1(x) dx + \dots + \int_{x_{n-1}}^{x_n} f_n(x) dx.$$

Beispiel 151.

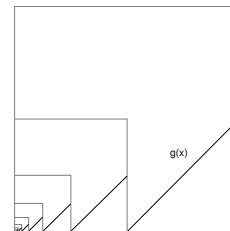
Die Funktion

$$f(x) := \begin{cases} 0 & \text{für } x = 0, \\ \frac{1}{2^k} + \frac{x}{2} & \text{für } \frac{1}{2^{k+1}} < x \leq \frac{1}{2^k}, k \in \mathbb{N} \end{cases}$$

ist eine monoton wachsende Funktion auf $[0, 1]$.



Die Funktion $g(x) = f(x) - x$ ist auf einer unendlichen Folge von Intervallen zwar monoton wachsend, aber nicht monoton und auch nicht stückweise monoton, eben weil man nicht mit endlich vielen Monotonie-Intervallen auskommt.



□

Integrierbare Funktionen. Wir haben damit das Integral definiert für alle stückweise stetigen oder stückweise monotonen Funktionen auf einem kompakten Intervall $[a, b]$. Wir vereinbaren, dass wir diese Funktionen im Rest des Skriptums kurz als *integrierbare Funktionen* bezeichnen, auch wenn sich das mit der mathematischen Standardterminologie nicht unbedingt verträgt: Dort ist die Klasse der integrierbaren Funktionen etwas größer.¹⁶

Integrierbare Funktionen sind immer beschränkt und auf einem kompakten Intervall definiert. Die Funktionen

$$f_1(x) := \begin{cases} \frac{1}{\sqrt{x}} & \text{für } 0 < x \leq 1 \\ 0 & \text{für } x = 0 \end{cases}$$

und

$$f_2(x) := e^{-x^2} \quad \text{für } -\infty < x < +\infty$$

¹⁶Das vorstehende Beispiel zeigt, dass unsere integrierbaren Funktionen keinen Vektorraum bilden, man müßte die Integraldefinition eigentlich auf den von ihnen erzeugten Vektorraum ausdehnen. Wir verzichten darauf.

sind deshalb nicht integrierbar, wohl aber ihre Einschränkungen

$$f_1|_{[a,1]} \quad \text{und} \quad f_2|_{[b,c]}$$

für $0 < a < 1$ bzw. beliebige reelle Zahlen $b < c$.

Allerdings möchte man auch Funktionen wie f_1 und f_2 doch gern integrieren, das Integral von f_2 zum Beispiel ist in der Wahrscheinlichkeitstheorie von zentraler Bedeutung. Vergleichen Sie dazu den Abschnitt 8.5 über *Uneigentliche Integrale*.

RECHENREGELN FÜR DAS INTEGRAL.

$$\int_a^b (f(x) + g(x))dx = \int_a^b f(x)dx + \int_a^b g(x)dx, \quad (54)$$

$$\int_a^b \lambda f(x)dx = \lambda \int_a^b f(x)dx, \quad \lambda \in \mathbb{R}, \quad (55)$$

$$\int_a^b f(x)dx \leq \int_a^b g(x)dx, \quad \text{falls } f \leq g, \quad (56)$$

$$\left| \int_a^b f(x)dx \right| \leq \int_a^b |f(x)|dx, \quad (57)$$

$$\int_a^b f(x)dx + \int_b^c f(x)dx = \int_a^c f(x)dx \quad \text{für } a < b < c. \quad (58)$$

Wir setzen dabei voraus, dass alle auftretenden Funktionen integrierbar sind. Die Beweise sind alle für approximierende Summen ziemlich klar und übertragen sich dann im Grenzwert auf das Integral.

Eine einfache Folgerung ist der

Satz 152 (Integralabschätzung). Seien $f : [a, b] \rightarrow \mathbb{R}$ integrierbar und $m \leq f(x) \leq M$ für alle $x \in [a, b]$. Dann gilt

$$m(b-a) \leq \int_a^b f(x)dx \leq M(b-a).$$

Zum Beweis betrachtet man die Funktion $g(x) := m$ für alle x und erhält aus $g \leq f$ die linke Ungleichung:

$$m(b-a) = \int_a^b g(x)dx \leq \int_a^b f(x)dx.$$

Die andere folgt ebenso.

Ist f sogar stetig und wählt man $m = \min_{a \leq x \leq b} f(x)$, $M = \max_{a \leq x \leq b} f(x)$, so ist also

$$m \leq \frac{1}{b-a} \int_a^b f(x)dx \leq M.$$

Nach dem Zwischenwertsatz für stetige Funktionen nimmt f alle Werte zwischen m und M wirklich an. Also gibt es ein $\xi \in]a, b[$ mit

$$f(\xi) = \frac{1}{b-a} \int_a^b f(x)dx.$$

Diesen Sachverhalt nennt man auch den **Mittelwertsatz der Integralrechnung**.

8.2 Das unbestimmte Integral

- Wir entdecken, dass Integration und Differentiation zueinander inverse Prozesse sind. Das ist oft hilfreich für die Berechnung von Integralen.

Wir kommen nun zum Zusammenhang zwischen Integration und Differentiation, der für die Berechnung von vielen Integralen überaus nützlich ist.

Wir betrachten dazu eine integrierbare Funktion $f : [a, b] \rightarrow \mathbb{R}$ und für eine variable obere Grenze $x \in [a, b]$ das Integral

$$F(x) := \int_a^x f(t) dt.$$

Dabei haben wir die Integrationsvariable jetzt mit t bezeichnet, weil wir x schon für die obere Grenze vergeben hatten. In Abhängigkeit von x definiert das also eine neue Funktion $F : [a, b] \rightarrow \mathbb{R}$, die man auch *das unbestimmte Integral* oder die *Integralfunktion* von f nennt. Wir versuchen F zu differenzieren:

$$\begin{aligned} \frac{F(x + \Delta x) - F(x)}{\Delta x} &= \frac{1}{\Delta x} \left(\int_a^{x+\Delta x} f(t) dt - \int_a^x f(t) dt \right) \\ &= \frac{1}{\Delta x} \left(\int_a^x f(t) dt + \int_x^{x+\Delta x} f(t) dt - \int_a^x f(t) dt \right) \\ &= \frac{1}{\Delta x} \int_x^{x+\Delta x} f(t) dt. \end{aligned}$$

Bemerkung. Diese Rechnung gilt zunächst nur, wenn $\Delta x > 0$. Offenbar ist es hilfreich, auch Integrale $\int_a^b f(x) dx$ zuzulassen, bei denen nicht $a < b$ ist. Dazu definieren wir

$$\int_a^a f(x) dx = 0 \quad \text{und} \quad \int_a^b f(x) dx = - \int_b^a f(x) dx, \quad \text{falls } a > b.$$

Dann bleibt die vorstehende Rechnung auch mit $\Delta x < 0$ richtig.

Wenn wir nun voraussetzen, dass f **stetig** ist, so können wir den Mittelwertsatz der Integralrechnung anwenden (vgl. 7.1) und finden, dass die rechte Seite sich schreiben lässt als $f(\xi)$, wobei ξ eine Zahl zwischen x und $x + \Delta x$ ist, über die wir zwar nichts Genaueres wissen, die aber für $\Delta x \rightarrow 0$ jedenfalls gegen x geht. Daher ist für **stetiges** f das unbestimmte Integral F differenzierbar und hat f als Ableitung:

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

Definition 153. Sei $f : I \rightarrow \mathbb{R}$ eine Funktion auf einem Intervall $I \subset \mathbb{R}$. Ist $F : I \rightarrow \mathbb{R}$ differenzierbar und gilt $F' = f$, so nennt man F eine *Stammfunktion* von f .

Nicht jede Funktion hat eine Stammfunktion:

Beispiel 154 (Heavisidefunktion). Die sogenannte Heaviside-Funktion $f = Y : \mathbb{R} \rightarrow \mathbb{R}$ mit $Y(x) := 0$ für $x < 0$ und $Y(x) := 1$ für $x \geq 0$ hat keine Stammfunktion: Ein solche müsste für $x \leq 0$ konstant und für $x > 0$ von der Form $x + b$ sein. Ist die Konstante $= b$, so ist die Funktion wenigstens stetig, hat aber bei 0 einen Knick und ist deshalb dort nicht differenzierbar. Also ist sie keine Stammfunktion. Zwar ist nur der Punkt 0 ein Problem, aber eben ein unlösbares. $\square$

Stetige Funktionen *haben* dagegen immer wenigstens eine Stammfunktion, nämlich das unbestimmte Integral. Das bedeutet aber nicht, dass Sie immer eine Stammfunktion „explizit“ *hinschreiben* können.

Stetige Funktionen haben sogar sehr viele Stammfunktionen. Ist nämlich F eine Stammfunktion von f und c eine beliebige Konstante, so ist auch $G(x) := F(x) + c$ eine Stammfunktion von f , weil additive Konstanten beim Differenzieren wegfallen:

$$G'(x) = F'(x) = f(x).$$

Das ist trivial.

Es gilt auch die Umkehrung: Sind F und G zwei Stammfunktionen von f auf dem Intervall I , so gilt $G = F + c$ für eine Konstante c . Je zwei Stammfunktionen von f unterscheiden sich also nur durch eine Konstante.

Beweis: Ist $G' = f = F'$, so folgt $(G - F)' = f - f = 0$, und eine Funktion mit Ableitung 0 auf einem Intervall ist konstant. (Konstanzkriterium der Differentialrechnung, Satz 107).

Wir fassen zusammen:

Satz 155 (Hauptsatz der Differential- und Integralrechnung). Sei $f : I \rightarrow \mathbb{R}$ stetig auf einem Intervall I und $a \in I$. Dann gilt:

(i) Die durch

$$F(x) := \int_a^x f(t)dt, \quad x \in I,$$

definierte Funktion, also das unbestimmte Integral von f , ist eine Stammfunktion von f , d.h.

$$F'(x) = f(x).$$

(ii) Umgekehrt gilt für jede Stammfunktion G von f auf I , dass

$$G(x) = \int_a^x f(t)dt + \text{const.} = \int_a^x f(t)dt + G(a).$$

(iii) Insbesondere ist für $I = [a, b]$ dann also

$$\int_a^b f(t)dt = G(b) - G(a) =: G(x)|_a^b.$$

Das hat phantastische Konsequenzen: Nehmen Sie an, Sie sollten $\int_0^\pi \sin x \, dx$ berechnen. Dann müssten Sie den Grenzwert

$$\lim_{n \rightarrow \infty} \frac{\pi}{n} \sum_{j=1}^n \sin\left((j-1)\frac{\pi}{n}\right)$$

berechnen. Versuchen Sie's mal!

Oder Sie erinnern sich daran, dass $\cos' x = -\sin x$, also $G(x) = -\cos x$ eine Stammfunktion von $f(x) = \sin x$ ist. Daher ist nach dem Hauptsatz

$$\int_0^\pi \sin x \, dx = -\cos \pi - (-\cos 0) = 1 + 1 = 2.$$

Keine Unterteilungen, keine Grenzwerte! Und es funktioniert für alle stetigen Funktionen, für die wir Stammfunktionen kennen!

Bemerkung zur Notation. Aus den nun ersichtlichen Gründen schreibt man oft

$$G(x) = \int f(x)dx + \text{const.},$$

wenn G eine Stammfunktion von f ist. Es ist hilfreich, dabei immer eine Integrationskonstante zu notieren. Sonst erhält man z.B. $-\cos x = \int \sin x \, dx = 23 - \cos x$, weil $G(x) = -\cos x$ aber genauso gut auch $H(x) = 23 - \cos x$ Stammfunktionen von $\sin x$ sind.

Beispiele.

$$\begin{aligned} \int_a^b e^x \, dx &= e^x \Big|_a^b = e^b - e^a, & \int_1^5 \frac{dx}{x} &= \ln 5 - \ln 1 = \ln 5, \\ \int_a^b \sin x \, dx &= -\cos x \Big|_a^b = \cos a - \cos b, & \int_0^b \frac{1}{1+x^2} \, dx &= \arctan b - \arctan 0 = \arctan b, \\ \int x^\alpha \, dx &= \frac{x^{\alpha+1}}{\alpha+1} + \text{const. für } \alpha \neq -1, & \int_a^b \frac{1}{\sqrt{1-x^2}} \, dx &= \arcsin x \Big|_a^b, \\ \int (a_0 + a_1 x + a_2 x^2 + \dots + a_n x^n) \, dx &= a_0 x + a_1 \frac{x^2}{2} + \dots + a_n \frac{x^{n+1}}{n+1} + \text{const.} \end{aligned}$$

Mathematische Software mit der Fähigkeit zum symbolischen Rechnen bietet auch die Möglichkeit, unbestimmte Integrale zu berechnen. Hier ein paar Zeilen Mathematica

```
In[1]:=Integrate[x^2(Sin[x])^2,x]
Out[1]=1/24(4x^3-6xCos[2x]+3Sin[2x]-6x^2Sin[2x])

In[2]:=Integrate[Sin[x^2],x]
Out[2]=Sqrt[Pi/2]FresnelS[Sqrt[2/Pi]x]
```

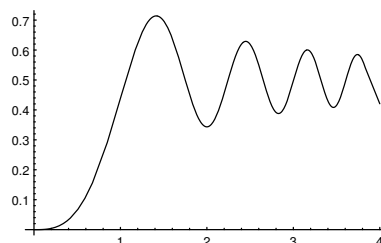
Während das erste Ergebnis sofort verständlich ist, bedarf das zweite einer Erklärung: Die Funktion $\sin(x^2)$ besitzt als stetige Funktion natürlich eine Stammfunktion. Sie besitzt aber *keine* Stammfunktion, die sich mit „elementaren“ Funktionen ausdrücken lässt. Die obigen Beispiele der Funktionen $\frac{1}{x}$ oder $\frac{1}{1+x^2}$ zeigen schon, dass relativ „einfache“ Funktionen zu komplizierten Stammfunktionen führen. Und für $\sin(x^2)$ ergibt sich als Stammfunktion ein sogenanntes *Fresnelintegral*, eine Funktion, die *definiert* ist durch die Gleichung

$$\text{FresnelS}(x) = \int_0^x \sin\left(\frac{\pi t^2}{2}\right) dt.$$

Das Fresnelintegral spielt eine Rolle in der geometrischen Optik. Die Werte bekommt man wie die Werte des Sinus zum Beispiel vom Rechner geliefert:

```
In[3]:=N[FresnelS[1.5],7]
Out[3]=0.697505

In[4]:=Plot[FresnelS[x],{x,0,4}]
Out[4]=
```



Ein wesentlicher Unterschied zwischen der Sinusfunktion und dem Fresnelintegral ist, dass viele Eigenschaften des Sinus einfacher zu beschreiben sind, zum Beispiel die Lage der lokalen Maxima. Auch hat man eine Fülle von Identitäten, die den Sinus mit anderen Funktionen verbinden.

Beispiel 156 (Normalverteilung, Fehlerfunktion).

Werkstoffe II (Abschnitt 8.3), Verfahrenstechnik I (Kapitel 2) oder Energie-, Impuls- und Stofftransport (Abschnitt 2.5)

Eine andere Funktion, deren Stammfunktion sich nicht mit den uns bekannten Funktionen ausdrücken lässt, ist e^{-x^2} . Die (mit einem Faktor normierte) Stammfunktion

$$\operatorname{erf}(x) := \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt$$

heißt die *Gaußsche Fehlerfunktion* oder *Normalverteilung*. Sie spielt eine wichtige Rolle bei der Theorie von Diffusionsprozessen und in der Statistik.

8.3 Integrationsregeln

- Integrieren ist schwieriger als Differenzieren, weil es keine “Integral-Kettenregel” und keine “Integral-Produktregel” gibt.
- Aber die Integration der normalen Kettenregel und Produktregel liefert Rechenregeln für das Integral, die trotzdem oft hilfreich für die Vereinfachung komplizierter Integrale sind.

Der Hauptsatz der Differential- und Integralrechnung ermöglicht die Berechnung von Integralen **stetiger** Funktionen mittels Stammfunktionen auf sehr elegante Weise, *wenn* man eine Stammfunktion kennt. Und weil wir früher eine Menge Funktionen differenziert haben, kennen wir für eine Menge von Funktionen auch die Stammfunktionen. Aber leider für viele Funktionen auch nicht: Was ist $\int e^{x^2} dx$ oder $\int \sqrt{1-x^2} dx$?

Wir wollen im folgenden einige Regeln kennen lernen, die uns helfen können, Integrale komplizierterer Funktionen auf einfachere zurückzuführen. So wie einem die Rechenregeln der Differentialrechnung, insbesondere die Produktregel und die Kettenregel geholfen haben, kompliziertere Funktionen zu differenzieren, ohne jedesmal wieder Grenzwerte von Differenzenquotienten auszurechnen.

Allerdings ist die Situation bei der Integration nicht so angenehm, wie bei der Differentiation, weil es eben für die Integration von Produkten und geschachtelten Funktionen keine oder keine einfachen Formeln gibt. Wenn man eine Stammfunktion von e^x und eine von x^2 kennt, kann man leider keine von e^{x^2} hinschreiben - tatsächlich gibt es für diese Funktion gar keine Stammfunktion, die sich durch „elementare Funktionen“ ausdrücken lässt. So wie das Integral der simplen rationalen Funktion $\frac{1}{x}$ auf die viel kompliziertere Logarithmusfunktion führt, führt die Integration von e^{x^2} auf eine „höhere“ Funktion als wir aus dem bisherigen mathematischen Alltag kennen.

Die hier zu besprechenden Integrationsregeln sind die Gegenstücke zur Kettenregel und Produktregel der Differentialrechnung.

SUBSTITUTIONSREGEL

Die Funktion $F(x)$ sei eine Stammfunktion der stetigen Funktion f , und die Funktion $x(t)$ sei differenzierbar mit stetiger Ableitung. Dann gilt nach der Kettenregel

$$f(x(t))\dot{x}(t) = F'(x(t))\dot{x}(t) = \frac{d}{dt}F(x(t)).$$

Aus dem Hauptsatz der Differential- und Integralrechnung folgt deshalb

$$\int_{\alpha}^{\beta} f(x(t))\dot{x}(t)dt = F(x(\beta)) - F(x(\alpha))$$

Andererseits ist wieder nach dem Hauptsatz $F(x(\beta)) - F(x(\alpha)) = \int_{x(\alpha)}^{x(\beta)} f(x)dx$. Zusammengefasst ergibt sich die **Substitutionsregel** (= Ersetzungsregel)

$$\boxed{\int_{\alpha}^{\beta} f(x(t))\dot{x}(t)dt = \int_{x(\alpha)}^{x(\beta)} f(x)dx.} \quad (59)$$

Natürlich hat man dabei vorausgesetzt, dass $x(t)$ auf dem Intervall zwischen α und β nur Werte annimmt, für die f auch definiert ist. Ist $x : [\alpha, \beta] \rightarrow [a, b]$ bijektiv und monoton wachsend, so kann man (59) auch schreiben als

$$\boxed{\int_a^b f(x)dx = \int_{x^{-1}(a)}^{x^{-1}(b)} f(x(t))\dot{x}(t)dt.} \quad (60)$$

Beispiel 157 (Kreisfläche). Die Fläche des Halbkreises vom Radius r um den Nullpunkt ist

$$\int_{-r}^{+r} \sqrt{r^2 - x^2} dx.$$

Wir betrachten $x : [-\frac{\pi}{2}, \frac{\pi}{2}] \rightarrow [-r, +r]$ mit

$$x(t) := r \sin t.$$

Dann liefert die Substitutionsregel:

$$\begin{aligned} \int_{-r}^{+r} \sqrt{r^2 - x^2} dx &= \int_{-\pi/2}^{\pi/2} \sqrt{r^2 - r^2 \sin^2 t} r \cos t dt = r^2 \int_{-\pi/2}^{\pi/2} \sqrt{1 - \sin^2 t} \cos t dt \\ &= r^2 \int_{-\pi/2}^{\pi/2} \cos^2 t dt = r^2 \int_{-\pi/2}^{\pi/2} \frac{1}{2} (\cos 2t + 1) dt = \frac{r^2}{2} \left(\frac{1}{2} \sin 2t + t \right) \Big|_{-\pi/2}^{\pi/2} = \frac{r^2}{2} \pi. \end{aligned}$$

Die ganze Kreisfläche ist also πr^2 . □

Jetzt wenden wir die Form (60) an:

Beispiel 158. Wir wollen

$$\int_1^7 \frac{dx}{(5x-2)^3} \tag{61}$$

berechnen und substituieren zur Vereinfachung

$$t = 5x - 2. \tag{62}$$

Dann ist

$$\frac{dt}{dx} = 5, \text{ also(?) } \frac{dx}{dt} = \frac{1}{5}, \tag{63}$$

und daher

$$\int_1^7 \frac{dx}{(5x-2)^3} = \int_3^{33} \frac{1}{t^3} \frac{dt}{5} = \frac{1}{5} \frac{1}{(-2)t^2} \Big|_3^{33} = -\frac{1}{10t^2} \Big|_3^{33} = -\frac{1}{10} \left(\frac{1}{33^2} - \frac{1}{3^2} \right) = 0.011.$$

Und meistens schreibt man (63) noch in der Form

$$dx = \frac{1}{5} dt.$$

Warum darf man so rechnen? Die Problematik liegt in der Zeile (62). Dort definiert man nicht $x(t)$, sondern die Umkehrfunktion $t(x)$. Sie sehen, dass $x(t) = \frac{t+2}{5}$. Für die Ableitung von $x(t)$ benutzt man in Gleichung (63) den Satz über die Ableitung der Umkehrfunktion und erhält einfach die Inverse der Ableitung an der entsprechenden Stelle¹⁷. Wichtig ist, dass die Funktion $t(x)$ wirklich invertierbar ist. □

Beispiel 159. Jetzt bearbeiten wir ebenso

$$\int_1^7 \frac{dx}{(5x^2-2)^3}.$$

Wir setzen

$$t = 5x^2 - 2$$

¹⁷Die taucht hier aber gar nicht auf, weil die Ableitung konstant ist.

und erhalten

$$\frac{dt}{dx} = 10x, \text{ also } \frac{dx}{dt} = \frac{1}{10x}$$

und

$$\int_1^7 \frac{dx}{(5x^2 - 2)^3} = \int_3^{243} \frac{1}{t^3} \frac{1}{10x} dt = \int_3^{243} \frac{1}{t^3} \frac{1}{10\sqrt{(t+2)/5}} dt.$$

Durch die Anwendung der Substitutionsregel ist das Integral also nicht einfacher, sondern komplizierter geworden. $\square$

Gelegentlich ist eine Substitution $x(t)$ mit fallendem x naheliegend. Die Substitutionsregel (59) bleibt auch richtig, wenn $\alpha > \beta$ und $x : [\beta, \alpha] \rightarrow [a, b]$.

Beispiel 160 (Ein Trick für ein wichtiges Integral).

Schwingungsenergie

Wir suchen $\int_0^{\pi/2} \cos^2 x dx$. Das haben wir im Beispiel 157 schon einmal berechnet, jetzt gehen wir einen anderen Weg. Wir setzen $x(t) = \pi/2 - t$, also $\dot{x}(t) = -1$:

$$\int_0^{\pi/2} \cos^2 x dx = \int_{\pi/2}^0 \cos^2(\pi/2 - t)(-1)dt = - \int_{\pi/2}^0 \sin^2 t dt = \int_0^{\pi/2} \sin^2 x dx.$$

Machen Sie sich diese Gleichheit anhand der Graphen anschaulich klar! Es sieht aus, als sei sie kein großer Fortschritt. Aber wegen $\sin^2 + \cos^2 = 1$ ist

$$\int_0^{\pi/2} \sin^2 x dx + \int_0^{\pi/2} \cos^2 x dx = \int_0^{\pi/2} 1 dx = \frac{\pi}{2}.$$

Aus beidem zusammen ergibt sich

$$\int_0^{\pi/2} \sin^2 x dx = \int_0^{\pi/2} \cos^2 x dx = \frac{\pi}{4}.$$

Das Integral ist zum Beispiel wichtig, weil die Energie einer harmonischen Schwingung

$$f(t) = A \cos t$$

gegeben wird durch das Periodenintegral

$$\int_0^{2\pi} A^2 \cos^2 t dt = 4 \int_0^{\pi/2} A^2 \cos^2 t dt = \pi A^2.$$

Beispiel 161 (Eine Falle). Wir berechnen $\int_0^\pi \sin x dx = 2$ nicht direkt, sondern mit der Substitution $\sin x = t$, also $\frac{dt}{dx} = \cos x = \sqrt{1 - t^2}$:

$$\int_0^\pi \sin x dx \underset{[\sin x = t]}{=} \int_0^0 t \frac{1}{\sqrt{1 - t^2}} dt = 0,$$

weil die obere und untere Integrationsgrenze gleich sind. Wo steckt der Fehler? Vielleicht finden Sie das heraus, wenn Sie das Integral bei $\pi/2$ unterteilen. $\square$

Man kann die Substitutionsregel natürlich auch zur Berechnung unbestimmter Integrale (= Stammfunktionen) benutzen. Sucht man eine Stammfunktion für $f(x)$ und findet eine, nämlich $G(t)$, für $f(x(t))\dot{x}(t)$, und ist $G(t) = F(x(t))$, so ist F eine Stammfunktion von f . Man kann also substituieren, das unbestimmte Integral berechnen und muss dann anschließend die Transformation $x \rightarrow t$ wieder rückgängig machen, d.h. t wieder durch x ausdrücken. (Dabei sollte $x(t)$ eine eindeutige Funktion sein.)

Beispiel 162. Zur Berechnung von $\int \sqrt{1-x^2} dx$ substituieren wir $x = \sin t$:

$$\begin{aligned}\int \sqrt{1-x^2} dx &= \int \sqrt{1-\sin^2 t} \cos t dt = \int \cos^2 t dt \\ &= \int \frac{1}{2}(\cos 2t + 1) dt = \frac{1}{2} \left(\frac{1}{2} \sin 2t + t \right) + \text{const.} \\ &= \frac{1}{2}(\sin t \cos t + t) + \text{const.} = \frac{1}{2}(x\sqrt{1-x^2} + \arcsin x) + \text{const.}\end{aligned}$$

Die Rechnung scheint ein wenig „großzügig“. Ist wirklich $\sqrt{1-\sin^2 t} = \cos t$? Oder ist es $= -\cos t$? Wenn wir davon ausgehen, dass $-1 \leq x \leq 1$, weil sonst der Integrand nicht reell ist, können wir für den Sinus das Definitionsintervall $-\frac{\pi}{2} \leq t \leq \frac{\pi}{2}$ nehmen. Dort ist der Cosinus dann positiv und der Arcussinus ist auch definiert. Zur Probe (und zur Übung) kann man natürlich die rechte Seite differenzieren. $\square$

PARTIELLE INTEGRATION

Seien $u, v : [a, b] \rightarrow \mathbb{R}$ differenzierbar mit stetigen Ableitungen. Aus der Produktregel

$$(uv)' = u'v + uv'$$

und dem Hauptsatz der Differential- und Integralrechnung folgt dann

$$(uv)|_a^b = \int_a^b u'(x)v(x) dx + \int_a^b u(x)v'(x) dx.$$

Das schreiben wir in der Form

$$\boxed{\int_a^b u'(x)v(x) dx = (uv)|_a^b - \int_a^b u(x)v'(x) dx.}$$

und nennen es die Regel der **partiellen Integration** (= Integration nach Teilen). Natürlich hat man wieder eine entsprechende Formel für die unbestimmten Integrale

$$\boxed{\int u'(x)v(x) dx = uv - \int u(x)v'(x) dx.}$$

Beispiel 163. Wir wollen $\int_a^b x e^x dx$ berechnen. Wir versuchen den Ansatz

$$u' = x, v = e^x.$$

Dann ist $u = \frac{1}{2}x^2, v' = e^x$ und die partielle Integration liefert

$$\int_a^b x e^x dx = \left(\frac{1}{2} x^2 e^x \right) \Big|_a^b - \int_a^b \frac{1}{2} x^2 e^x dx.$$

Das hat nichts gebracht! Wir versuchen nun

$$u' = e^x, v = x.$$

Dann bekommen wir $u = e^x, v' = 1$ und

$$\int_a^b x e^x dx = (x e^x) \Big|_a^b - \int_a^b e^x = ((x-1)e^x) \Big|_a^b$$

$\square$

Beispiel 164. Mit einem hübschen Trick kann man das Integral von $\ln x$ berechnen: Wir setzen $u' = 1, v = \ln x$ und erhalten mit $u = x, v' = \frac{1}{x}$:

$$\int \ln x \, dx = x \ln x - \int 1 \, dx = x(\ln x - 1) + \text{const.}$$

□

SCHLUSSBEMERKUNG ZUM „KNACKEN“ VON INTEGRALEN.

Die geschickte Anwendung der Integrationsregeln zum expliziten Ausrechnen von Integralen ist eine Kunst, es gibt dafür keine Rezepte. Mittlerweile sind aber die Computerprogramme kaum noch zu schlagen. Wenn man (oder der Computer) erfolglos ist, kann das verschiedene Gründe haben: Möglicherweise findet man (oder der Computer) nicht den richtigen Trick, möglicherweise ist aber das Integral auch nicht durch eine elementare Funktion lösbar, sondern definiert eine neue Funktion, wie wir es oben beim Fresnel-Integral gesehen haben. Wir schließen mit zwei Beispielen solcher Integralen, die sich aus relativ einfachen Anwendungsproblemen ergeben.

Im ersten Beispiel berechnen wir die Länge einer Ellipse. Für eine beliebige stetig differenzierbare Kurve

$$C : (x(t), y(t)) \text{ mit } a \leq t \leq b$$

in der Ebene ist die Länge einfach das Integral über die Geschwindigkeit:

$$L = \int_a^b \sqrt{\dot{x}^2(t) + \dot{y}^2(t)} dt.$$

Für Kurven

$$(x(t), y(t), z(t)) \text{ mit } a \leq t \leq b$$

im 3-dimensionalen Raum ergibt sich entsprechend

$$L = \int_a^b \sqrt{\dot{x}^2(t) + \dot{y}^2(t) + \dot{z}^2(t)} dt.$$

Beispiel 165 (Elliptisches Integral und Kurvenlänge). Die Parametrisierung

$$C : x(t) = a \sin t, y(t) = b \cos t \quad 0 \leq t \leq 2\pi$$

mit $a > b > 0$ beschreibt eine Ellipse mit den Halbachsen a und b . Ihre Länge ist

$$\begin{aligned} L &= \int_0^{2\pi} \sqrt{a^2 \cos^2 t + b^2 \sin^2 t} \, dt = \int_0^{2\pi} \sqrt{a^2 - (a^2 - b^2) \sin^2 t} \, dt \\ &= a \int_0^{2\pi} \sqrt{1 - \frac{a^2 - b^2}{a^2} \sin^2 t} \, dt. \end{aligned}$$

Für $0 < k \leq 1$ heißt

$$E(k) := \int_0^{\pi/2} \sqrt{1 - k^2 \sin^2 x} \, dx$$

das *vollständige elliptische Integral 2. Art*. Die Länge der Ellipse ist

$$L = 4aE(k), \quad k = \sqrt{\frac{a^2 - b^2}{a^2}},$$

daher der Name elliptisches Integral. Für $k \neq 1$ besitzt $\sqrt{1 - k^2 \sin^2 x}$ (wie der Fresnel-Integrand im letzten Abschnitt) keine elementar auszudrückende Stammfunktion. □

Beispiel 166 (Schwingungsdauer des Pendels).

Müller: *Mechanik II*, Abschnitt 15.2

Die Pendelgleichung ist

$$\ddot{\phi} + \frac{g}{L} \sin \phi = 0. \quad (64)$$

Für kleine Auslenkungen ist $\sin \phi \approx \phi$ und man betrachtet die linearisierte Gleichung

$$\ddot{\phi} + \frac{g}{L} \phi = 0.$$

Die Lösung ist $\phi(t) = \alpha \sin \sqrt{\frac{g}{L}}(t - t_0)$. Das ist eine harmonische Schwingung mit der Frequenz $\omega = \sqrt{\frac{g}{L}}$, also mit der Schwingungsdauer

$$T = \frac{2\pi}{\omega} = 4 \sqrt{\frac{L}{g}} \frac{\pi}{2}.$$

Insbesondere hängt T nicht von der Schwingungsamplitude ab – solange diese so klein ist, dass die Linearisierung der Gleichung sinnvoll ist. Für große Pendelschwingungen muss man aber die Differentialgleichung (64) verwenden. Diese ist nicht mehr so einfach lösbar, genauer gesagt gar nicht mit elementaren Funktionen. Ein längere Rechnung ergibt eine von der Amplitude α abhängige Schwingungsdauer

$$T = 4 \sqrt{\frac{L}{g}} K\left(\sin \frac{\alpha}{2}\right).$$

Dabei ist

$$K(k) := \int_0^{\pi/2} \frac{dx}{\sqrt{1 - k^2 \sin^2 x}}$$

das sogenannte *vollständige elliptische Integral 1. Art*, das auch nicht über eine elementare Stammfunktion auszuwerten ist.

8.4 Integration komplexer und rationaler Funktionen

- Komplexwertige Funktionen eines reellen Arguments integriert man nach Real- und Imaginärteil getrennt.

Für komplexwertige Funktionen $f : [a, b] \rightarrow \mathbb{C}$ einer reellen Variablen definiert man das Integral durch Zerlegung in Real- und Imaginärteil: Ist $f = u + iv$ mit $u, v : [a, b] \rightarrow \mathbb{R}$, so sei

$$\int_a^b f(x) dx := \int_a^b u(x) dx + i \int_a^b v(x) dx.$$

Die Integrationsregeln (54), (55), (57) und (58) sowie der Hauptsatz der Differential- und Integralrechnung, die Substitutionsregel und die partielle Integration übertragen sich sofort auf diese Situation, während Monotonie-Aussagen wie (56) oder der Satz 152 natürlich im Komplexen keinen Sinn mehr machen. Auch der Mittelwertsatz der Integralrechnung gilt nicht mehr.

Beispiel 167. Für $k \geq 2$ und $z \in \mathbb{C}$ hat

$$f(x) := \frac{1}{(x-z)^k}$$

die Stammfunktion

$$F(x) = \frac{-1}{(k-1)(x-z)^{k-1}}.$$

Daher ist, wenn $z \notin [a, b]$

$$\int_a^b \frac{dx}{(x-z)^k} = \frac{1}{(k-1)(a-z)^{k-1}} - \frac{1}{(k-1)(b-z)^{k-1}}.$$

□

Vorsicht aber bei der Integration von $1/(x-z)$ mit komplexem $z = a + ib$ oder ähnlichem. Die Funktion $\ln|x-z|$ ist *keine* Stammfunktion. Ihre Ableitung ist wie sie selbst nämlich reellwertig, also sicher nicht gleich $1/(x-z)$. Wir betrachten dazu das

Beispiel 168. Es gilt

$$\begin{aligned} \int_0^1 \frac{dx}{x-i} &= \int_0^1 \frac{(x+i)}{(x+i)(x-i)} dx = \int_0^1 \left(\frac{x}{x^2+1} + i \frac{1}{x^2+1} \right) dx \\ &= \left(\frac{1}{2} \ln(1+x^2) + i \arctan x \right) \Big|_0^1 = \frac{1}{2} \ln 2 + i \frac{\pi}{4}. \end{aligned}$$

Das ist aber verschieden von $\ln|1-i| - \ln|0-i| = \frac{1}{2} \ln 2$.

□

Wozu brauchen Sie das? Die Integration komplexwertiger Funktionen eines reellen Arguments ist vor allem bei der Fourieranalyse von Schwingungsvorgängen in der Mechanik, Akustik oder Elektrotechnik wichtig. Darauf zielt das folgende Beispiel.

Beispiel 169. Seien $k, l \in \mathbb{Z}$ ganze Zahlen, $T \in \mathbb{R}$ positiv und $\omega = \frac{2\pi}{T}$. Dann gilt

$$\frac{1}{T} \int_0^T e^{ik\omega t} e^{-il\omega t} dt = \frac{1}{T} \int_0^T e^{i(k-l)\omega t} dt = \frac{1}{T} \frac{e^{i(k-l)\omega t}}{i(k-l)\omega} \Big|_0^T = \frac{1}{T} \frac{e^{i(k-l)2\pi} - e^0}{i(k-l)\omega} = 0,$$

falls $k \neq l$. Ist $k = l$, so ist der Integrand $= 1$ und das Integral $= T$. Also finden wir

$$\frac{1}{T} \int_0^T e^{ik\omega t} e^{-il\omega t} dt = \begin{cases} 0 & \text{falls } k \neq l, \\ 1 & \text{falls } k = l. \end{cases}$$

Andrerseits ist

$$\begin{aligned} \frac{1}{T} \int_0^T e^{ik\omega t} e^{-il\omega t} dt &= \frac{1}{T} \int_0^T (\cos k\omega t + i \sin k\omega t)(\cos l\omega t - i \sin l\omega t) dt \\ &= \frac{1}{T} \int_0^T (\cos k\omega t \cos l\omega t + \sin k\omega t \sin l\omega t) dt \\ &\quad + \frac{i}{T} \int_0^T (-\cos k\omega t \sin l\omega t + \sin k\omega t \cos l\omega t) dt. \end{aligned}$$

Durch Vergleich mit dem vorigen Ergebnis sieht man, dass der Imaginärteil 0 ist, also

$$\frac{1}{T} \int_0^T \cos k\omega t \sin l\omega t dt = \frac{1}{T} \int_0^T \sin k\omega t \cos l\omega t dt$$

Ersetzt man k durch $-k$, so bleibt die linke Seite unverändert, die rechte ändert aber ihr Vorzeichen. Daher sind beide Seiten 0:

$$\boxed{\frac{1}{T} \int_0^T \cos k\omega t \sin l\omega t dt = 0 \quad \text{für alle } k, l \in \mathbb{Z}.}$$

Ein analoges Argument liefert

$$\boxed{\frac{1}{T} \int_0^T \cos k\omega t \cos l\omega t dt = \frac{1}{T} \int_0^T \sin k\omega t \sin l\omega t dt = 0 \quad \text{für alle } k \neq l \in \mathbb{N}.}$$

□

Rationale Funktionen. Rationale Funktionen der Form $\frac{p(x)}{q(x)}$ lassen sich explizit integrieren: Man zerlegt sie mittels Partialbruchzerlegung und muss dann nur die dabei auftretenden Typen von Integralen beherrschen. Wir beschränken uns hier auf den Fall, dass p und q reell sind, und der Nennergrad höher als der Zählergrad ist. Eine einfache reelle Nullstelle x_0 des Nenners führt auf ein Integral

$$\int \frac{a dx}{x - x_0} = a \ln |x - x_0| + \text{const.}, \quad x_0 \text{ reell!}$$

Eine komplexe Nullstelle $\xi + i\eta$ kommt zusammen mit einem konjugiert-komplexen Partner $\xi - i\eta$ und

$$\begin{aligned} \frac{A}{x - (\xi + i\eta)} + \frac{B}{x - (\xi - i\eta)} &= \frac{(A + B)(x - \xi) + i(A - B)\eta}{\eta^2 + (x - \xi)^2} \\ &= \frac{a(x - \xi)}{\eta^2 + (x - \xi)^2} + \frac{b\eta}{\eta^2 + (x - \xi)^2}. \end{aligned}$$

Aber dann ist mit der Substitutionsregel

$$\begin{aligned} \int \frac{x - \xi}{\eta^2 + (x - \xi)^2} dx &= \int \frac{\frac{x - \xi}{\eta}}{1 + \left(\frac{x - \xi}{\eta}\right)^2} \frac{dx}{\eta} = \int \frac{t}{1 + t^2} dt = \frac{1}{2} \ln(t^2 + 1) + \text{const.} \\ &= \ln \sqrt{1 + \left(\frac{x - \xi}{\eta}\right)^2} + \text{const.} \end{aligned}$$

und

$$\begin{aligned}\int \frac{\eta \, dx}{\eta^2 + (x - \xi)^2} &= \int \frac{1}{1 + \left(\frac{x-\xi}{\eta}\right)^2} \frac{dx}{\eta} = \int \frac{dt}{1+t^2} = \arctan t + \text{const.} \\ &= \arctan \frac{x - \xi}{\eta} + \text{const.}\end{aligned}$$

Bei mehrfachen Nullstellen des Nenners ergibt die (komplexe) Partialbruchzerlegung zusätzlich Terme der Form

$$\int \frac{dx}{(x-a)^k}$$

mit $k \geq 2$ und möglicherweise komplexem a . Diese besitzen eine Stammfunktion

$$-\frac{1/(k-1)}{(x-a)^{k-1}}.$$

Beispiel 170. Was ist

$$\int \frac{4x^2 - 4x}{(x^2 + 1)^2} dx = \int \frac{4x^2 - 4x}{(x+i)^2(x-i)^2} dx?$$

Hier haben wir zwei doppelte Nullstellen. Die Partialbruchzerlegung liefert

$$\frac{4x^2 - 4x}{(x+i)^2(x-i)^2} = \frac{A}{x+i} + \frac{B}{(x+i)^2} + \frac{\bar{A}}{x-i} + \frac{\bar{B}}{(x-i)^2}$$

mit $A = i, B = 1 - i$, also, wenn man die beiden Terme mit linearem Nenner zusammenfaßt,

$$\frac{4x^2 - 4x}{(x+i)^2(x-i)^2} = \frac{2}{x^2 + 1} + \frac{1-i}{(x+i)^2} + \frac{1+i}{(x-i)^2}$$

Integration liefert

$$\int \frac{4x^2 - 4x}{(x+i)^2(x-i)^2} dx = 2 \arctan x - \frac{1-i}{x+i} - \frac{1+i}{x-i} = 2 \arctan x - \frac{2x-2}{x^2+1}.$$

□

8.5 Uneigentliche Integrale

- Wie integriert man unbeschränkte Funktionen und wie Funktionen über unbeschränkte Intervalle?

Wir haben die Integration bisher nur über kompakte Intervalle $[a, b]$ definiert und unsere integrierbaren Funktionen waren immer beschränkt. In diesem Abschnitt wollen wir Integrale für den Fall definieren, dass der Integrationsbereich oder die Funktion unbeschränkt ist.

Definition 171. Sei $f : [a, \infty[\rightarrow \mathbb{R}$ eine Funktion, die über jedes Intervall $[a, b]$ integrierbar ist. Dann definieren wir das *uneigentliche Integral*

$$\int_a^\infty f(x) dx := \lim_{b \rightarrow \infty} \int_a^b f(x) dx. \quad (65)$$

Dieses uneigentliche Integral kann existieren - oder auch nicht.

Beispiel 172.

$$\int_0^\infty \frac{dx}{1+x^2} = \lim_{b \rightarrow \infty} \int_0^b \frac{1}{1+x^2} dx = \lim_{b \rightarrow \infty} \arctan x \Big|_0^b = \lim_{b \rightarrow \infty} (\arctan b - \arctan 0) = \frac{\pi}{2}.$$

□

Beispiel 173.

$$\int_1^\infty \frac{dx}{x} = \lim_{b \rightarrow \infty} \ln x \Big|_1^b = \lim_{b \rightarrow \infty} \ln b,$$

und das existiert nicht.

□

Beispiel 174. Sei $\alpha > 0$, $\alpha \neq 1$. Dann ist

$$\begin{aligned} \int_1^\infty \frac{dx}{x^\alpha} &= \lim_{b \rightarrow \infty} \frac{1}{(1-\alpha)x^{\alpha-1}} \Big|_1^b = \lim_{b \rightarrow \infty} \left(\frac{1}{(1-\alpha)b^{\alpha-1}} - \frac{1}{1-\alpha} \right) \\ &= \begin{cases} \frac{1}{\alpha-1}, & \text{falls } \alpha > 1, \\ \text{nicht existent,} & \text{falls } \alpha < 1. \end{cases} \end{aligned}$$

Aus Beispiel 173 wissen wir, dass das uneigentliche Integral für $\alpha = 1$ ebenfalls nicht existiert. Vergleichen Sie die Skizze in Beispiel 179.

□

Beispiel 175. Weil e^{-x^2} keine elementare Stammfunktion besitzt, kann man die Existenz von

$$\int_0^\infty e^{-x^2} dx \quad (66)$$

nicht wie bei den vorstehenden Beispielen nachweisen. Wir benutzen eine andere Methode. Weil $e^{-x^2} > 0$, ist offenbar

$$b \mapsto \int_0^b e^{-x^2} dx$$

in b monoton steigend. Wenn wir zeigen können, dass diese Funktion außerdem beschränkt ist, ist sie für $b \rightarrow \infty$ konvergent, und (66) existiert. Für $x \geq 1$ ist nun $e^{-x^2} \leq e^{-x}$ und daher für $b > 1$

$$\begin{aligned} \int_0^b e^{-x^2} dx &= \int_0^1 e^{-x^2} dx + \underbrace{\int_1^b e^{-x^2} dx}_{\leq \int_1^b e^{-x} dx = e^{-1} - e^{-b} < e^{-1}} < \int_0^1 e^{-x^2} dx + e^{-1}. \end{aligned}$$

Das zeigt die Beschränktheit und damit die Existenz des uneigentlichen Integrals. Seinen Wert können wir mit unseren bisherigen Methoden nicht berechnen. Der Computer liefert $\sqrt{\pi}/2$, und das werden wir in der Analysis II bestätigen.

□

Das uneigentliche Integral $\int_{-\infty}^b f(x)dx$ definiert man analog zum obigen, und man setzt

$$\int_{-\infty}^{\infty} f(x)dx = \int_{-\infty}^0 f(x)dx + \int_0^{\infty} f(x)dx,$$

falls die beiden rechten Integrale existieren. Das ist eine stärkere Forderung als die Bedingung, dass $\lim_{b \rightarrow \infty} \int_{-b}^b f(x)dx$ existiert, wie man z.B. an $\int_{-\infty}^{\infty} x dx$ sieht.

Beispiel 176. Nach dem letzten Beispiel und aus Symmetriegründen existiert $\int_{-\infty}^0 e^{-x^2} dx$ und deshalb

$$\int_{-\infty}^{\infty} e^{-x^2} dx = (\sqrt{\pi}).$$

□

Beispiel 177 (Elektronenkonzentration und Fermi-Integral).

Werkstoffe II, Abschnitt 1.12

Bei der Berechnung der Elektronenkonzentration in einem Leitungsband tritt das Integral

$$\int_{W_L}^{\infty} \frac{\sqrt{W - W_L}}{e^{\frac{W - W_F}{kT}} + 1} dW$$

auf. Dabei sind die Energieniveaus W_L und W_F ebenso konstant, wie k (=Boltzmannkonstante) und T (=Temperatur). Das uneigentliche Integral existiert aus ähnlichen Gründen wie das Integral $\int_0^{\infty} e^{-x^2} dx$, und so wie dieses lässt es sich nicht mit einer elementaren Stammfunktion lösen. Wir versuchen trotzdem die Substitution $\eta = \frac{W - W_L}{kT}$ und erhalten

$$\int_{W_L}^{\infty} \frac{\sqrt{W - W_L}}{e^{\frac{W - W_F}{kT}} + 1} dW = \int_0^{\infty} \frac{\sqrt{\eta} \sqrt{kT}}{e^{\eta+s} + 1} kT d\eta = (kT)^{\frac{3}{2}} \int_0^{\infty} \frac{\sqrt{\eta}}{e^{\eta+s} + 1} d\eta$$

mit $s = \frac{W_L - W_F}{kT}$. Das uneigentliche Integral

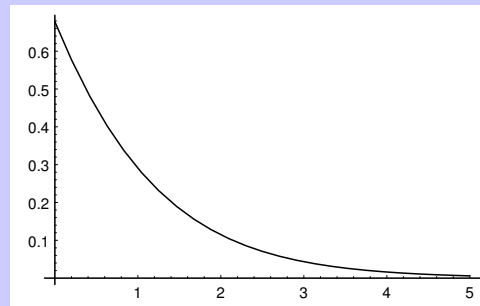
$$F_{\frac{1}{2}}(s) = \int_0^{\infty} \frac{\eta^{1/2}}{e^{\eta+s} + 1} d\eta$$

kann man auch nicht mit einer Stammfunktion explizit berechnen, aber es hängt nur noch von einem Parameter s ab, gegenüber anfänglich vieren: W_L, W_F, k und T . Die Funktion $F_{\frac{1}{2}}$ heißt Fermi-Integral. Sie lässt sich – ebenso wie z.B. die trigonometrischen Funktionen – numerisch berechnen.

Mit den Mathematica-Befehlen

```
F[s] := NIntegrate[  $\frac{\sqrt{\eta}}{(\text{Exp}[\eta+s]+1)}$ , {η, 0, ∞}];  
Plot[F[s], {s, 0, 5}]
```

erhält man den Graphen von $F_{\frac{1}{2}}$.



Wir kommen nun zur **Integration unbeschränkter Funktionen**.

Definition 178. Sei $f :]a, b] \rightarrow \mathbb{R}$ eine Funktion, so dass für alle $c \in]a, b]$ das Integral $\int_c^b f(x)dx$ existiert. Dann definieren wir das uneigentliche Integral

$$\int_a^b f(x)dx = \lim_{c \searrow a} \int_c^b f(x)dx.$$

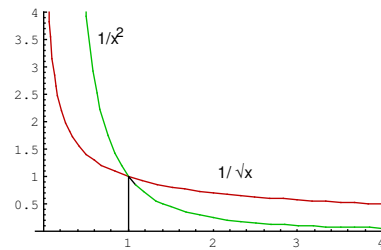
Es kann existieren - oder auch nicht.

Weil $\lim_{c \rightarrow 0} \ln c = -\infty$, existiert zum Beispiel $\int_0^1 \frac{dx}{x}$ nicht, vgl. die Rechnung im analogen Beispiel oben.

Beispiel 179. Sei $\alpha > 0, \alpha \neq 1$.

$$\begin{aligned} \int_0^1 \frac{dx}{x^\alpha} &= \lim_{a \searrow 0} \frac{1}{(1-\alpha)x^{\alpha-1}} \Big|_a^1 = \lim_{a \searrow 0} \left(\frac{1}{(1-\alpha)} - \frac{1}{(1-\alpha)a^{\alpha-1}} \right) \\ &= \begin{cases} \text{nicht existent,} & \text{falls } \alpha > 1 \\ \frac{1}{1-\alpha}, & \text{falls } \alpha < 1. \end{cases} \end{aligned}$$

Für die Integrale $\int_0^1 \frac{dx}{x^\alpha}$ und $\int_1^\infty \frac{dx}{x^\alpha}$ ist also $\alpha = 1$ der kritische Parameterwert: für $\alpha = 1$ existieren die uneigentlichen Integrale beide nicht. Ist die Funktion aber $< \frac{1}{x}$, so existieren sie. Das ist in $]0, 1]$ für $\alpha < 1$ und in $[1, \infty[$ für $\alpha > 1$ der Fall.



□

9 Fourieranalysis

- Differenzierbare Funktionen lassen sich lokal gut durch (Taylor-)Polynome approximieren.
- Wir lernen, wie man periodische Funktionen sogar global durch Sinus- und Cosinusfunktionen approximiert.
- Die Fourierapproximation ist das mathematische Werkzeug für die Frequenzanalyse von periodischen Schwingungen.

Die Taylorformel gibt einem die Möglichkeit, eine (genügend oft differenzierbare) Funktion f in der Nähe eines Punktes x_0 durch ein Polynom zu approximieren:

$$f(x) \approx \sum_{k=0}^n a_k (x - x_0)^k \quad \text{mit } a_k = \frac{f^{(k)}(x_0)}{k!}.$$

Die Formel für die Koeffizienten hatten wir gefunden, indem wir uns überlegt hatten, wie man die Koeffizienten eines *Polynoms* $f(x)$ aus $f(x)$ herausbekommt: Nämlich durch Auswertung der Ableitungen im Entwicklungspunkt, vgl. Beispiel 112. Die Taylor-Approximation eines *Polynoms* liefert also wieder dasselbe Polynom. Für andere Funktionen gibt der Satz von Taylor (Satz 114) Auskunft über den Approximationsfehler.

In der Taylor-Theorie approximiert man also komplizierte Funktionen mit einfachen Bausteinen, nämlich Potenzen von x , genauer mit Linearkombinationen von Potenzen von x , also mit Polynomen.

Bei den in der Mechanik und Elektrotechnik häufig auftretenden periodischen Funktionen (Schwingungen, Wechselstrom) ist es naheliegend, andere Bausteine zu nehmen, nämlich solche die selber schon periodisch sind: Sinus- und Cosinusschwingungen zum Beispiel, sogenannte harmonische Schwingungen. Wir bezeichnen die Variable in diesem Abschnitt mit t , weil sie oft eine Zeit repräsentiert.

9.1 Reelle Fourieranalysis

TRIGONOMETRISCHE POLYNOME

Definition 180 (Periodizität). Wir nennen eine Funktion T -periodisch (mit $T > 0$), wenn $f(t + T) = f(t)$ für alle t .¹⁸

Beispiel 181. Sei $T > 0$ und $\omega := \frac{2\pi}{T}$. Die Funktionen

$$\cos k\omega t, \quad \sin k\omega t$$

sind für jedes $k \in \mathbb{N}$ periodisch mit der Periode T :

$$\cos k\omega(t + T) = \cos(k\omega t + k2\pi) = \cos k\omega t$$

Beachten Sie, dass $\cos 0\omega t = 1$ natürlich auch T -periodisch ist.

Entsprechendes gilt für den Sinus. Dann sind aber auch Linearkombinationen

$$\sum_{k=0}^n (a_k \cos k\omega t + b_k \sin k\omega t) \tag{67}$$

¹⁸ T ist nicht unbedingt die kleinste Periode von f , siehe das folgende Beispiel.

mit reellen Koeffizienten a_k und b_k T -periodische Funktionen.

Ebenso sind die komplexwertigen Funktionen

$$\sum_{k=-n}^n c_k e^{ik\omega t}, \quad c_k \in \mathbb{C} \quad (68)$$

für $n \in \mathbb{N}$ periodisch mit Periode $T = \frac{2\pi}{\omega}$. □

Definition 182 (Trigonometrische Polynome). Funktionen der Form (67) oder (68) heißen *trigonometrische Polynome* vom Grad oder von der Ordnung n .

Anschauliche Interpretation. Physikalisch besteht ein solches Polynom also aus der Überlagerung einer Grundschiwingung der Frequenz $\omega = \frac{2\pi}{T}$ und Oberschwingungen der Frequenzen $k\omega$. Die Koeffizienten geben an, mit welcher Amplitude, d.h. wie stark die Oberschwingungen vertreten sind.

Wir beschränken uns einstweilen auf den reellen cos-sin-Fall. Es ist klar, dass die trigonometrischen Polynome ähnlich wie die Polynome eine Klasse „einfacher“ Funktionen bilden und daher als Ausgangsbasis für Approximationsvorhaben dienen können.

ORTHOGONALITÄTSRELATIONEN UND FOURIERKOEFFIZIENTEN

Die trigonometrischen „Bausteine“ $\cos k\omega t$ und $\sin k\omega t$ besitzen eine Eigenschaft, die wir von den Potenzen x^k nicht kennen, und durch die die Fourieranalysis eine ganz eigene Prägung erhält:

Satz 183 (Orthogonalitätsrelationen). Für alle $k, l \in \mathbb{N}$ und $\omega = \frac{2\pi}{T} > 0$ gilt:

$$\begin{aligned} \frac{2}{T} \int_0^T \cos k\omega t \cos l\omega t \, dt &= \begin{cases} 2, & \text{falls } k = l = 0 \\ 1, & \text{falls } k = l > 0 \\ 0, & \text{falls } k \neq l \end{cases} \\ \frac{2}{T} \int_0^T \sin k\omega t \sin l\omega t \, dt &= \begin{cases} 1, & \text{falls } k = l > 0 \\ 0, & \text{sonst} \end{cases} \\ \frac{2}{T} \int_0^T \cos k\omega t \sin l\omega t \, dt &= 0 \end{aligned}$$

Zum Beweis kann man von trigonometrischen Formeln für $\cos k\omega t \cos l\omega t$ usw. ausgehen, oder – viel einfacher – das Beispiel 169 benutzen. Dann bleibt nur der Fall $k = l$ und der ist einfach.

Im Anhang Abschnitt 11.9 erklären wir, was Orthogonalrelationen mit „orthogonal“ zu tun haben.

REELLE FOURIERAPPROXIMATION

Nun stellen wir uns in Anlehnung an die Taylorentwicklung die folgende Frage: Gegeben eine Funktion, von der wir wissen, dass sie ein trigonometrisches Polynom ist, von der wir aber nicht die Darstellung (67) kennen. Wie können wir die Koeffizienten a_k, b_k herausbekommen? Bei den Polynomen hatte uns die Differentiation geholfen, weil Potenzen beim Differenzieren schließlich verschwinden. Bei trigonometrischen Polynomen hilft dagegen die Integration:

Beispiel 184. Gegeben seien $T = \frac{2\pi}{\omega}$ und ein trigonometrisches Polynom

$$\phi(t) = \sum_{k=0}^n (a_k \cos k\omega t + b_k \sin k\omega t).$$

Wir multiplizieren diese Gleichung mit $\sin l\omega t$ für irgendein $0 < l \leq n$ und integrieren:

$$\begin{aligned} \frac{2}{T} \int_0^T \phi(t) \sin l\omega t dt &= \sum_{k=0}^n \left(\frac{2}{T} \int_0^T a_k \cos k\omega t \sin l\omega t dt + \frac{2}{T} \int_0^T b_k \sin k\omega t \sin l\omega t dt \right) \\ &= \sum_{k=0}^n \left(\underbrace{a_k \frac{2}{T} \int_0^T \cos k\omega t \sin l\omega t dt}_{=0} + \underbrace{b_k \frac{2}{T} \int_0^T \sin k\omega t \sin l\omega t dt}_{=0, \text{ außer für } k=l} \right) \\ &= b_l. \end{aligned}$$

Eine ähnliche Rechnung kann man für den Cosinus machen, wobei der Fall $l = 0$ eine kleine Sonderrolle spielt. Wir erhalten $\frac{2}{T} \int_0^T \phi(t) \cos l\omega t dt = a_l$ für $l > 0$ und $= 2a_0$ für $l = 0$. Dieses Problem verschieben wir an eine andere Stelle, indem wir a_0 durch $\frac{a_0}{2}$ ersetzen, und trigonometrische Polynome in Zukunft in der Form

$$\phi(t) = \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t)$$

schreiben. □

Stückweise Monotonie. Im folgenden haben wir es oft mit Integralen zu tun, an denen eine T -periodische Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ beteiligt ist. Damit diese Integrale existieren, machen wir die Generalvoraussetzung, dass $f : \mathbb{R} \rightarrow \mathbb{R}$ *stückweise monoton*¹⁹ ist, vgl. Abschnitt 8.1.

Das impliziert insbesondere, dass f beschränkt ist. Beispiele sind natürlich die Sinus- und Cosinusfunktionen aber auch Sägezahn- oder Rechteckfunktionen.

Alle im Zusammenhang mit der Fourieranalysis betrachteten reellen Funktionen seien auf dem Periodenintervall $[0, T]$ stückweise monoton. Komplexwertige Funktionen sollen stückweise monotonen Real- und Imaginärteil haben.

Nun betrachten wir eine T -periodische Funktion $f(t)$, die (vielleicht) kein trigonometrisches Polynom ist, die wir aber durch ein solches approximieren wollen. Dazu definieren wir

Definition 185 (Fourierkoeffizienten). Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ (oder $\mathbb{C}$) eine Funktion der Periode $T = \frac{2\pi}{\omega}$. Dann definieren wir für alle $k \in \mathbb{N}$ die *Fourierkoeffizienten* von f als

$$\begin{aligned} a_k &= \frac{2}{T} \int_0^T f(t) \cos k\omega t dt, \\ b_k &= \frac{2}{T} \int_0^T f(t) \sin k\omega t dt. \end{aligned}$$

Wir nennen das mit diesen Koeffizienten gebildete trigonometrische Polynom

$$\phi(t) = \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t)$$

das n -te *Fourierpolynom* oder das *Fourierpolynom n -ter Ordnung* zur Funktion $f(t)$. Manchmal schreibt man dafür genauer auch $\phi_n(t)$ oder $\phi_n^f(t)$ oder ähnliches.

¹⁹Im Zusammenhang mit der Fouriertheorie ist das eine bessere Voraussetzung als die stückweise Stetigkeit. Vgl. dazu den Satz 217.

Es ist natürlich a priori nicht klar, ob das Fourierpolynom n -ter Ordnung (für hinreichend großes n) wirklich eine gute Approximation für $f(t)$ darstellt. Wie groß ist der Fehler? Wir kommen auf diese Frage später zurück und unterstellen einstweilen, dass er klein ist:

$$f(t) \approx \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t).$$

Die *Fourieranalyse* einer periodischen Funktion $f(t)$, d.h. die Berechnung der a_k und b_k , und die *Fouriersynthese*, das Zusammenfügen der Koeffizienten zu einem trigonometrischen Polynom, haben viele Anwendungen.

- Eine benutzen Sie ständig: Im menschlichen Ohr sind die Haarzellen des Cortischen Organs jeweils für bestimmte Frequenzen empfindlich. Das Ohr übermittelt dem Gehirn also die Fourierkoeffizienten der von ihm aufgenommenen akustischen Signale.
- Rauschunterdrückungs- oder Kompressionsverfahren (etwa für MP3) zerlegen Signale mit der Fourieranalyse in ihr Frequenzspektrum, filtern die unerwünschten oder überflüssigen Frequenzen heraus und setzen das Signal dann wieder zusammen.
- Die Wirkung linearer Systeme etwa in der Regelungstechnik läßt sich an harmonischen Schwingungen testen und mit der Fourieranalyse für beliebigen periodischen Input vorhersagen.
- Während Generatoren natürlicherweise harmonische Spannungen liefern, ist man bei technischen Anwendungen zum Beispiel an linearen Sägezahnspannungen interessiert (etwa für die Zeilensteuerung des Elektronenstrahls in einer Bildröhre). Die Fourieranalyse liefert Auskunft darüber, wie man durch Überlagerung harmonischer Schwingungen solche „willkürlichen“ Spannungen „synthetisieren“ kann.
- Auf überraschend andere Weise dient die Fourieranalyse bei Rand-Anfangswert-Problemen partieller Differentialgleichungen und damit bei sehr vielen Problemen der Verfahrens-, Energie- oder Elektrotechnik als wichtiges Hilfsmittel. Darauf gehen wir etwas später noch ein, vgl. Beispiel 192.

9.2 Anwendungsbeispiele

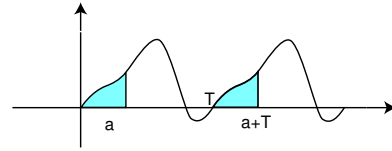
- Rauschunterdrückung, Klirrfaktor und ein Ausblick auf die Anwendung der Fourierapproximation für Anfangswertprobleme partieller Differentialgleichungen.

Bevor wir Beispiele betrachten, machen wir noch eine einfache Feststellung, die einem unnötigen Rechenaufwand ersparen kann.

Zunächst ist für T -periodische Funktionen und beliebiges $a \in \mathbb{R}$

$$\int_0^T f(t) dt = \int_a^{a+T} f(t) dt.$$

Machen Sie sich das am Graphen anschaulich klar!



Deshalb kann man $\int_0^T$ zum Beispiel auch durch $\int_{-\frac{T}{2}}^{\frac{T}{2}}$ ersetzen.

Ist andererseits g eine *gerade* Funktion, also $g(-t) = g(t)$, so ist für alle $a > 0$

$$\int_{-a}^{+a} g(t) dt = 2 \int_0^{+a} g(t) dt,$$

während für *ungerade* Funktionen $u(t) = -u(-t)$ das Integral verschwindet:

$$\int_{-a}^{+a} u(t) dt = 0.$$

Auch dies macht man sich sofort graphisch klar. Weil aber für gerades f die Funktion $f(t) \cos k\omega t$ auch gerade, aber $f(t) \sin k\omega t$ ungerade ist, und für ungerades f das Umgekehrte gilt, folgt

$$\begin{aligned} a_k &= \begin{cases} \frac{4}{T} \int_0^{\frac{T}{2}} f(t) \cos k\omega t dt & \text{falls } f \text{ gerade} \\ 0 & \text{falls } f \text{ ungerade} \end{cases} \\ b_k &= \begin{cases} \frac{4}{T} \int_0^{\frac{T}{2}} f(t) \sin k\omega t dt & \text{falls } f \text{ ungerade} \\ 0 & \text{falls } f \text{ gerade.} \end{cases} \end{aligned}$$

In den Fourierpolynomen einer geraden Funktion treten also nur Cosinusterme auf, in denen einer ungeraden Funktion nur Sinusterme.

Eine allgemeine Funktion (wie $1 + \sin x$) ist natürlich weder gerade noch ungerade!

Beispiel 186 (Rechteckspannung). Wir betrachten die 2π -periodische Funktion (also $\omega = 1$) mit

$$f(t) := \begin{cases} -1 & \text{für } -\pi < t < 0 \\ 1 & \text{für } 0 < t < \pi. \end{cases}$$

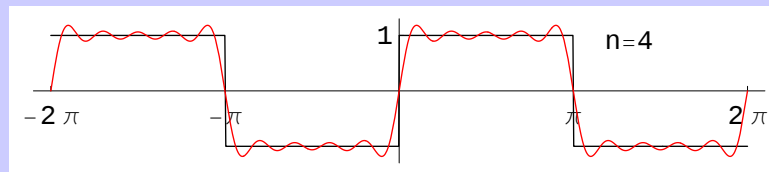
Die Werte an den Stellen $k\pi$ ($k \in \mathbb{Z}$) sind ohne Bedeutung. Setzen wir $f(k\pi) = 0$ für alle k , so wird die Funktion ungerade. Also müssen wir nur die b_k berechnen und finden

$$b_k = \frac{2}{\pi} \int_0^\pi \sin kt \, dt = -\frac{2}{\pi} \frac{\cos kt}{k} \Big|_0^\pi = \frac{2}{\pi k} (1 - (-1)^k) = \begin{cases} \frac{4}{\pi k} & \text{falls } k \text{ ungerade} \\ 0 & \text{sonst} \end{cases}.$$

Es ergibt sich

$$f(t) \approx \frac{4}{\pi} \sum_{m=0}^n \frac{\sin(2m+1)t}{2m+1} = \frac{4}{\pi} \left(\sin t + \frac{1}{3} \sin 3t + \frac{1}{5} \sin 5t + \dots + \frac{1}{2n+1} \sin(2n+1)t \right).$$

So kann man also Rechteckspannungen erzeugen. Die folgende Abbildung zeigt das $2n+1$ -te Fourierpolynom für $n=4$:



Der „Überschuss“ an den Sprungstellen ist ein nach Gibbs benanntes typisches Phänomen bei der Approximation von Sprungfunktionen durch Fourierpolynome. Er beträgt knapp 9% des Sprungs.

Beispiel 187. Die Funktion

$$f(t) = |\sin t|$$

ist eine gerade Funktion, die z.B. bei der Gegentakt-Gleichrichtung von Sinusspannungen auftritt. In den Fourierpolynomen gibt es nur Cosinusterme. Wir berechnen für $k \neq 1$ mit Hilfe des Additionstheorems $2 \sin x \cos y = \sin(y+x) - \sin(y-x)$

$$\begin{aligned} a_k &:= \frac{2}{\pi} \int_0^\pi |\sin t| \cos kt \, dt = \frac{2}{\pi} \int_0^\pi \sin t \cos kt \, dt \\ &= \frac{1}{\pi} \int_0^\pi (\sin(k+1)t - \sin(k-1)t) \, dt = \frac{1}{\pi} \left(-\frac{\cos(k+1)t}{k+1} + \frac{\cos(k-1)t}{k-1} \right) \Big|_0^\pi \\ &= \frac{1}{\pi} \left(-\frac{(-1)^{k+1} - 1}{k+1} + \frac{(-1)^{k-1} - 1}{k-1} \right) \\ &= \begin{cases} \frac{-4}{\pi} \frac{1}{k^2-1}, & \text{falls } k \text{ gerade} \\ 0 & \text{sonst.} \end{cases} \end{aligned}$$

Für $k=1$ erhält man auch 0. Also ist

$$|\sin t| \approx \frac{2}{\pi} - \frac{4}{\pi} \sum_{m=1}^n \frac{\cos 2mt}{4m^2-1} = \frac{2}{\pi} - \frac{4}{3\pi} \cos 2t - \frac{4}{15\pi} \cos 4t - \dots - \frac{4}{(4n^2-1)\pi} \cos 2nt.$$

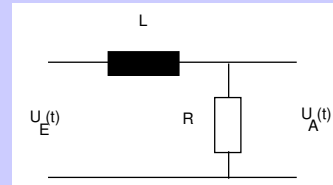
Beispiel 188 (Übertragung eines Sägezahns). Wir betrachten die Übertragung durch ein RL-Glied.

Für die Eingangsspannung

$$U_E(t) = U \sin \omega t$$

finden wir durch Lösen der Differentialgleichung

$$LI' + RI = U_E(t)$$



am Widerstand R den Strom

$$I(t) = e^{-\frac{R}{L}t} \int^t e^{\frac{R}{L}t} \sin(\omega t) dt = \frac{U}{R^2 + \omega^2 L^2} (R \sin(\omega t) - L\omega \cos(\omega t))$$

und damit die Ausgangsspannung

$$U_A(t) = \frac{UR^2}{R^2 + \omega^2 L^2} \sin(\omega t - \phi), \quad \tan \phi = \frac{\omega L}{R}.$$

Wir wollen nun aber eine Sägezahnspannung an den Eingang legen und fragen nach der Ausgangsspannung. Für die Sägezahnfunktion mit Periode $T = \frac{2\pi}{\omega}$ und

$$U_E(t) = \begin{cases} 0 & \text{falls } t = 0 \\ U \left(1 - \frac{2t}{T}\right) & \text{für } 0 < t < T \end{cases}$$

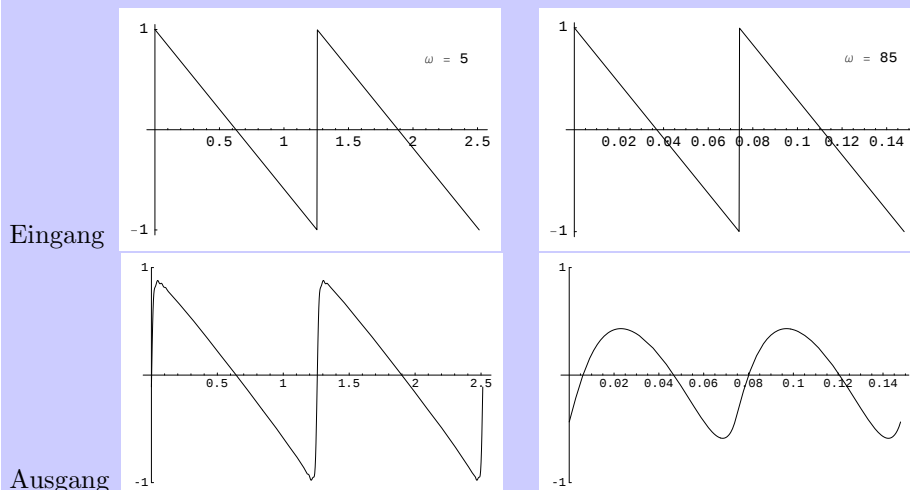
finden wir

$$U_E(t) \approx \frac{2U}{\pi} \sum_{k=1}^n \frac{\sin k\omega t}{k}.$$

Aus der Linearität des RL-Gliedes ergibt sich für den Ausgang

$$U_A(t) \approx \frac{2U}{\pi} \sum_{k=1}^n \frac{R^2}{R^2 + k^2 \omega^2 L^2} \sin(k\omega t - \phi_k), \quad \tan \phi_k = \frac{k\omega L}{R}$$

Für $L = 0$ ist $U_A = U_E$. Für kleines L hat man $U_A \approx U_E$, aber für großes L werden die höheren Frequenzen deutlich abgeschwächt und die Übertragung ist stark verzerrt.

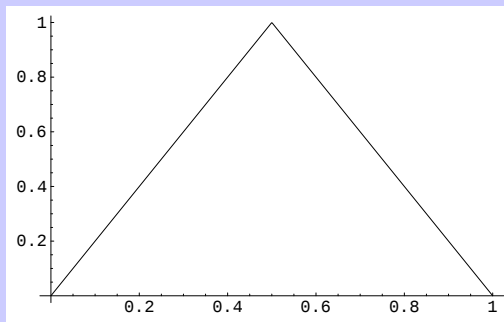


Beispiel 189 (Signalverarbeitung, Rauschunterdrückung). Wir betrachten ein periodisches Signal

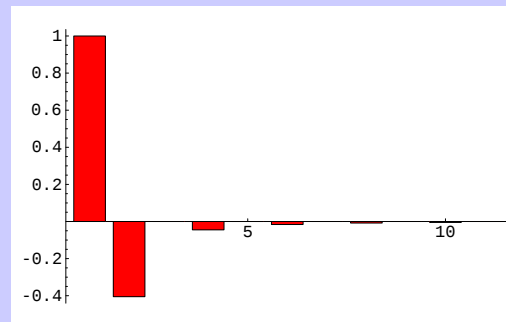
$$f(t) := 1 - |2t - 1|, \quad 0 \leq t \leq 1$$

mit Periode $T = 1$.

Hier ist das Signal ...



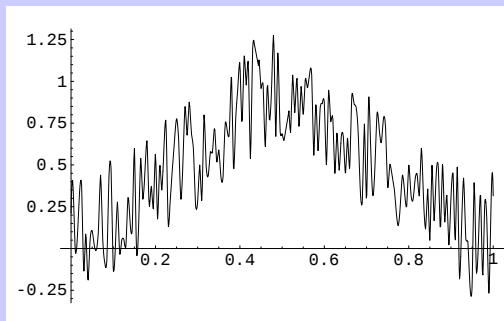
und hier sein Frequenzspektrum.



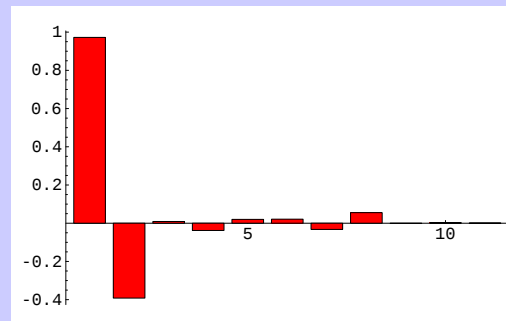
Nun stören wir das Signal durch Überlagerung einer zufälligen Funktion mit Werten zwischen -0.35 und $.35$ (Rauschen)

$$\tilde{f}(t) := 1 - |2t - 1| + 0.7(\text{Random}(t) - \frac{1}{2}).$$

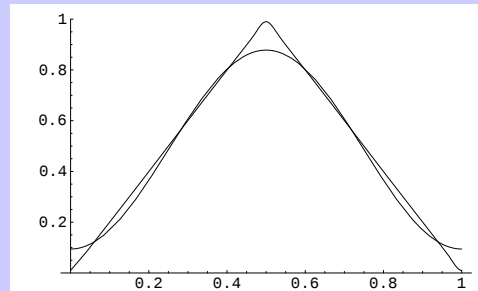
Hier ist das gestörte Signal ...



und hier sein Frequenzspektrum.



Nun filtern wir die Frequenzen mit Amplituden vom Betrag < 0.06 heraus. Die entsprechende Spektralfunktion transformieren wir mit der umgekehrten Fouriertransformation zurück. Die nebenstehende Abbildung zeigt die Rekonstruktion des Signals aus dem ungestörten Spektrum und aus dem gefilterten gestörten Spektrum



Beispiel 190 (Klirrfaktor). Wir betrachten die Verstärkung einer reinen Sinusspannung $U_E(t) = U \sin \omega t$. Ein idealer Verstärker hat eine lineare Kennlinie der Steilheit λ , und die Ausgangsspannung ist

$$U_A(t) = f(U_E) = \lambda U_E(t).$$

Technisch sind aber lineare Kennlinien nur näherungsweise zu realisieren. Im nicht-linearen Fall machen wir eine Fourieranalyse des Ausgangssignals um zu sehen, wie weit es von der angestrebten linearen Verstärkung abweicht. d.h. wir approximieren die Ausgangsspannung:

$$U_A(t) \approx \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t).$$

Wir benutzen nun

$$\begin{aligned} a_1 \cos \omega t + b_1 \sin \omega t &= \sqrt{a_1^2 + b_1^2} \left(\frac{a_1}{\sqrt{a_1^2 + b_1^2}} \cos \omega t + \frac{b_1}{\sqrt{a_1^2 + b_1^2}} \sin \omega t \right) \\ &= \sqrt{a_1^2 + b_1^2} (\sin \phi \cos \omega t + \cos \phi \sin \omega t) = \sqrt{a_1^2 + b_1^2} \sin(\omega t + \phi) \end{aligned}$$

und schreiben

$$U_A(t) \approx \frac{a_0}{2} + \sqrt{a_1^2 + b_1^2} \sin(\omega t + \phi) + \sum_{k=2}^n (a_k \cos k\omega t + b_k \sin k\omega t).$$

Dies interpretieren wir so: Der erste Term ist eine Hintergrund-Gleichspannung, die man ignorieren kann. Der zweite, separat gestellte Term ist die gewünschte Verstärkung

$$\sqrt{a_1^2 + b_1^2} \sin(\omega t + \phi) = \lambda U \sin(\omega t + \phi)$$

der Eingangsspannung um den Faktor $\lambda = \frac{\sqrt{a_1^2 + b_1^2}}{U}$, allerdings mit einer Phasenverschiebung um ϕ . Die restlichen Terme sind unerwünschte Oberschwingungen.

Ein Maß für ihre Größe ist ihre Energie relativ zur Gesamtenergie. Für eine T -periodische Spannung $U(t)$ ist der in einer Periode an einem Ohmschen Widerstand R freigesetzte Energie

$$E = \frac{1}{R} \int_0^T U(t)^2 dt.$$

Für

$$U(t) = \sum_{k=2}^n (a_k \cos k\omega t + b_k \sin k\omega t)$$

folgt aus den Orthogonalitätsrelationen unmittelbar $E = \frac{2}{T} \frac{1}{R} \sum_{k=2}^n (a_k^2 + b_k^2)$ und relativ zur Gesamtenergie

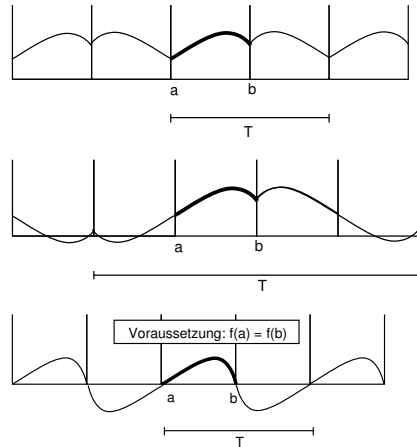
$$K^2 := \frac{\sum_{k=2}^n (a_k^2 + b_k^2)}{\sum_{k=1}^n (a_k^2 + b_k^2)}.$$

Die Wurzel K aus diesem Ausdruck heißt der *Klirrfaktor*. Beachten Sie, dass mit U_A auch die a_k, b_k und damit der Klirrfaktor auf nichtlineare Weise von der Aussteuerung U des Eingangs U_E abhängen.

Beispiel 191 (Periodische Fortsetzung). Im nächsten Beispiel gehen wir auf die Anwen-

derung der Fourieranalyse auf Rand-Anfangswert-Probleme ein. Dazu brauchen wir folgende Überlegung:

Eine beliebige stetige Funktion auf einem kompakten Intervall lässt sich zu einer stetigen periodischen Funktion fortsetzen, und das eventuell auf verschiedene Weise, vgl. die Abbildungen. Damit lässt sich also jede „anständige“ stetige Funktion auf einem Intervall $[a, b]$ durch Fourierpolynome approximieren.



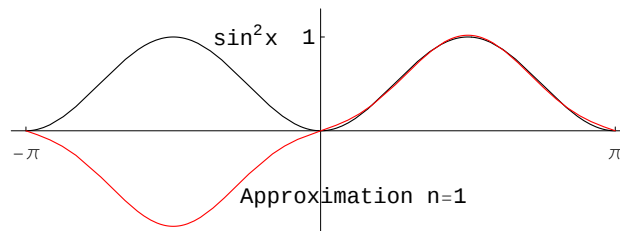
Zum Beispiel ist nach einer elementaren trigonometrischen Identität

$$\sin^2 x = \frac{1}{2} - \frac{1}{2} \cos 2x.$$

Das können wir als Fourierentwicklung auf ganz $\mathbb{R}$ verstehen. Betrachten wir andererseits $f(x) = \sin^2 x$ nur auf dem Intervall $[0, \pi]$, so können wir das mit $f(x) := -\sin^2 x$ auf $[-\pi, 0]$ und dann 2π -periodisch fortsetzen. Die zugehörige Fourierapproximation enthält also nur Sinus(!)terme, und Sie können ausrechnen, dass

$$\sin^2 x \approx \sum_{m=0}^n \frac{-8}{\pi(2m-1)(2m+3)} \sin((2m+1)x) \text{ für } 0 \leq x \leq \pi.$$

(Falls Sie Zweifel mit dem Vorzeichen haben: Nicht alle Koeffizienten sind negativ!) Auf dem Intervall $[0, \pi]$ besitzt $\sin^2 x$ also zwei ganz verschiedene “Fourierapproximationen”. Die erste ist identisch mit $\sin^2 x$ auf ganz $\mathbb{R}$, die zweite approximiert nur auf dem Intervall $[0, \pi]$ (und allen Intervallen $[2k\pi, (2k+1)\pi]$). Die folgende Abbildung zeigt, dass die Approximation bereits für $n = 1$ ausserordentlich gut ist:



□

Beispiel 192 (Rand-Anfangswert-Probleme).

Theoretische Elektrotechnik, Abschnitt 7.1, EIS, Abschnitt 2.3, Popov: Mechanik III, Vorlesung 2

Die Schwingung einer Saite der Länge π mag man sich vielleicht so vorstellen:



und das ausdrücken durch $u(x) = A \sin 5x$. Aber das ist natürlich nur eine Momentaufnahme: eigentlich ist der Bewegungsvorgang abhängig von Ort *und Zeit*, etwa

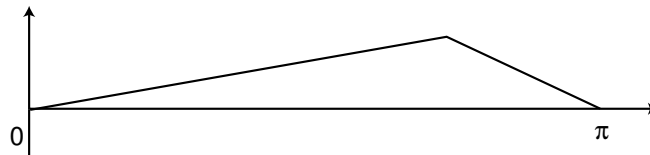
$$u(x, t) = A \sin kx \cos k\omega t,$$

wobei k ganzzahlig sein muss, damit $u(0, t) = u(\pi, t) = 0$, und wobei die Schwingungsgleichung, die wir hier nicht angeben wollen, die zeitliche Frequenz $k\omega$ mit der räumlichen Frequenz k koppelt. Dieselbe Schwingungsgleichung zeigt, dass auch die Überlagerung solcher Schwingungen einen möglichen Verlauf ergibt:

$$u(x, t) = \sum_k A_k \sin kx \cos k\omega t.$$

Dabei kann die Summe endlich oder auch unendlich sein, vgl. den Abschnitt 10 über unendliche Reihen.

Aber die Schwingung muss keineswegs so sinusförmig verlaufen. Die Saite eines Zupfinstruments sieht im Moment des Anreißens etwa so aus:



Um die weitere Bewegung zu beschreiben, muss man eine Überlagerung von Sinusschwingungen mit den richtigen Gewichten (=Amplituden) A_k finden, die für $t = 0$, d.h. wenn man die Cosinusterme = 1 setzt, zum Beispiel diese Form

$$u_0(x) = \begin{cases} mx & \text{für } 0 \leq x \leq 2 \\ 2m(\pi - x)/(\pi - 2) & \text{für } 2 \leq x \leq \pi \end{cases}$$

hat. Also sucht man Koeffizienten A_k , so dass

$$u_0(x) = \sum_k A_k \sin kx.$$

Die Funktion $u_0(x)$ ist eigentlich nicht periodisch, aber wir können sie wie in der dritten Abbildung von Beispiel 191 beschrieben zu einer ungeraden 2π -periodischen Funktion fortsetzen und diese dann wie gewünscht durch eine Sinussumme approximieren.

Genauereres darüber erfahren Sie in den Vorlesungen *Differentialgleichungen für Ingenieure* oder *Integraltransformationen und partielle Differentialgleichungen für Ingenieure*. $\square$

9.3 Approximation im quadratischen Mittel

- Wie gut approximiert die Fourierapproximation?
- Wir lernen sinnvolles Maß für die Approximationsgüte kennen.

Wir haben bei unserer bisherigen Behandlung der Fourierapproximation immer unterstellt, dass die Fourierpolynome einer „anständigen“ periodischen Funktion diese wirklich gut approximieren, dass also das Restglied

$$R_n(t) = f(t) - \left(\frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t) \right)$$

klein ist für großes n . Leider gibt es hier aber keine so handliche Restgliedabschätzung wie im Taylorschen Satz. Deshalb betrachten wir die Frage nach „guter Approximation“ von einer anderen Seite:

Man kann der Ansicht sein, dass so ein trigonometrisches Polynom die untersuchte Funktion gut darstellt, wenn es sich überall nur sehr wenig von ihr unterscheidet. So wie man der Ansicht sein kann, dass eine Messreihe (x_i, y_i) eine gegebene Funktion $f(x)$ sehr gut bestätigt, wenn sie sich überall nur sehr wenig von ihr unterscheidet, d.h. wenn $|f(x_i) - y_i|$ für alle i klein ist.

Aber beim Messen gibt es Messfehler und insbesondere „Ausreißer“, die man eigentlich gern ignorieren würde. Deshalb kann man nicht nur zwangsläufig, sondern nach einigem Nachdenken sogar voller Überzeugung zufrieden sein, wenn die Messwerte „insgesamt“, und das heißt im Mittel über den ganzen Bereich, nah bei der Kurve liegen. Als Maß für diese mittlere Abweichung könnte man

$$\frac{1}{n} \sum_{k=1}^n |f(x_i) - y_i| \quad (69)$$

nehmen, wobei die y_i die Messwerte an den Stellen x_i sind. Vielfach verwendet man aber (aus Gründen, auf die wir gleich eingehen) einen anderen Mittelwert, das *quadratische Mittel*:

$$\sqrt{\frac{1}{n} \sum_{k=1}^n (f(x_i) - y_i)^2} \quad (70)$$

oder einfacher

$$\sum_{k=1}^n (f(x_i) - y_i)^2. \quad (71)$$

Bei der Fourierapproximation wollen wir nicht einen Funktionsverlauf durch diskrete Messungen ermitteln, sondern eine gegebene periodische Funktion $f(t)$ durch ein trigonometrisches Polynom $\phi(t)$ möglichst gut approximieren. An die Stelle der endlichen Summe tritt dann ein Integral, und wir können unsere Zielvorstellung so präzisieren: Wir wollen zu $f(t)$ ein $\phi(t)$ mit vorgegebener Ordnung n finden, so dass das quadratische Mittel der Abweichung

$$\sqrt{\frac{1}{T} \int_0^T (f(t) - \phi(t))^2 dt}$$

möglichst klein wird. Dann sagt man, man hat $f(t)$ im quadratischen Mittel durch trigonometrische Polynome vorgegebener Ordnung optimal approximiert. Natürlich kann man sowohl die Wurzel wie den Faktor $\frac{1}{T}$ dabei auch weglassen, man ändert damit nur die Skala der Fehlerbewertung, das optimale $\phi(t)$, wenn es denn ein solches gibt, bleibt dasselbe.

Warum gibt man der Approximation im quadratischen Mittel – zum Beispiel vor der im betraglichen Mittel – den Vorzug?

- Ganz banal rechnet es sich mit Quadraten wie in (71) einfacher als mit Absolutbeträgen.
- Ein wichtigerer Grund ist, dass (70) bis auf den Faktor $\frac{1}{n}$ eng verwandt ist mit dem üblichen Abstand

$$\sqrt{\sum_{i=1}^3 (x_i - y_i)^2}$$

von Punkten (x_1, x_2, x_3) und (y_1, y_2, y_3) im dreidimensionalen Euklidischen Raum.

Nach diesen Vorbemerkungen allgemeiner Art machen wir uns jetzt an die Approximation periodischer Funktionen im quadratischen Mittel durch trigonometrische Polynome. Wir wählen die reelle Schreibweise, beschränken uns einstweilen aber auf *ungerade* Funktionen, weil das den Schreibaufwand reduziert und die Argumentation übersichtlicher macht.

Sei also $f : \mathbb{R} \rightarrow \mathbb{R}$ eine ungerade Funktion mit der Periode $T = \frac{2\pi}{\omega}$. Wir fixieren eine Zahl n und suchen unter allen ungeraden trigonometrischen Polynomen der Ordnung n

$$\phi(t) = \sum_{k=1}^n \beta_k \sin k\omega t$$

dasjenige, für das

$$\int_0^T (f(t) - \phi(t))^2 dt = \min.$$

Wegen der Orthogonalitätsrelationen ist

$$\frac{2}{T} \int_0^T \phi(t)^2 dt = \sum_{k,l=1}^n \beta_k \beta_l \frac{2}{T} \int_0^T \sin k\omega t \sin l\omega t dt = \sum_{k=1}^n \beta_k^2$$

und, wenn wir wie üblich mit b_k die Fourierkoeffizienten von f bezeichnen,

$$\frac{2}{T} \int_0^T f(t) \phi(t) dt = \sum_{k=1}^n \beta_k \frac{2}{T} \int_0^T f(t) \sin k\omega t dt = \sum_{k=1}^n \beta_k b_k.$$

Wir erhalten

$$\begin{aligned} \frac{2}{T} \int_0^T (f(t) - \phi(t))^2 dt &= \frac{2}{T} \int_0^T f(t)^2 dt - 2 \frac{2}{T} \int_0^T f(t) \phi(t) dt + \frac{2}{T} \int_0^T \phi(t)^2 dt \\ &= \frac{2}{T} \int_0^T f(t)^2 dt - 2 \sum_{k=1}^n \beta_k b_k + \sum_{k=1}^n \beta_k^2 \\ &= \frac{2}{T} \int_0^T f(t)^2 dt - \sum_{k=1}^n b_k^2 + \sum_{k=1}^n (b_k - \beta_k)^2. \end{aligned}$$

Daran kann man sofort erkennen, wann die rechte Seite am kleinsten wird, nämlich wenn $\beta_k = b_k$ für alle k , d.h. wenn

$$\phi(t) = \sum_{k=1}^n b_k \sin k\omega t$$

das Fourierpolynom ist.

Ein entsprechendes Resultat erhält man auch ohne die Voraussetzung, dass f ungerade ist:

Satz 193 (Approximation im quadratischen Mittel). *Sei f eine stückweise monotone Funktion mit Periode $T = \frac{2\pi}{\omega}$ und den Fourierkoeffizienten a_k, b_k . Unter allen trigonometrischen Polynomen der Ordnung n liefert das n -te Fourierpolynom von f*

$$\phi(t) = \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t)$$

die beste Approximation im quadratischen Mittel. Für dieses ist der „quadratische Fehler“

$$\frac{2}{T} \int_0^T (f(t) - \phi(t))^2 dt = \frac{2}{T} \int_0^T f(t)^2 dt - \left(\frac{a_0^2}{2} + \sum_{k=1}^n (a_k^2 + b_k^2) \right).$$

Man kann beweisen, dass dieser Fehler für $n \rightarrow \infty$ gegen 0 konvergiert, d.h. dass man $f(t)$ im quadratischen Mittel beliebig gut durch Fourierpolynome approximieren kann.

9.4 Komplexe Fourieranalysis

- Komplexe Fourieranalysis ist im wesentlichen eine Vereinfachung der Notation mit Hilfe der Eulerschen Formel für $e^{i\phi}$.

Wir haben schon gesehen, dass sich Sinus- und Cosinusfunktionen oft bequemer mittels der Eulerschen Formel durch

$$e^{i\phi} = \cos \phi + i \sin \phi$$

ersetzen lassen. Es ist also

$$\cos \phi = \frac{e^{i\phi} + e^{-i\phi}}{2}, \quad \sin \phi = \frac{e^{i\phi} - e^{-i\phi}}{2i}.$$

Damit lassen sich trigonometrische Polynome sehr einfach komplex schreiben:

$$\begin{aligned} \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos k\omega t + b_k \sin k\omega t) \\ &= \frac{a_0}{2} + \sum_{k=1}^n \left(\frac{a_k}{2} e^{ik\omega t} + \frac{a_k}{2} e^{-ik\omega t} - i \frac{b_k}{2} e^{ik\omega t} + i \frac{b_k}{2} e^{-ik\omega t} \right) \\ &= \underbrace{\frac{a_0}{2}}_{=: c_0} + \sum_{k=1}^n \left(\underbrace{\left(\frac{a_k}{2} - i \frac{b_k}{2} \right)}_{=: c_k} e^{ik\omega t} + \underbrace{\left(\frac{a_k}{2} + i \frac{b_k}{2} \right)}_{=: c_{-k}} e^{-ik\omega t} \right) \\ &= \sum_{k=-n}^{+n} c_k e^{ik\omega t}. \end{aligned}$$

Sind die a_k und b_k die Fourierkoeffizienten einer $T = \frac{2\pi}{\omega}$ -periodischen Funktion, so erhalten wir für $k \geq 0$

$$\begin{aligned} c_k &= \frac{1}{2}(a_k - ib_k) = \frac{1}{T} \int_0^T f(t) \cos k\omega t \, dt - \frac{i}{T} \int_0^T f(t) \sin k\omega t \, dt \\ &= \frac{1}{T} \int_0^T f(t) e^{-ik\omega t} \, dt. \end{aligned}$$

Rechnen Sie nach, dass dieselbe Formel auch für $k < 0$ gilt.

Definieren wir also für $k \in \mathbb{Z}$ die *komplexen Fourierkoeffizienten* der $T = \frac{2\pi}{\omega}$ -periodischen Funktion $f(t)$ durch

$$c_k := \frac{1}{T} \int_0^T f(t) e^{-ik\omega t} \, dt,$$

so schreibt sich die n -te Fourierapproximation von f als

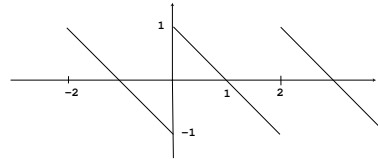
$$f(t) \approx \sum_{k=-n}^{+n} c_k e^{ik\omega t}.$$

Sie sehen, dass die Formeln sehr viel einfacher sind. Der Faktor vor dem Koeffizientenintegral ist $\frac{1}{T}$ statt $\frac{2}{T}$, sogar die Sonderrolle von a_0 ist verschwunden. Für reellwertige Funktionen ist $c_{-k} = \overline{c_k}$, was die Berechnung der Koeffizienten manchmal vereinfacht.

Die Interpretation der reellen Fourierpolynome sollte man auch hier beibehalten: Man hat eine Superposition (= Überlagerung) von – nun komplexwertigen – harmonischen Schwingungen der Frequenzen $k\omega$ mit Amplituden, die durch die Fourierkoeffizienten angegeben werden. Umgekehrt liefert die Folge der Fourierkoeffizienten uns die Information, welche Frequenzen wie stark an der Schwingung $f(t)$ beteiligt sind.

Beispiel 194. Die 2-periodische Funktion mit

$$f(t) = \begin{cases} 1-t & \text{für } 0 < t < 2 \\ 0 & \text{für } t = 0 \end{cases}$$



hat $\omega = 2\pi/2 = \pi$ und die komplexen Fourierkoeffizienten $c_0 = 0$ und

$$\begin{aligned} c_k &= \frac{1}{2} \int_0^2 (1-t) e^{-ik\pi t} dt = \frac{1}{2} \int_0^2 e^{-ik\pi t} dt - \frac{1}{2} \int_0^2 t e^{-ik\pi t} dt \\ &= \underbrace{\frac{e^{-ik\pi t}}{-2ik\pi} \Big|_0^2}_{=0} - \frac{t e^{-ik\pi t}}{-2ik\pi} \Big|_0^2 + \frac{1}{2} \int_0^2 \frac{e^{-ik\pi t}}{-ik\pi} dt = \frac{1}{ik\pi} + \frac{e^{-ik\pi t}}{-2k^2\pi^2} \Big|_0^2 = \frac{1}{ik\pi} \end{aligned}$$

für $k \neq 0$. Damit erhalten wir

$$f(t) \approx \frac{1}{\pi} \sum_{\substack{k=-n \\ k \neq 0}}^{+n} \frac{1}{ik} e^{ik\pi t}.$$

Durch Zusammenfassen der Terme mit negativem und positivem Index folgt

$$\begin{aligned} f(t) &\approx \frac{1}{\pi} \sum_{k=1}^n \frac{1}{ik} e^{ik\pi t} + \frac{1}{\pi} \sum_{k=-n}^{-1} \frac{1}{ik} e^{ik\pi t} = \frac{1}{\pi} \sum_{k=1}^n \frac{1}{ik} e^{ik\pi t} + \frac{1}{\pi} \sum_{k=1}^n \frac{-1}{ik} e^{-ik\pi t} \\ &= \frac{2}{\pi} \sum_{k=1}^n \frac{e^{ik\pi t} - e^{-ik\pi t}}{2ki} = \frac{2}{\pi} \sum_{k=1}^n \frac{\sin k\pi t}{k}. \end{aligned}$$

Natürlich hätte man das auch mit der reellen Fourierreihe erhalten. □

10 Unendliche Reihen

Unendliche Reihen sind Verallgemeinerungen endlicher Summen auf den Fall unendlich vieler Glieder. Von besonderer Bedeutung sind „Polynome von unendlichem Grad“, die sogenannten *Potenzreihen*

$$\sum_{k=0}^{\infty} a_k (x - x_0)^k$$

und „trigonometrische Polynome von unendlicher Ordnung“, die sogenannten *Fourierreihen*

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos k\omega t + b_k \sin k\omega t).$$

Zunächst aber betrachten wir die einfachere Situation unendlicher Reihen, deren Glieder Konstanten sind.

10.1 Reihen mit konstanten Gliedern

- Was sind unendliche Reihen, was sind Partialsummen?
- Wann ist eine unendliche Reihe konvergent?
- Die geometrische Reihe als „allerwichtigste“ konvergente Reihe
- Rechenregeln für Reihen

Eine *unendliche Reihe* hat die Form

$$a_0 + a_1 + a_2 + \dots = \sum_{k=0}^{\infty} a_k,$$

wobei die a_k reelle Zahlen sind. Die endlichen Summen

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

heißen die *Partialsummen* der Reihe.

Definition 195 (Reihenkonvergenz). Wenn die Folge $(s_n)_{n \in \mathbb{N}}$ der Partialsummen gegen A konvergiert, nennt man die unendliche Reihe gegen A *konvergent* und schreibt

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

Man nennt A dann auch *die Summe* der unendlichen Reihe.

Statt $\mathbb{N}$ betrachtet man auch andere Summationsbereiche, z.B. $\sum_{k=1}^{\infty} \frac{1}{k}$.

Beispiel 196 (Die geometrische Reihe). Sei $q \in \mathbb{R}$ und sei $a_k = q^k$. Die zugehörige Reihe

$$\sum_{k=0}^{\infty} q^k = 1 + q + q^2 + q^3 + \dots$$

heißt die **geometrische Reihe**. Sie ist die wichtigste Reihe überhaupt. Die Partialsummen sind gegeben durch

$$s_n = \sum_{k=0}^n q^k = \begin{cases} n+1 & \text{falls } q = 1 \\ \frac{1-q^{n+1}}{1-q} & \text{sonst.} \end{cases}$$

Für $q = 1$ ist die Reihe also divergent. Für $|q| > 1$ ist die Folge $(q^{n+1})_{n \in \mathbb{N}}$ und daher auch die Partialsummenfolge $s_n = \frac{1-q^{n+1}}{1-q}$ divergent. Dasselbe gilt für $q = -1$. Die geometrische Reihe ist also für alle q mit $|q| \geq 1$ divergent. Andererseits ist $\lim_{n \rightarrow \infty} q^{n+1} = 0$ für $|q| < 1$, und deshalb erhalten wir

$$\sum_{k=0}^{\infty} q^k = \frac{1}{1-q} \quad \text{für } |q| < 1.$$

Insbesondere folgt daraus z.B.

$$0,999\dots = \sum_{k=1}^{\infty} 9 \cdot \left(\frac{1}{10}\right)^k = \frac{9}{10} \sum_{k=0}^{\infty} \left(\frac{1}{10}\right)^k = \frac{9}{10} \frac{1}{1 - \frac{1}{10}} = 1.$$

□

Beispiel 197 (Dopplereffekt). Eine ruhende Schallquelle sendet Wellen mit der Frequenz ν , der Schwingungsdauer $T = 1/\nu$ und der Ausbreitungsgeschwindigkeit c aus. Die Wellenlänge ist dann $\lambda = cT = c/\nu$.

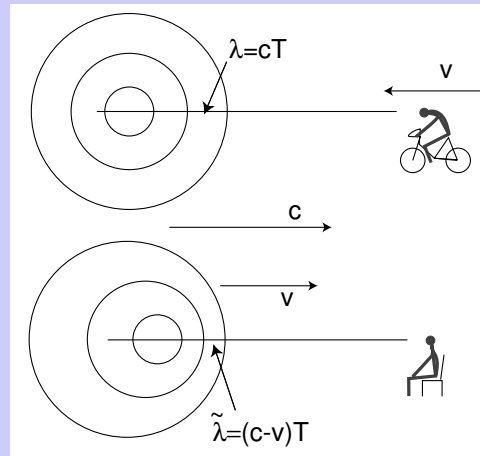
Bewegt sich der Empfänger mit der Geschwindigkeit $v < c$ auf die ruhende Quelle zu, so verkürzt sich die empfangene Frequenz auf

$$\hat{\nu} = \frac{c+v}{\lambda} = \nu \left(1 + \frac{v}{c}\right).$$

Bewegt sich andererseits die Quelle mit der Geschwindigkeit $v < c$ auf den Empfänger zu, so verkürzt sich die Wellenlänge auf $\tilde{\lambda} = (c-v)/\nu = c/\tilde{\nu}$, die empfangene Frequenz ist

$$\tilde{\nu} = \nu \frac{1}{1-v/c} = \nu \left(1 + \frac{v}{c} + \left(\frac{v}{c}\right)^2 + \dots\right).$$

In diesem Fall ist die Frequenz also höher als bei bewegtem Empfänger.



Rechenregeln für konvergente unendliche Reihen. Sind $\sum a_k = A$ und $\sum b_k = B$ konvergente unendliche Reihen, so konvergiert $\sum (a_k + b_k)$ gegen $A + B$:

$$\sum (a_k + b_k) = \sum a_k + \sum b_k. \quad (72)$$

Ist weiter $c \in \mathbb{R}$ so folgt

$$\sum (ca_k) = c \sum a_k. \quad (73)$$

Mit Produkten unendlicher Reihen ist es komplizierter. Man muss jedes Glied der einen Reihe mit jedem Glied der anderen Reihe malnehmen, was bei unendlich vielen Gliedern Probleme macht. Eine einleuchtende Anordnung der Produkte gibt die sogenannte Produktformel von

Cauchy:

$$\left(\sum_{i=0}^{\infty} a_i\right)\left(\sum_{j=0}^{\infty} b_j\right) = \sum_{k=0}^{\infty} c_k, \text{ wobei } c_k := a_0 b_k + a_1 b_{k-1} + \dots + a_k b_0. \quad (74)$$

Die Gleichung stimmt, wenn alle drei Reihen konvergent sind. Aber aus der Konvergenz der Reihen $\sum_{i=0}^{\infty} a_i$ und $\sum_{j=0}^{\infty} b_j$ folgt im allgemeinen nicht die Konvergenz der Reihe $\sum_{k=0}^{\infty} c_k$. Vergleiche dazu aber die Bemerkungen über *absolut konvergente* Reihen weiter unten.

KONVERGENZKRITERIEN.

Wie prüft man, ob eine Reihe konvergiert? Ein einfaches notwendiges Kriterium ist das folgende: Wenn die Reihe $\sum a_k$ konvergiert, wenn also die Partialsummenfolge (s_n) gegen einen Wert A konvergiert, dann konvergiert auch die Folge (s_{n-1}) gegen A , sie hat ja dieselben Glieder, nur in der Nummerierung um eins verschoben. Aber dann gilt $0 = \lim s_n - \lim s_{n-1} = \lim(s_n - s_{n-1}) = \lim a_n$. Nur wenn $\lim a_k = 0$ hat die Reihe also eine Chance, konvergent zu sein:

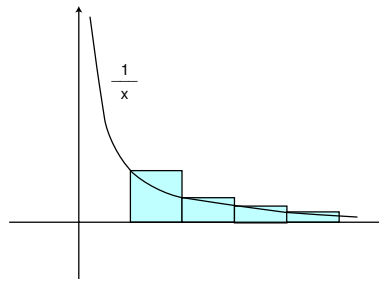
Satz 198 (Notwendiges Konvergenzkriterium). $\sum a_k$ konvergent $\implies \lim a_k = 0$.

Dass aber eine Reihe mit dieser Eigenschaft nicht konvergent sein muss, zeigt das sehr berühmte „Gegenstück“ zur geometrischen Reihe:

Beispiel 199 (Divergenz der harmonischen Reihe). Die harmonische Reihe ist die Reihe

$$\sum_{k=1}^{\infty} \frac{1}{k} = 1 + \frac{1}{2} + \frac{1}{3} + \dots$$

Wir erinnern daran, dass nach Beispiel 173 das uneigentliche Integral $\int_1^{\infty} \frac{1}{x} dx$ nicht existiert. Aus dem Vergleich von Integral und Summe in der Abbildung folgt daher die Divergenz der harmonischen Reihe.



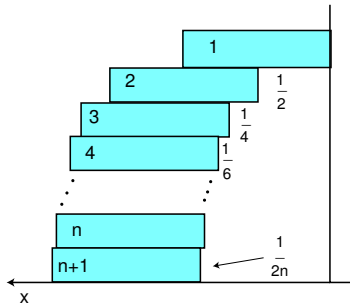
□

Die Divergenz der harmonischen Reihe hat folgende „praktische Anwendung“:

Wir bauen einen Turm aus Ziegelsteinen der Länge 1 „von oben nach unten“, indem wir den bereits gebauten Turm so auf den nächsten Stein setzen, dass sein Schwerpunkt gerade über der Kante des neuen untersten Steins liegt, der Turm also gerade eben nicht umkippt. Ist S_{n-1} die x -Koordinate des Schwerpunktes der ersten $n-1$ Steine, so liegt der Schwerpunkt des n -ten Steins also bei $S_{n-1} + \frac{1}{2}$, und der Schwerpunkt des erweiterten Turms bei

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

Wegen der Divergenz der harmonischen Reihe kann man also den Überhang beliebig groß machen. Zum Beispiel ist $\frac{1}{2} \sum_{k=1}^4 \frac{1}{k} = 1.04$, mit fünf Steinen kann man einen Überhang von mehr als einem Stein und mit 32 Steinen einen von mehr als zwei Steinen realisieren.



Beispiel 200 (Dipolketten).

Werkstoffe I, Abschnitt 2.5

Für $s > 1$ existiert das Integral

$$\int_1^\infty \frac{dx}{x^s}.$$

Am nebenstehenden Bild sieht man deshalb, dass

$$\sum_{k=1}^\infty \frac{1}{k^s}$$

für $s > 1$ konvergiert.

Die für $s > 1$ definierte Funktion

$$\zeta(s) := \sum_{n=1}^\infty 1/n^s$$

heißt die *Riemannsche Zetafunktion*.

In der Theorie der Dielektrika untersucht man eindimensionale Dipolketten. Für jedes $k \in \mathbb{Z} \setminus \{0\}$ sitzt an der Stelle $\vec{x} = k\vec{a}$ ein Dipol mit Dipolmoment $\vec{p} = p \frac{\vec{a}}{|\vec{a}|}$, der im Nullpunkt ein elektrisches Feld der Stärke

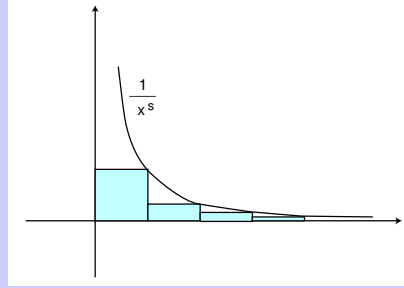
$$\vec{E}_k = \frac{1}{4\pi\epsilon} \left(3 \frac{\vec{p} \cdot \vec{x}}{|\vec{x}|^5} \vec{x} - \frac{1}{|\vec{x}|^3} \vec{p} \right) = \frac{1}{4\pi\epsilon} \left(3 \frac{p}{|k|^3 |\vec{a}|^3} \frac{\vec{a}}{|\vec{a}|} - \frac{1}{|k|^3 |\vec{a}|^3} p \frac{\vec{a}}{|\vec{a}|} \right) = \frac{p}{2\pi\epsilon |\vec{a}|^3} \frac{\vec{a}}{|\vec{a}|} \frac{1}{|k|^3}$$

erzeugt. Das gesamte Feld ist dann

$$\vec{E} = \frac{p}{2\pi\epsilon |\vec{a}|^3} \frac{\vec{a}}{|\vec{a}|} \sum_{k \in \mathbb{Z}, k \neq 0} \frac{1}{|k|^3} = \frac{p}{\pi\epsilon |\vec{a}|^3} \sum_{k=1}^\infty \frac{1}{k^3} = \frac{\zeta(3)p}{\pi\epsilon |\vec{a}|^3} \frac{\vec{a}}{|\vec{a}|} = \frac{1.202p}{\pi\epsilon |\vec{a}|^3} \frac{\vec{a}}{|\vec{a}|}.$$

Die Überlegungen der vorangehenden Beispiele kann man ausbauen zu einem Satz, der eine Beziehung zwischen uneigentlichen Integralen und unendlichen Reihen herstellt:

Satz 201 (Reihen-Integral-Kriterium). Sei $f: [m, \infty[\rightarrow \mathbb{R}$, $m \in \mathbb{N}$ eine monoton fallende Funktion mit $f(x) \geq 0$ für alle $x \in [m, \infty[$. Dann existiert das uneigentliche Integral $\int_m^\infty f(x) dx$ genau dann, wenn die Reihe $\sum_{k=m}^\infty f(k)$ konvergiert.



10.2 Weitere Konvergenzkriterien

Zurück zur Frage: **Wie prüft man, ob eine unendliche Reihe konvergent ist?** Wir wollen dafür jetzt zwei hinreichende Kriterien geben. „Hinreichend“ bedeutet: Wenn dies und das erfüllt ist, *dann* ist die Reihe konvergent. Aber sie kann auch konvergent sein, wenn diese Voraussetzungen nicht vorliegen.

Zunächst vergleichen wir zwei Reihen $\sum a_k$ und $\sum b_k$. Von der zweiten Reihe sei schon bekannt, dass sie konvergiert, und wir wollen annehmen, dass

$$0 \leq a_k \leq b_k \quad \text{für alle } k \quad (75)$$

gilt. Weil die Summanden ≥ 0 sind, sind die Folgen der Partialsummen monoton wachsend und

$$\sum_{k=0}^n a_k \leq \sum_{k=0}^n b_k \leq \sum_{k=0}^{\infty} b_k =: M.$$

Die Partialsummenfolge der Reihe $\sum a_k$ ist also beschränkt, und daher nach dem Monotoniekriterium konvergent.

Dieses Argument kann man nur auf Reihen ohne negative Glieder anwenden. Aber mit einem Trick läßt es sich verallgemeinern: Hat man statt (75) die allgemeinere Bedingung

$$|a_k| \leq b_k \quad \text{für alle } k,$$

so kann man die Reihe $\sum c_k = \sum (a_k + |a_k|)$ betrachten und mit der Reihe $2\sum b_k = \sum 2b_k$ vergleichen. Die zweite Reihe ist natürlich wie die Reihe $\sum b_k$ konvergent, und weil

$$0 \leq a_k + |a_k| \leq 2b_k,$$

ist die Reihe $\sum c_k$ nach unserer gerade gemachten Überlegung konvergent. Natürlich konvergiert danach auch die Reihe $\sum |a_k|$ und damit schließlich die Reihe

$$\sum a_k = \sum ((a_k + |a_k|) - |a_k|) = \sum (a_k + |a_k|) - \sum |a_k|.$$

Wir fassen das zusammen:

Satz 202 (Vergleichskriterium=Majorantenkriterium). *Die Reihe $\sum b_k$ sei konvergent, und es gelte*

$$|a_k| \leq b_k \quad \text{für alle } k.$$

(Hier genügt: Für alle k von einem gewissen $k_0 \in \mathbb{N}$ an.) Dann ist Reihe

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

konvergent.

Beispiel 203. Wir vergleichen die Reihen

$$\sum_{k=1}^{\infty} \frac{1}{k^2} \quad \text{und} \quad \sum_{k=2}^{\infty} \left(\frac{1}{k-1} - \frac{1}{k} \right)$$

Für $k > 1$ ist

$$\left| \frac{1}{k^2} \right| = \frac{1}{k^2} < \frac{1}{k(k-1)} = \frac{1}{k-1} - \frac{1}{k}.$$

Die Konvergenz der zweiten Reihe kann man aber leicht einsehen, weil man ihre Partialsummen explizit bestimmen kann:

$$\sum_{k=2}^n \left(\frac{1}{k-1} - \frac{1}{k} \right) = \left(\frac{1}{1} - \frac{1}{2} \right) + \left(\frac{1}{2} - \frac{1}{3} \right) + \dots + \left(\frac{1}{n-1} - \frac{1}{n} \right) = 1 - \frac{1}{n} \rightarrow 1.$$

Damit haben wir noch einmal die Konvergenz von $\sum_{k=1}^{\infty} \frac{1}{k^2}$ bewiesen, vgl. Beispiel 200. Schwieriger ist, den Grenzwert der Reihe zu bestimmen. Er ist $\zeta(2) = \sum \frac{1}{k^2} = \pi^2/6$. $\square$

Der Nachteil beim Vergleichskriterium ist, dass man schon eine konvergente Majorante $\sum b_k$ haben muss. Das folgende Kriterium benutzt nur die zu untersuchende Reihe und ist deshalb meistens die **erste Wahl**, wenn man eine Reihe auf Konvergenz untersuchen will. Erst wenn das Quotientenkriterium keine Auskunft über die Konvergenz gibt, versucht man andere Kriterien.

Satz 204 (Quotientenkriterium). *Ist die Reihe $\sum_{k=0}^{\infty} a_k$ konvergent? Berechne dazu*

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right|.$$

- *Ist der Grenzwert < 1 , so ist die Reihe konvergent.*
- *Ist der Grenzwert > 1 , so ist die Reihe divergent.*
- *Ist der Grenzwert $= 1$ oder existiert er gar nicht, so gibt dieses Kriterium keine Auskunft.*

Das ist so zu verstehen, dass die a_k alle (oder wenigstens alle von einer Stelle k_0 an) ungleich null sein müssen und der Limes auch existiert.

Beweis. Zum Beweis wählen wir ein q echt zwischen $\lim \left| \frac{a_{k+1}}{a_k} \right|$ und 1:

$$\lim \left| \frac{a_{k+1}}{a_k} \right| < q < 1.$$

Weil $\lim \left| \frac{a_{k+1}}{a_k} \right| < q$ ist dann (nach der Definition der Konvergenz)

$$\left| \frac{a_{k+1}}{a_k} \right| < q \tag{76}$$

für alle k von einem gewissen k_0 an. Daraus folgt für $k \geq k_0$:

$$|a_k| < q|a_{k-1}| < q^2|a_{k-2}| < \dots < q^{k-k_0}|a_{k_0}| = q^k \frac{|a_{k_0}|}{q^{k_0}}.$$

Aber nach dem Resultat über die geometrische Reihe ist

$$\sum_{k=0}^{\infty} q^k \frac{|a_{k_0}|}{q^{k_0}} = \frac{|a_{k_0}|}{q^{k_0}} \sum_{k=0}^{\infty} q^k = \frac{|a_{k_0}|}{q^{k_0}} \frac{1}{1-q}.$$

Die erste Behauptung folgt daher aus dem Vergleichssatz angewendet auf die Reihe $\sum b_k$ mit $b_k = q^k \frac{|a_{k_0}|}{q^{k_0}}$.

Die zweite Behauptung ergibt sich wie folgt: Ist der Limes > 1 , so bedeutet das für große k , dass $|a_{k+1}| > |a_k|$. Also konvergieren die a_k sicher nicht gegen null. $\square$

Bemerkung: Beim Beweis haben wir nicht wirklich die Existenz des Limes benutzt, sondern nur gebraucht, dass es eine Zahl q **kleiner als** 1 gibt, so dass (76) für alle $k \geq k_0$ gilt. Das Quotientenkriterium gilt also auch mit dieser schwächeren Voraussetzung:

Gibt es $q < 1$ und k_0 , so dass $\left|\frac{a_{k+1}}{a_k}\right| \leq q$ für alle $k \geq k_0$, so ist die Reihe $\sum a_k$ konvergent.

Beispiel 205. Die Reihe

$$\sum_{k=0}^{\infty} \frac{k}{2^k}$$

ist konvergent nach dem Quotientenkriterium: Es ist

$$\left|\frac{a_{k+1}}{a_k}\right| = \left|\frac{\frac{k+1}{2^{k+1}}}{\frac{k}{2^k}}\right| = \frac{2^k(k+1)}{2^{k+1}k} = \frac{1}{2}\left(1 + \frac{1}{k}\right) \rightarrow \frac{1}{2} < 1.$$

□

Beispiel 206 (Wo das Quotientenkriterium versagt). Für die beiden Reihen

$$\sum_{k=1}^{\infty} \frac{1}{k} \quad \text{und} \quad \sum_{k=1}^{\infty} \frac{1}{k^2}$$

ist der Limes des Quotienten aufeinanderfolgender Glieder beide Male = 1. Das Quotientenkriterium macht in diesen Fällen keine Aussage. Aber wie wir gesehen haben, ist die erste Reihe divergent, die zweite (zum Beispiel nach dem Vergleichskriterium) konvergent. □

Zum Abschluss betrachten wir noch eine Reihe, für welche die vorstehenden Kriterien ebenfalls versagen, nämlich die sogenannte **alternierende harmonische Reihe**:

$$\sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{k} = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} \pm \dots$$

Diese Reihe ist konvergent. Die Tatsache der Konvergenz kann man relativ leicht einsehen. Bei jedem Schritt wird abwechselnd die Partialsumme erhöht oder erniedrigt, wobei die Differenz monoton kleiner wird. Die ungeraden Partialsummen $1, 1 - \frac{1}{2} + \frac{1}{3}, \dots$ sind daher eine monoton fallende, die geraden $1 - \frac{1}{2}, 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4}, \dots$ eine monoton wachsende Folge. Beide sind beschränkt und daher konvergent. Und weil sich die benachbarten s_n und s_{n+1} nur um $\pm \frac{1}{n+1}$ unterscheiden, konvergieren sie gegen denselben Grenzwert. Also konvergiert die gesamte Folge der Partialsummen. Diese Argumentation liefert ein weiteres Konvergenzkriterium:

Satz 207 (Leibniz-Kriterium). Eine alternierende Reihe deren Glieder absolut genommen eine monoton gegen null konvergierende Folgen bilden, ist konvergent.

Der Grenzwert der obigen Reihe ist wiederum schwierig zu bestimmen. Er ist $\ln 2$.

An dieser Reihe kann man noch ein wichtiges Phänomen deutlich machen: Bei unendlichen Reihen ist die Reihenfolge der Glieder im allgemeinen nicht mehr gleichgültig, es gilt **kein „Kommutativgesetz“**. Das liegt daran, dass eine Umordnung der Glieder die Partialsummenfolge völlig durcheinander bringt. Bei der alternierenden harmonischen Reihe sind nämlich die Reihen der positiven bzw. negativen Glieder

$$1 + \frac{1}{3} + \frac{1}{5} + \dots \quad \frac{1}{2} + \frac{1}{4} + \frac{1}{6} + \dots$$

beide divergent. (Die zweite ist einfach die halbe harmonischen Reihe und deshalb divergent; die Glieder und darum die Partialsummen der ersten Reihe sind aber offensichtlich größer als die der zweiten, weshalb auch die erste Reihe divergent ist.) Darum kann man von der alternierenden Reihe zunächst so viele positive Glieder (der Reihe nach) addieren, bis man z.B. über 27 ist. Dann addiert man negative Glieder solange, bis man unter 27 kommt. Dann addiert man weiter positive Glieder, bis man wieder über 27 ist. Auf diese Weise erwischt man schließlich alle Glieder der Reihe und hat sie so umgeordnet, dass die neue Reihe nun gegen 27 konvergiert.

Diese Zauberei funktionierte, weil man Glieder verschiedenen Vorzeichens hat, die zwar gegen 0 konvergieren, die aber für sich genommen divergente Reihen bilden. Wenn nicht nur die Reihe $\sum a_k$, sondern auch die Reihe $\sum |a_k|$ konvergiert, kann das wohl nicht mehr passieren.

Definition 208 (Absolute Konvergenz). Wenn die Reihe $\sum |a_k|$ konvergiert, nennt man die Reihe $\sum a_k$ *absolut konvergent*. (Nach dem Vergleichskriterium ist sie dann auch im gewöhnlichen Sinne konvergent.)

Vorteile absolut konvergenter Reihen. Vergleichskriterium und Quotientenkriterium liefern *absolut* konvergente Reihen. Diese haben den Vorteil, dass man mit ihnen „ziemlich“ so wie mit endlichen Summen rechnen darf. Genauer:

- Absolut konvergente Reihen darf man beliebig umordnen, die Konvergenz und der Grenzwert ändern sich nicht.
- Das Cauchyprodukt zweier absolut konvergenter Reihen ist wieder (absolut) konvergent, vgl. (74).²⁰

Komplexe Reihen. Für Reihen mit komplexen Gliedern definiert man die Konvergenz genauso wie im reellen Fall. Alle die obigen Regeln – mit Ausnahme des Leibniz-Kriteriums – gelten auch im Komplexen. Zum Beispiel gilt

$$\sum_{k=0}^{\infty} z^k = \frac{1}{1-z} \quad \text{für alle } z \in \mathbb{C} \text{ mit } |z| < 1.$$

²⁰Es genügt, wenn nur eine der beiden Reihen absolut konvergent und die andere konvergent ist.

10.3 Funktionenreihen

- Reihen, deren Glieder nicht konstant, sondern Funktionen sind, nennt man Funktionenreihen.
- Wir lernen die beiden wichtigsten Typen solcher Reihen kennen, die Potenzreihen und die Fourierreihen.

POTENZ- UND TAYLORREIHEN

Die wichtigsten Funktionen erhält man durch eine Verallgemeinerung der Polynome auf unendliche Summen:

Definition 209. Eine *Potenzreihe* ist eine unendliche Reihe der Form

$$\sum_{k=0}^{\infty} a_k (z - z_0)^k. \quad (77)$$

Wir wollen diese Reihen gleich im Komplexen betrachten. Das ist nicht komplizierter, aber die Begriffe „Konvergenzradius“ und „Konvergenzkreis“ werden viel anschaulicher. Wir nehmen deshalb an, dass die Koeffizienten a_k und der Entwicklungspunkt z_0 komplexe Zahlen sind. Weiter ist z eine komplexe Variable.

Auf der Menge $D \subset \mathbb{C}$ aller z , für die die Reihe konvergiert, liefert $f(z) := \sum_{k=0}^{\infty} a_k (z - z_0)^k$ also eine Funktion

$$f : D \rightarrow \mathbb{C}.$$

Wir untersuchen nun die Frage, für welche Werte von z die Reihe konvergiert. Es ist einleuchtend, dass sie konvergiert, wenn $|z - z_0|$ „klein“ ist, und divergiert, wenn diese Zahl „groß“ ist. Dabei hängt die Bedeutung von „klein“ und „groß“ vermutlich von den a_k ab. Wir versuchen das Quotientenkriterium. Beachten Sie aber, dass jetzt $a_k (z - z_0)^k$ die Rolle übernimmt, die im Satz 204 der Term a_k spielte. Wir müssen also

$$\left| \frac{a_{k+1} (z - z_0)^{k+1}}{a_k (z - z_0)^k} \right| = \left| \frac{a_{k+1}}{a_k} \right| \cdot |z - z_0|$$

betrachten.

Wir nehmen an, dass der Grenzwert $A := \lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right|$ existiert.²¹ Dann ist die Reihe für $A|z - z_0| < 1$ nach dem Quotientenkriterium (absolut) konvergent, für $A|z - z_0| > 1$ divergent. Bei $A = 0$ ist die Reihe also für alle $z \in \mathbb{C}$ konvergent, bei $0 < A < \infty$ ist die Reihe für alle z mit

$$|z - z_0| < R := \frac{1}{A} \quad \text{konvergent,} \quad (78)$$

$$|z - z_0| > R := \frac{1}{A} \quad \text{divergent.} \quad (79)$$

Ist $A = \infty$, so ist die Reihe nur für $z = z_0$ konvergent.

Es gilt:

²¹Das muss nicht so sein. Dann kommen wir mit diesem Kriterium nicht weiter.

Satz 210 (Konvergenz von Potenzreihen). Die Potenzreihe

$$\sum_{k=0}^{\infty} a_k (z - z_0)^k$$

ist für alle z in einem offenen Kreis (dem Konvergenzkreis) vom Radius R um den Mittelpunkt z_0 (absolut) konvergent und für alle z außerhalb des abgeschlossenen Kreises divergent. Das Konvergenzverhalten auf dem Rand des Kreises muss man bei jeder speziellen Reihe in jedem Punkt einzeln untersuchen.

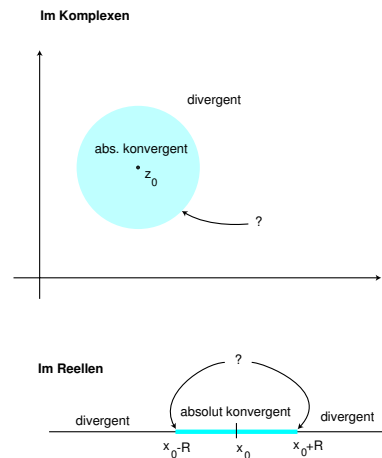
Für den sogenannten Konvergenzradius R gilt

$$R = \lim_{k \rightarrow \infty} \frac{1}{\left| \frac{a_{k+1}}{a_k} \right|},$$

falls der Limes existiert. Dabei ist auch der Wert $R = \infty$ zugelassen, ein Kreis mit unendlichem Radius ist die ganze Ebene $\mathbb{C}$.

Ein „offener Kreis“ ist ein Kreis ohne die berandende Kreislinie, beim „abgeschlossenen Kreis“ ist diese einbegriffen. Beachten Sie, dass der Quotient der a_k hier das Reziproke von der Formel im Quotientenkriterium ist. Beachten Sie weiter, dass wir die obige Aussage nur unter der Annahme bewiesen haben, dass der Grenzwert existiert. Sie bleibt aber richtig, auch wenn der Grenzwert nicht existiert, nur hat man dann keine so einfache Formel mehr für R .

Ein Wort zur Sprache: Sagen Sie nicht, die Reihe sei innerhalb des Konvergenzradius konvergent. Der Konvergenzradius ist eine Zahl, z.B. 7. Was soll es bedeuten, dass die Reihe innerhalb von 7 konvergiert?



Beispiel 211. Die geometrische Reihe $\sum z^k$ ist eine Potenzreihe mit Konvergenzradius $R = 1$. Es gilt

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

□

Beispiel 212. Die Potenzreihe

$$\sum_{k=0}^{\infty} \frac{5^k}{k+1} z^{2k+1} \tag{80}$$

enthält nur ungerade Potenzen, es ist $a_k = 0$ für alle geraden k . Darum sind die Quotienten aufeinanderfolgender Koeffizienten abwechselnd null oder nicht definiert. Der obige Grenzwert existiert also nicht. Aber man kann natürlich das Quotientenkriterium direkt versuchen:

$$\lim_{k \rightarrow \infty} \left| \frac{\frac{5^{k+1}}{k+2} z^{2k+3}}{\frac{5^k}{k+1} z^{2k+1}} \right| = \lim_{k \rightarrow \infty} \left| \frac{5(k+1)z^2}{k+2} \right| = 5|z|^2.$$

Also hat man Konvergenz für $|z| < \frac{1}{\sqrt{5}}$ und Divergenz für $|z| > \frac{1}{\sqrt{5}}$. Der Konvergenzradius ist $R = \frac{1}{\sqrt{5}}$.

Eine andere Möglichkeit mit dieser Reihe umzugehen, benutzt die Umformung

$$\sum_{k=0}^{\infty} \frac{5^k}{k+1} z^{2k+1} = z \sum_{k=0}^{\infty} \frac{5^k}{k+1} (z^2)^k$$

Die Potenzreihe

$$\sum_{k=0}^{\infty} \frac{5^k}{k+1} w^k$$

hat den Konvergenzradius

$$\lim_{k \rightarrow \infty} \frac{\frac{5^k}{k+1}}{\frac{5^{k+1}}{k+2}} = \lim_{k \rightarrow \infty} \frac{5^k(k+2)}{5^{k+1}(k+1)} = \frac{1}{5},$$

konvergiert also für $|w| < \frac{1}{5}$. Die Originalreihe konvergiert daher für $|z^2| < \frac{1}{5}$ also für

$$|z| < \frac{1}{\sqrt{5}}.$$

□

Im folgenden Satz beschränken wir uns wieder auf reelle Reihen.

Satz 213 (Differentiation von Potenzreihen). *Die reelle Potenzreihe*

$$\sum_{k=0}^{\infty} a_k (x - x_0)^k$$

habe Konvergenzradius $R > 0$. Sie konvergiert dann also auf einem Intervall $]x_0 - R, x_0 + R[$ und definiert dort eine Funktion $f(x)$. Diese Funktion ist differenzierbar, und es gilt

$$f'(x) = \sum_{k=1}^{\infty} k a_k (x - x_0)^{k-1} = \sum_{k=0}^{\infty} (k+1) a_{k+1} (x - x_0)^k.$$

Der Konvergenzradius der abgeleiteten Reihe ist wieder R .

Kurz: *Potenzreihen darf man gliedweise differenzieren. Der Konvergenzradius ändert sich dabei nicht. Ebenso darf man Potenzreihen gliedweise integrieren.*

Beispiel 214. Die Reihe $\sum_{k=0}^{\infty} \frac{z^{k+1}}{k+1}$ hat Konvergenzradius $R = 1$. Sie ist

- für $z = 1$ divergent (harmonische Reihe)
- für $z = -1$ konvergent (alternierende harmonische Reihe).

Für reelles $x \in [-1, 1[$ gilt

$$\sum_{k=0}^{\infty} \frac{x^{k+1}}{k+1} = -\ln(1-x). \quad (81)$$

Die Ableitungen der beiden Seiten sind nämlich $\sum_{k=0}^{\infty} x^k$ bzw. $\frac{1}{1-x}$, also gleich. Deshalb sind die beiden Seiten von (81) nach dem Konstanzkriterium gleich bis auf eine additive

Konstante c . Einsetzen von $x = 0$ liefert $c = 0$. Dieses liefert die Gleichheit (81) nur für $|x| < 1$. Sie ist aber auch für $x = -1$ wahr und liefert, wenn man mit (-1) multipliziert, den Wert der alternierenden harmonischen Reihe:

$$\sum_{k=0}^{\infty} \frac{(-1)^k}{k+1} = 1 - \frac{1}{2} + \frac{1}{3} - + \dots = \ln(2).$$

□

In den Abschnitten 6.2 und 6.3 haben wir die Taylorpolynome der trigonometrischen Funktionen und der Exponentialfunktion ausgerechnet und diese Funktionen dadurch approximiert. Die Vermutung liegt nahe, dass man sie durch unendliche Taylorpolynome, d.h. durch die bei $n \rightarrow \infty$ entstehenden Potenzreihen exakt darstellen kann.

Beispiel 215 (Exponentialreihe). Die Potenzreihe

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

hat den Konvergenzradius $R = \infty$. (Nachrechnen!) Also definiert sie eine differenzierbare Funktion $y : \mathbb{R} \rightarrow \mathbb{R}$, für die gilt

$$y(0) = 1 + \frac{0^1}{1!} + \frac{0^2}{2!} + \dots = 1$$

und

$$y'(x) = \sum_{k=0}^{\infty} k \frac{x^{k-1}}{k!} = \sum_{k=1}^{\infty} \frac{x^{k-1}}{(k-1)!} = \sum_{k=0}^{\infty} \frac{x^k}{k!} = y(x).$$

Damit ist der noch offene Existenzbeweis für die Exponentialfunktion geführt:

$$\exp x = \sum_{k=0}^{\infty} \frac{x^k}{k!}.$$

□

Ebenso schließt man für die trigonometrischen Funktionen und erhält deren Reihendarstellungen:

Beispiel 216 (Sinus- und Cosinusreihe).

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

□

Die eben betrachteten Reihen, die Exponentialreihe und die trigonometrischen Reihen, sind Beispiele für sogenannte *Taylorreihen*:

Ist die Funktion $f(x)$ auf einem Intervall um x_0 unendlich oft differenzierbar, so kann man die Potenzreihe

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

bilden, die Taylorreihe zu $f(x)$ im Punkt x_0 .

Wie im Beispiel 112 sieht man: Ist $f(x)$ durch eine Potenzreihe um x_0 gegeben, so ist diese auch gerade die Taylorreihe im Punkt x_0 .

Aber im allgemeinen ist es nicht klar, ob das Restglied für $n \rightarrow \infty$ wirklich gegen 0 geht, d.h. ob die Taylorreihe von f wirklich gegen f konvergiert. Und es gibt Funktionen, bei denen das falsch ist.

FOURIERREIHEN

Wir betrachten eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ mit Periode $T = \frac{2\pi}{\omega}$ mit Fourierkoeffizienten a_k und b_k . Dann heißt

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos k\omega t + b_k \sin k\omega t)$$

die *Fourierreihe* von $f(t)$.

Satz 217 (Konvergenz von Fourierreihen²²). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ sei $T = \frac{2\pi}{\omega}$ -periodisch und stückweise monoton. Dann gilt:

- An allen Stetigkeitsstellen t von f konvergiert die Fourierreihe gegen $f(t)$.
- An den Unstetigkeitsstellen existieren wegen der stückweisen Monotonie der links- und rechtsseitige Grenzwert $f(t-)$ und $f(t+)$ von f , und die Fourierreihe konvergiert gegen das daraus gebildete arithmetische Mittel.

Man hat also für alle t :

$$\frac{f(t-) + f(t+)}{2} = \frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos k\omega t + b_k \sin k\omega t).$$

Beispiel 218. Die 2-periodische Funktion

$$f(t) = \begin{cases} 1-t & \text{für } 0 < t < 2 \\ 0 & \text{für } t = 0 \end{cases}$$

aus Beispiel 194 hat nach Definition an den Sprungstellen gerade den Mittelwert 0 als Funktionswert. Daher wird sie überall durch ihre Fourierreihe dargestellt:

$$f(t) = \sum_{k=1}^{\infty} \frac{2}{\pi k} \sin k\pi t.$$

□

Beispiel 219. Ähnlich erhält man für die 2π -periodische Funktion mit

$$g(t) := \begin{cases} -1 & \text{für } -\pi < t < 0 \\ 1 & \text{für } 0 < t < \pi. \\ 0 & \text{für } t = 0, \pi \end{cases}$$

(vgl. Beispiel 186), dass

$$g(t) = \frac{4}{\pi} \sum_{k=0}^{\infty} \frac{1}{2k+1} \sin(2k+1)t.$$

□

Aus der letzten Behauptung von Satz 193 über die Approximation im quadratischen Mittel folgt

²²Einen Beweis für den Satz in dieser bequemen, aber nicht sehr verbreiteten Version findet man etwa in *Mangoldt-Knopp, Einführung in die Höhere Mathematik, Band III*.

Satz 220 (Parsevalsche Gleichung). Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ sei $T = \frac{2\pi}{\omega}$ -periodisch und stückweise monoton mit Fourierkoeffizienten a_k und b_k . Dann gilt:

$$\frac{2}{T} \int_0^T f(t)^2 dt = \frac{a_0^2}{2} + \sum_{k=1}^{\infty} (a_k^2 + b_k^2).$$

Beispiel 221. Wir betrachten die beiden Funktionen aus den vorangehenden Beispielen 218 und 219. Wir erhalten

$$\sum_{k=1}^{\infty} \frac{4}{\pi^2} \frac{1}{k^2} = \frac{2}{2} \int_0^2 (1-t)^2 dt = - \left. \frac{(1-t)^3}{3} \right|_0^2 = \frac{2}{3}.$$

Also ist

$$\sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{\pi^2}{6}.$$

Die Konvergenz der Reihe $\sum \frac{1}{k^2}$ haben wir früher mit dem Reihen-Integral-Kriterium bzw. mit dem Majorantenkriterium nachgewiesen. Jetzt haben wir sozusagen zufällig auch den Reihenwert bestimmt.

Ebenso ergibt sich

$$\frac{16}{\pi^2} \sum_{k=0}^{\infty} \left(\frac{1}{2k+1} \right)^2 = \frac{2}{2\pi} \int_{-\pi}^{+\pi} g(t)^2 dt = 2,$$

und damit

$$\sum_{k=0}^{\infty} \left(\frac{1}{2k+1} \right)^2 = \frac{\pi^2}{8}.$$

□

Wie Potenzreihen, so darf man auch Fourierreihen gliedweise integrieren, aber man darf sie i.a. nicht gliedweise differenzieren: Nicht nur die Sprungstellen machen Probleme, wie man denken könnte. Nach dem Beispiel 218 ist

$$1-t = 2 \sum_{k=1}^{\infty} \frac{\sin k\pi t}{k\pi} \quad \text{für } 0 < t < 2.$$

Durch gliedweises Differenzieren der Reihe bekommt man die Reihe $2 \sum_{k=1}^{\infty} \cos k\pi t$, und an der Stelle $t = 1$, an der f keinen Sprung hat, ist die Reihe $2 \sum_{k=1}^{\infty} \cos k\pi = 2 \sum_{k=1}^{\infty} (-1)^k$ gar nicht konvergent. Dasselbe gilt zum Beispiel für alle irrationalen Werte von t .

Konvergente Fourierreihen darf man im allgemeinen NICHT gliedweise differenzieren.

11 Anhang

11.1 Inverse der RSA-Kodierung

Wir zeigen unter Verwendung von ein bißchen elementarer Zahlentheorie, dass die im Beispiel 8 angegebene Abbildung $x \mapsto x^{17} \pmod{2773}$ injektiv ist.

Wir verdeutlichen damit einen von ungezählten mathematischen Sachverhalten, denen sie in Ihrem Alltag ständig begegnen, ohne sich darüber Rechenschaft abzulegen.

Seien p und q zwei Primzahlen und $n = pq$. Sei weiter $e \in \mathbb{N}$ teilerfremd zu $f := (p-1)(q-1)$. Dann gibt es nach elementarer Teilbarkeitstheorie ganze Zahlen $a, b \in \mathbb{Z}$ mit

$$ae + bf = 1.$$

Diese Zahlen lassen sich mittels euklidischem Algorithmus leicht ermitteln, wenn man p und q , also e und f kennt. Die Faktorisierung von n in zwei Primzahlen ist der „Knackpunkt“ bei der RSA-Codierung.

Für alle $x \in \mathbb{N}$ ist

$$x^{ae+bf} = x^{ae}(x^{b(q-1)})^{p-1} = x^1.$$

Falls p ein Teiler von x ist, folgt daraus

$$x^{ae} = 0 = x \pmod{p}.$$

Falls p kein Teiler x ist, ist es auch keiner von $x^{b(q-1)}$, und daher gilt nach dem *Kleinen Satz von Fermat*

$$(x^{b(q-1)})^{p-1} = 1 \pmod{p}.$$

Also ist in jedem Falle

$$x^{ae} = x \pmod{p}$$

und natürlich ebenso

$$x^{ae} = x \pmod{q}.$$

Aus beidem zusammen folgt

$$x^{ae} = x \pmod{pq}.$$

Das bedeutet $(x^e)^a = x \pmod{n}$, die Abbildungen

$$x \mapsto x^e \pmod{n}$$

und

$$x \mapsto x^a \pmod{n}$$

sind auf $\{0, 1, \dots, n-1\}$ invers zueinander.

In den Beispielen 8 und 13 ist $p = 59, q = 47, e = 17, a = 157$.

11.2 Komplexe Zahlen in der Elektrotechnik: Beispiele

Wir geben noch einige Beispiel für die Bedeutung komplexer Zahlen in der Elektrotechnik

Beispiel 222 (Komplexer Widerstand).

Grundlagen der Elektrotechnik, Teil A, 7. Übung

Für einen komplexen Widerstand Z nennt man

- $\operatorname{Re} Z$ den Wirkwiderstand,
- $\operatorname{Im} Z$ den Blindwiderstand und
- $|Z|$ die Impedanz oder den Scheinwiderstand.

Beispiel 223 (Debye-Gleichungen).

Werkstoffe I, Abschnitt 2.8

Bei der Polarisation von Dielektrika im Wechselfeld verwendet man die komplexe Dielektrizitätskonstante²³

$$\underline{\epsilon} = 1 + \frac{(\epsilon_{r,s} - 1)}{1 + i\omega\tau}$$

Die Zerlegung von $\underline{\epsilon} = \epsilon' + i\epsilon''$ führt zu den sogenannten Debye-Gleichungen:

$$\begin{aligned}\epsilon' + i\epsilon'' &= \frac{\epsilon_{r,s} + i\omega\tau}{1 + i\omega\tau} = \frac{(\epsilon_{r,s} + i\omega\tau)(1 - i\omega\tau)}{(1 + i\omega\tau)(1 - i\omega\tau)} \\ \Rightarrow \\ \epsilon' &= \frac{\epsilon_{r,s} + \omega^2\tau^2}{1 + \omega^2\tau^2}, \quad \epsilon'' = \frac{(1 - \epsilon_{r,s})\omega\tau}{1 + \omega^2\tau^2}.\end{aligned}$$

Beispiel 224 (Wechselstrom-Netzwerke).

Grundlagen der Elektrotechnik, Teil A, 7. Übung

Der Zusammenhang von Spannung und Strom in einem Gleichstromkreis mit Widerstand R wird durch das Ohmsche Gesetz $U = IR$ geliefert, und bei verzweigten Stromkreisen (Netzwerken) ergeben sich die Widerstände aus den Kirchhoffschen Gesetzen: Bei Hintereinanderschaltung addieren sich die Widerstände, bei Parallelschaltung addieren sich deren Reziproke: $\frac{1}{R} = \frac{1}{R_1} + \frac{1}{R_2}$. Für Wechselstromkreise mit Ohmschen Widerständen, Spulen und Kondensatoren sehen die entsprechenden Gesetze grauslich kompliziert aus, – wenn man sie reell schreibt. Man kann sie aber durch Verwendung komplexer Zahlen enorm vereinfachen. Man betrachtet die Spannung $U = U_0 \cos \omega t$ als Realteil einer „komplexen Spannung“

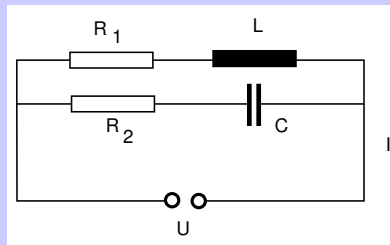
$$U_{\text{komplex}} := U_0(\cos \omega t + i \sin \omega t).$$

Analog verfährt man mit dem Strom und dem Widerstand. Der (frequenzabhängige) komplexe Widerstand ist dabei wie folgt gewählt:

- ein Ohmscher Widerstand R hat den komplexen Widerstand $Z = R$.
- eine Induktivität L hat den komplexen Widerstand $Z = i\omega L$.
- eine Kapazität C hat den komplexen Widerstand $Z = \frac{1}{i\omega C}$.

Mit diesen Vereinbarungen gelten die Gesetze von Ohm und Kirchhoff auch für Wechselstromkreise!

Konkret betrachten wir einen komplexen Wechselstromkreis mit $U = U_0 \cos \omega t$.



Dann gilt $I_{\text{komplex}} = \frac{U_{\text{komplex}}}{Z}$, wobei der komplexe Widerstand Z sich wie folgt errechnet:

$$\frac{1}{Z} = \frac{1}{R_1 + i\omega L} + \frac{1}{R_2 + \frac{1}{i\omega C}}$$

Das ist eine einfache und leicht zu merkende Sache. Wenn man den reellen Strom wirklich haben will, muss man allerdings rechnen. Das Ergebnis der Rechnung kann man sich nicht mehr merken:

$$I = U_0 \frac{\sqrt{(R_1 + R_2)^2 + (\omega L - \frac{1}{\omega C})^2}}{\sqrt{R_1^2 + \omega^2 L^2} \sqrt{R_2^2 + \frac{1}{\omega^2 C^2}}} \cos(\omega t - \phi),$$

$$\phi = \arctan \frac{1 - \omega^2 LC}{\omega C(R_1 + R_2)} + \arctan \frac{\omega^2 LCR_2 - R_1}{\omega CR_1 R_2 + \omega L}.$$

Man muss dazu vor allem den Quotienten komplexer Zahlen berechnen und wieder in der Form $x + iy$ darstellen. Darauf gehen wir gleich ein. Weiter ist die sogenannte Polarkoordinaten-Darstellung komplexer Zahlen hilfreich. Diese behandeln wir im nächsten Abschnitt.

11.3 Konvergenz numerisch

Wie wichtig sind Konvergenzbeweise wie im Abschnitt 3.1? Im Zeitalter der programmierbaren Rechner kann man blitzartig die ersten 1000 Glieder einer Folge ausrechnen und sieht oft sogar schon viel eher, ob sie konvergent ist, vgl. Beispiel 41 oben. Aber endlich viele Glieder – gleich wieviele – bilden immer nur einen verschwindend kleinen Abschnitt der unendlichen Folge. Die Frage der Konvergenz auf der Basis dieser unvollständigen Information entscheiden zu wollen ist deshalb sehr kühn, wie das nächste Beispiel zeigt.

Beispiel 225. Wir untersuchen die Folge

$$x_0 = 1, x_1 = \left(\frac{p-1}{p}\right)^{2k}, \quad x_{n+1} = (2 \sqrt[2k]{x_n} - \sqrt[2k]{x_{n-1}})^{2k}$$

für $k = 6, p = 100$ mit Mathematica.

```
In[1]:= k=6; p=100;
a=1; myc=1;

Print["0: ",1]
Print["1: ",N[b=((p-1)/p)^(2 k)]]
Do[z=(2 Abs[b]^(1/(2k))- Abs[a]^(1/(2k)))^(2k);
a=b ;b=z; myc=myc+1; Print[myc,": ",N[z]],{101}]
```

Hier der Output:

0 :	1	60 :	0.0000167772	83 :	5.82622 10 ⁻¹⁰
1 :	0.886385	61 :	0.0000123816	84 :	2.81475 10 ⁻¹⁰
2 :	0.784717	62 :	9.06574 10 ⁻⁶	85 :	1.29746 10 ⁻¹⁰
3 :	0.693842	63 :	6.58295 10 ⁻⁶	86 :	5.66939 10 ⁻¹¹
4 :	0.61271	64 :	4.73838 10 ⁻⁶	87 :	2.32981 10 ⁻¹¹
5 :	0.54036	65 :	3.37922 10 ⁻⁶	88 :	8.9161 10 ⁻¹²
6 :	0.47592	66 :	2.38642 10 ⁻⁶	89 :	3.13843 10 ⁻¹²
7 :	0.418596	67 :	1.66789 10 ⁻⁶	90 :	1. 10 ⁻¹²
8 :	0.367666	68 :	1.15292 10 ⁻⁶	91 :	2.8243 10 ⁻¹³
9 :	0.322475	69 :	7.87663 10 ⁻⁷	92 :	6.87195 10 ⁻¹⁴
10 :	0.28243	70 :	5.31441 10 ⁻⁷	93 :	1.38413 10 ⁻¹⁴
11 :	0.24699	71 :	3.53815 10 ⁻⁷	94 :	2.17678 10 ⁻¹⁵
12 :	0.215671	72 :	2.32218 10 ⁻⁷	95 :	2.44141 10 ⁻¹⁶
13 :	0.188032	73 :	1.50095 10 ⁻⁷	96 :	1.67772 10 ⁻¹⁷
14 :	0.163675	74 :	9.5429 10 ⁻⁸	97 :	5.31441 10 ⁻¹⁹
15 :	0.142242	75 :	5.96046 10 ⁻⁸	98 :	4.096 10 ⁻²¹
16 :	0.12341	76 :	3.65203 10 ⁻⁸	99 :	1. 10 ⁻²⁴
17 :	0.10689	77 :	2.19146 10 ⁻⁸	100 :	0
...		...			

Sieht konvergent aus, ist es aber nicht. Die nächsten beiden Glieder sind

$$\begin{aligned} 101 : & \quad 1. 10^{-24} \\ 102 : & \quad 4.096 10^{-21}. \end{aligned}$$

Tatsächlich kann man zeigen, dass für natürliche Zahlen p

$$x_n := \left(\frac{p-n}{p}\right)^{2k},$$

und diese Folge ist offensichtlich divergent. Der “Konvergenztest mit dem Rechner“ ist also **nicht verlässlich**. Trotzdem kann der Rechner nützliche Hinweise geben (z.B. auf ein bestimmtes Monotonieverhalten), denen man dann aber rigoros nachgehen muss: Der Rechner liefert Vermutungen, keine Beweise. $\square$

11.4 Lineare Differentialgleichungen mit konstanten Koeffizienten

Eine lineare Differentialgleichung mit konstanten Koeffizienten ist eine Differentialgleichung der Form

$$a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_1 y' + a_0 y = b(x). \quad (82)$$

Dabei sind $a_0, \dots, a_n$ reelle oder komplexe Zahlen.

Die Differentialgleichung heißt vom Grad n , wenn der höchste Koeffizient $a_n \neq 0$. Oft normiert man die Gleichung so, dass $a_n = 1$. In den Anwendungen kommen sehr häufig Differentialgleichung vom Grad 2 vor.

Beispiel 226. Die Schwingungsgleichung $m\ddot{x} = -kx$ eines Federpendels kann man schreiben als

$$\ddot{x} + \frac{k}{m}x = 0$$

oder als

$$\ddot{x} + \frac{a}{m} + \frac{k}{m}x = 0,$$

wenn Reibung im Spiel ist.

In der Elektrotechnik wird der Strom I in einem RLC-Schwingkreis beschrieben durch die Differentialgleichung

$$L\ddot{I} + R\dot{I} + \frac{1}{C}I = 0.$$

Ist in dem Kreis eine (zeitlich variable) Spannungsquelle $U(t)$ vorhanden, so erhält man die Gleichung

$$L\ddot{I} + R\dot{I} + \frac{1}{C}I = U(t).$$

Es ist sinnvoll, zwei Fälle von (82) zu unterscheiden:

- Wenn $b = 0$ die Nullfunktion ist, nennt man die Gleichung *homogen*.
- Andernfalls nennt man die Gleichung *inhomogen*.

HOMOGENE LINEARE DIFFERENTIALGLEICHUNGEN

Diese haben wir für den Fall $n = 2$ schon im Abschnitt über die komplexen Zahlen behandelt.

Für die Gleichung

$$a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_1 y' + a_0 y = 0 \quad (83)$$

führt der Exponentialansatz $y(x) = e^{\lambda x}$ auf die zugehörige *charakteristische* Gleichung

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

Jede Nullstelle λ produziert eine Lösung $e^{\lambda x}$. Ist λ eine k -fache Nullstelle, so produziert es k verschiedene Lösungen

$$e^{\lambda x}, x e^{\lambda x}, \dots, x^{k-1} e^{\lambda x}.$$

Die Linearkombinationen dieser so gefundenen insgesamt n Lösungen sind die sämtlichen Lösungen von (83).²⁴ Eine *homogene* lineare Differentialgleichung vom Grad n hat n linear unabhängige Lösungen.

²⁴Allerdings können die Nullstellen natürlich komplex sein, auch wenn die Koeffizienten reell sind. Wir erhalten dann *komplexe Lösungen*. Wir gehen hier nicht auf die Frage ein, wie man daraus reelle Lösungen gewinnt. Vergleichen Sie den Abschnitt über die komplexen Zahlen und die Module über Differentialgleichungen.

Bei n verschiedenen Nullstellen sieht das also so aus:

$$y(x) = c_1 e^{\lambda_1 x} + \dots + c_n e^{\lambda_n x}.$$

Den Fall einer doppelten Nullstelle beschreibt das folgende

Beispiel 227. Die charakteristische Gleichung von

$$y'' + 6y' + 9y = 0$$

hat -3 als doppelte Nullstelle. Die allgemeine Lösung ist

$$y(x) = c_1 e^{-3x} + c_2 x e^{-3x}.$$

□

Wenn man an einer Stelle x_0 Anfangswerte in Gestalt der Ableitungen bis zur Ordnung $n-1$ vorgibt, also Werte für

$$y(x_0), \dots, y^{(n-1)}(x_0),$$

so sind dadurch die Koeffizienten c_j bestimmt, das *Anfangswertproblem* besitzt eine eindeutige Lösung.

Beispiel 228. Wir bleiben bei der Differentialgleichung aus dem letzten Beispiel. Verlangt man $y(0) = 2, y'(0) = -1$, so ergibt das ein lineares Gleichungssystem für c_1, c_2 :

$$\begin{aligned} y(0) &= c_1 + c_2 = 2 \\ y'(0) &= -3c_1 + c_2 = -1 \end{aligned}$$

Man findet $c_1 = \frac{3}{4}, c_2 = \frac{5}{4}$ und die Lösung

$$y(x) = \frac{1}{4}(3 + 5x)e^{-3x}.$$

□

INHOMOGENE LINEARE DIFFERENTIALGLEICHUNGEN

Wenn die rechte Seite $b(x)$ nicht die Nullfunktion ist, nennt man die Differentialgleichung inhomogen. Manchmal kann man eine Lösung raten. Wenn man zum Beispiel in die linke Seite ein Polynom $y(x)$ einsetzt, kommt offenbar ein Polynom heraus. Das kann man benutzen, wenn die rechte Seite als Polynom vorgegeben ist: Man setzt für $y(x)$ ein allgemeines Polynom an und bestimmt dann die Koeffizienten (durch Koeffizientenvergleich) so, dass die richtige rechte Seite herauskommt.²⁵

Beispiel 229. Wir betrachten

$$y'' + 3y' + y = x^2 - 7. \quad (84)$$

Wir versuchen den Ansatz $y(x) = Ax^2 + Bx + C$. Einsetzen liefert

$$2A + 3(2Ax + B) + Ax^2 + Bx + C = x^2 - 7$$

oder

$$Ax^2 + (6A + B)x + 2A + 3B + C = x^2 - 7.$$

²⁵ Ähnlich kann man zum Beispiel vorgehen, wenn die rechte Seite von der Form $a \cos \omega x + b \sin \omega x$ ist. Fallen Ihnen andere Situationen ein, für die Sie auf diese Weise Lösungen „raten“ können?

Also erhalten wir eine Lösung, wenn $A = 1$, $B = -6$ und $C = 9$ ist. Die Funktion

$$\bar{y}(x) = x^2 - 6x + 9$$

ist eine Lösung. Aber sie ist nicht die einzige Lösung der Differentialgleichung! Hat man eine Funktion $y_h(x)$ mit

$$y_h'' + 3y_h' + y_h = 0,$$

so ist auch $y(x) + y_h(x)$ eine Lösung von (84). Warum? $\square$

Der zuletzt geschilderte Sachverhalt ist ein allgemeiner: Hat man (zum Beispiel durch intelligentes Raten) eine spezielle Lösung y_{speziell} der inhomogenen Differentialgleichung (82) gefunden, so ergibt sich die allgemeine Lösung von (82), indem man zu y_{speziell} alle Lösungen der *zugehörigen homogenen Gleichung* (83) addiert.

Beispiel 230. Wir schließen an das letzte Beispiel an, und suchen eine Lösung von

$$y'' + 3y' + y = x^2 - 7$$

mit den Anfangsbedingungen $y(0) = 0$, $y'(0) = 5$.

Die gefundene Lösung $y(x) = x^2 - 6x + 9$ hat offenbar andere Anfangsbedingungen. Wir lösen daher die zugehörige homogene Differentialgleichung

$$y'' + 3y' + y = 0.$$

Die Wurzeln der charakteristischen Gleichung $\lambda^2 + 3\lambda + 1 = 0$ sind

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

Daher ist die allgemeine Lösung der inhomogenen Gleichung gegeben durch

$$y(x) = x^2 - 6x + 9 + c_1 e^{(-\frac{3}{2} + \frac{1}{2}\sqrt{5})x} + c_2 e^{(-\frac{3}{2} - \frac{1}{2}\sqrt{5})x}.$$

Wir finden

$$y(0) = 9 + c_1 + c_2 = 0,$$

$$y'(0) = -6 + c_1(-\frac{3}{2} + \frac{1}{2}\sqrt{5}) + c_2(-\frac{3}{2} - \frac{1}{2}\sqrt{5}) = 5.$$

Auflösen dieses linearen Gleichungssystems liefert $c_{1,2} = -\frac{1}{2}(9 \mp \sqrt{5})$. Die gesuchte Lösung ist also

$$y(x) = x^2 - 6x + 9 - \frac{1}{2}(9 - \sqrt{5})e^{(-\frac{3}{2} + \frac{1}{2}\sqrt{5})x} - \frac{1}{2}(9 + \sqrt{5})e^{(-\frac{3}{2} - \frac{1}{2}\sqrt{5})x}.$$

$\square$

11.5 Regel von Bernoulli/de l' Hospital

Die Regel von Bernoulli/de l' Hospital ist eine sehr praktische Methode zur Berechnung gewisser Grenzwerte. Zur Vorbereitung zeigen wir eine Verallgemeinerung des Mittelwertsatzes:

Satz 231 (Verallgemeinerter Mittelwertsatz). Seien $f, g : I \rightarrow \mathbb{R}$ differenzierbar auf dem Intervall $I \subset \mathbb{R}$. Die Ableitung von g sei nirgends 0. Sind $a, b \in I$ beliebig mit $a \neq b$, so ist $g(a) \neq g(b)$, und zwischen a und b gibt es ein ξ mit

$$\frac{f(b) - f(a)}{g(b) - g(a)} = \frac{f'(\xi)}{g'(\xi)}. \quad (85)$$

(Für $g(x) = x$ erhält man wieder den alten Mittelwertsatz.)

Beweis. Nach dem Mittelwertsatz ist $g(a) \neq g(b)$, sonst hätte die Ableitung dazwischen eine Nullstelle.

Betrachten Sie die Funktion

$$h(x) = (f(b) - f(a))g(x) - (g(b) - g(a))f(x).$$

Rechnen Sie nach, dass $h(a) = h(b)$. Also gibt es nach dem Mittelwertsatz ein ξ zwischen a und b mit

$$h'(\xi) = \frac{h(b) - h(a)}{b - a} = 0.$$

Aber

$$h'(\xi) = (f(b) - f(a))g'(\xi) - (g(b) - g(a))f'(\xi),$$

und nach Division mit $g(b) - g(a) \neq 0$ folgt die Behauptung. $\square$

Satz 232 (Regel von Bernoulli/de l' Hospital). Seien $I \subset \mathbb{R}$ ein Intervall, $f, g : I \rightarrow \mathbb{R}$ differenzierbar und $x_0 \in \mathbb{R} \cup \{-\infty, +\infty\}$. Es gelte

$$\lim_{x \rightarrow x_0} f(x) = \lim_{x \rightarrow x_0} g(x) = 0 \quad (86)$$

oder

$$\lim_{x \rightarrow x_0} g(x) \in \{-\infty, +\infty\}. \quad (87)$$

Weiter sei $g'(x) \neq 0$ für alle $x \in I$, und es existiere

$$\lim_{x \rightarrow x_0} \frac{f'(x)}{g'(x)} =: A \in \mathbb{R} \cup \{-\infty, +\infty\}. \quad (88)$$

Dann existiert auch

$$\lim_{x \rightarrow x_0} \frac{f(x)}{g(x)} \quad (89)$$

und hat denselben Wert A .

Beweis. Wir beweisen das nur für die Voraussetzung (86). Der Einfachheit halber nehmen wir an, dass $x_0 \in I$. (In dem Satz könnte es auch einer der Endpunkte des Intervalls sein,

der nicht selbst zum Intervall gehört, z.B. $b = +\infty$.) Für $x \in I \setminus \{x_0\}$ gibt es dann nach dem verallgemeinerten Mittelwertsatz 231 ein ξ zwischen x und x_0 , so dass

$$\frac{f(x)}{g(x)} = \frac{f(x) - f(x_0)}{g(x) - g(x_0)} = \frac{f'(\xi)}{g'(\xi)} \rightarrow \frac{f'(x_0)}{g'(x_0)} \text{ für } x \rightarrow x_0.$$

□

Beispiel 233. Sei $f(x) = \cos x$, $g(x) = 2x - \pi$, $x_0 = \pi/2$. Die Voraussetzungen des Satzes sind erfüllt, und weil

$$\lim_{x \rightarrow \frac{\pi}{2}} \frac{f'(x)}{g'(x)} = \lim_{x \rightarrow \frac{\pi}{2}} \frac{-\sin x}{2} = \frac{-\sin \pi/2}{2} = -\frac{1}{2},$$

ist auch

$$\lim_{x \rightarrow \frac{\pi}{2}} \frac{\cos x}{2x - \pi} = -\frac{1}{2}.$$

□

Beispiel 234. Sei $f(x) = \sinh x$, $g(x) = 2x + 3$. Dann gilt **nicht**

$$\lim_{x \rightarrow 0} \frac{\sinh x}{2x + 3} = \lim_{x \rightarrow 0} \frac{\cosh x}{2} = \frac{1}{2}.$$

Warum nicht?

□

Beispiel 235. Für $f(x) = x^k$, $g(x) = e^x$, $x_0 = +\infty$ sind die Voraussetzungen des Satzes erfüllt, aber was ist

$$\lim_{x \rightarrow \infty} \frac{kx^{k-1}}{e^x}?$$

Dafür gelten wiederum die Voraussetzungen des Satzes usw. Schließlich erhalten wir

$$\lim_{x \rightarrow \infty} \frac{x^k}{e^x} = \lim_{x \rightarrow \infty} \frac{kx^{k-1}}{e^x} = \dots = \lim_{x \rightarrow \infty} \frac{k!}{e^x},$$

falls der letzte Limes rechts existiert. Tut er aber und ist null. Also

$$\boxed{\lim_{x \rightarrow \infty} \frac{x^k}{e^x} = 0.}$$

Die Exponentialfunktion wächst schneller als jede Potenz.

□

Bemerkung. Eine solche iterierte Anwendung der Regel von Bernoulli/de l'Hospital kommt häufiger vor. Wenn

$$\lim_{x \rightarrow x_0} \frac{f^{(k)}(x)}{g^{(k)}(x)}$$

existiert und $g^{(k)}(x)$ keine Nullstelle hat, braucht man die entsprechende Bedingung für die niedrigeren Ableitungen nicht mehr zu verifizieren: sie haben dann höchstens endlich viele Nullstellen, und dafür gilt der Satz auch.

Beispiel 236 (Langevin-Funktion).*Werkstoffe I, Abschnitt 2.7*

Vergleichen Sie Beispiel 143.

Die Langevin-Funktion

$$L(x) = \frac{\cosh x}{\sinh x} - \frac{1}{x},$$

spielt im Zusammenhang mit paramagnetischem Verhalten eine Rolle. Für sie gilt

$$\begin{aligned} \lim_{x \searrow 0} L(x) &= \lim_{x \searrow 0} \left(\frac{\cosh x}{\sinh x} - \frac{1}{x} \right) = \lim_{x \searrow 0} \frac{x \cosh x - \sinh x}{x \sinh x} \\ &= \lim_{x \searrow 0} \frac{\cosh x + x \sinh x - \cosh x}{\sinh x + x \cosh x} = \lim_{x \searrow 0} \frac{x \sinh x}{\sinh x + x \cosh x} \\ &= \lim_{x \searrow 0} \frac{\sinh x + x \cosh x}{\cosh x + \cosh x + x \sinh x} = 0. \end{aligned}$$

Beispiel 237. Es gilt

$$\lim_{x \rightarrow +\infty} x \ln\left(1 + \frac{1}{x}\right) = \lim_{x \rightarrow +\infty} \frac{\ln\left(1 + \frac{1}{x}\right)}{\frac{1}{x}} = \lim_{y \searrow 0} \frac{\ln(1+y)}{y} = \lim_{y \searrow 0} \frac{\frac{1}{1+y}}{1} = 1.$$

Dabei folgt das vorletzte Gleichheitszeichen aus der Regel von Bernoulli/de l'Hospital. Weil $\exp$ stetig ist, ergibt sich daraus

$$\lim_{x \rightarrow +\infty} \exp\left(x \ln\left(1 + \frac{1}{x}\right)\right) = \exp(1),$$

also

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

□

11.6 Zum Integralbegriff

Dieser Abschnitt ist nur für Mathematiker bestimmt.

Die Integraldefinition im Satz 148 ist technisch sehr einfach, sie erfordert keine Ober- und Untersummen wie das Riemannintegral und keine gleichmäßige Approximation wie das Regelintegral. Deshalb muss man auch keinen Unabhängigkeitsnachweis führen. Der Existenzbeweis für den Grenzwert der Riemannschen Summen bei stetigen Funktionen ist Standard, für monotone Funktionen verdanke ich ihn Ekki Tjaden:

Sei $f : [a, b] \rightarrow \mathbb{R}$ monoton wachsend. Sei

$$T : a = x_0 < x_1 < \dots < x_n = b$$

eine Zerlegung und

$$\delta_T := \max_j (x_j - x_{j-1}).$$

Definiere

$$\underline{S}_T = \sum_{j=1}^n f(x_{j-1})(x_j - x_{j-1}), \quad \overline{S}_T = \sum_{j=1}^n f(x_j)(x_j - x_{j-1}).$$

Dann ist

$$\overline{S}_T - \underline{S}_T = \sum_{j=1}^n (f(x_j) - f(x_{j-1}))(x_j - x_{j-1}) \leq (f(b) - f(a))\delta_T.$$

Weiter ist jede Untersumme kleiner als jede Obersumme:

$$\underline{S}_{T_1} \leq \overline{S}_{T_2}$$

für beliebige Zerlegungen T_1 und T_2 . Es folgt

$$\begin{aligned} \underline{S}_{T_1} &\leq \overline{S}_{T_2} \leq \underline{S}_{T_2} + (f(b) - f(a))\delta_{T_2} \\ \underline{S}_{T_2} &\leq \overline{S}_{T_1} \leq \underline{S}_{T_1} + (f(b) - f(a))\delta_{T_1} \end{aligned}$$

und

$$|\underline{S}_{T_1} - \underline{S}_{T_2}| \leq (f(b) - f(a)) \max(\delta_{T_1}, \delta_{T_2}).$$

Wählt man daher eine Zerlegungsfolge, für die $\delta_{T_j} \rightarrow 0$, so ist $(\underline{S}_{T_j})_{j \in \mathbb{N}}$ eine Cauchyfolge und daher konvergent. Weiter konvergiert die Folge der Obersummen $(\overline{S}_{T_j})_{j \in \mathbb{N}}$ gegen denselben Grenzwert. Also ist der Grenzwert das Riemannintegral $\int_a^b f(x)dx$. Die „linksanschlagenden“ Untersummen $\underline{S}_{T_j}$ konvergieren also für jede, insbesondere für die äquidistante, Zerlegungsfolge mit $\delta \rightarrow 0$, gegen das Riemannintegral.

Für eine mathematisch befriedigende Theorie müßte man das Integral auf den von den stückweise stetigen und monotonen Funktionen erzeugten Vektorraum linear fortsetzen – und landet beim Regelintegral. Ich verzichte aus offensichtlichen Gründen darauf. Der einzige Nachteil, den man sich einhandelt, sind etwas stärkere Voraussetzungen für die Rechenregeln.

11.7 Von der Summe zum Integral

Wir erläutern den Übergang von der Summe zum Integral noch einmal an einem Beispiel aus der Mechanik.

Beispiel 238 (Drehmoment). Ein an einem Ende eingespannter horizontaler Balken übt auf die Befestigung ein Drehmoment aus. Wie groß ist dieses?

Wir messen die horizontale Entfernung von der Befestigung mit der Variablen x mit $0 \leq x \leq L$. Wir teilen den Balken durch nah beieinanderliegende Punkte $x_i, i \leq n$. Dann ist das Drehmoment des Stücks von x_{i-1} bis $x_i = x_{i-1} + \Delta x_i$ annähernd gegeben durch

$$N_i = g\rho q(x_i)x_i\Delta x_i,$$

wobei g die Erdbeschleunigung, ρ die Dichte und $q(x_i)$ den Querschnitt des Balkens an der Stelle x_i bezeichnet. (Es ist nämlich $q(x_i)\Delta x_i$ das Volumen, $g\rho q(x_i)\Delta x_i$ die Gewichtskraft und x_i der Hebelarm – alle Werte nur angenähert.)

Das gesamte Drehmoment ist dann annähernd

$$N \approx \sum_{i=1}^n g\rho q(x_i)x_i\Delta x_i.$$

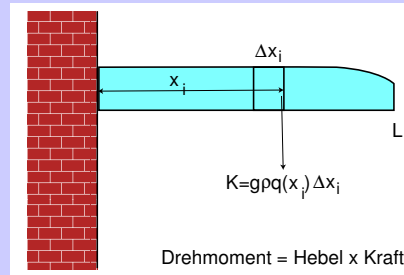
Den genauen Wert für das Drehmoment erhält man durch immer feinere Unterteilung, das heißt durch Übergang zum Integral:

$$N = \int_0^L g\rho q(x)x \, dx.$$

Ist der Querschnitt q konstant, so ergibt sich

$$N = g\rho q \int_0^L x \, dx = \frac{1}{2}g\rho qL^2 = \frac{g}{2}ML,$$

wobei $M = \rho qL$ die Masse des Balkens ist.



11.8 Numerische Integration.

Im Prinzip ist die Approximation durch Summen gemäß unserer Integraldefinition ein brauchbares numerisches Verfahren zur Berechnung von Integralen, das im Grenzwert ja sogar gegen den genauen Wert konvergiert. Aber dafür muss man natürlich (unendlich) viele Zerlegungen betrachten. Für die schnelle näherungsweise Berechnung bestimmter Integrale verwendet man die Approximation durch andere, „genauere“ Funktionen, die sich aber auch noch gut beherrschen lassen:

Trapezregel. Eine erste Näherung für $\int_a^b f(x)dx$ ist

$$\int_a^b f(x)dx \approx \frac{f(a) + f(b)}{2} (b - a).$$

Der Ausdruck rechts ist gerade der Flächeninhalt des Trapezes, welches entsteht, wenn wir den Graphen der Funktion durch die Sekante zwischen den Endpunkten $(a, f(a))$ und $(b, f(b))$ ersetzen. Diese Regel kann man auch mehrfach anwenden, indem man $[a, b]$ in n gleiche Teilintervalle der Länge $h = \frac{b-a}{n}$ unterteilt, für jedes Teilintervall die Trapezregel anwendet und summiert. Man erhält

$$\int_a^b f(x)dx \approx \frac{b-a}{2n} \left(f(a) + f(b) + 2 \sum_{i=1}^{n-1} f(a + ih) \right).$$

Simpsonregel. Statt den Graphen von $f : [a, b] \rightarrow \mathbb{R}$ durch eine Gerade, nämlich die Sekante zu ersetzen, kann man ihn auch ersetzen durch einen Parabelbogen durch die Graphenpunkte zu a , $\frac{a+b}{2}$ und b . Die entstehende Fläche, die die gesuchte vermutlich besser approximiert als das Trapez, lässt sich explizit berechnen. Es folgt

$$\int_a^b f(x)dx \approx \frac{b-a}{6} \left(f(a) + f(b) + 4f\left(\frac{a+b}{2}\right) \right).$$

Auch diese Regel kann man mehrfach anwenden. Man unterteilt $[a, b]$ in eine gerade Anzahl $2n$ von Teilintervallen gleicher Länge $h = \frac{b-a}{2n}$ und wendet auf jedes ungerade Teilintervall und das folgende gerade die Simpsonregel an. Man erhält

$$\int_a^b f(x)dx \approx \frac{h}{3} \left(f(a) + f(b) + 4 \sum_{i=1}^n f(a + (2i-1)h) + 2 \sum_{i=1}^{n-1} f(a + 2ih) \right).$$

Für subtilere Verfahren der numerischen Integration schauen Sie in bessere Formelsammlungen.

Mathematische Software verfügt über Programme zur symbolischen wie zur numerischen Integration. Die letzteren beruhen auf solchen Verfahren.

11.9 Was sind Orthogonalitätsrelationen?

Im Satz 183 hatten wir gezeigt, dass

$$\begin{aligned}\frac{2}{T} \int_0^T \cos k\omega t \cos m\omega t dt &= \begin{cases} 2, & \text{falls } m = k = 0 \\ 1, & \text{falls } m = k > 0 \\ 0, & \text{falls } m \neq k \end{cases} \\ \frac{2}{T} \int_0^T \sin k\omega t \sin m\omega t dt &= \begin{cases} 1, & \text{falls } m = k > 0 \\ 0, & \text{falls } m \neq k \end{cases} \\ \frac{2}{T} \int_0^{2\pi} \cos k\omega t \sin m\omega t dt &= 0 \end{aligned}$$

für alle $k, m \in \mathbb{N}$ und $\omega = \frac{2\pi}{T} > 0$, und wir hatten das *Orthogonalitätsrelationen* genannt. „Orthogonal“ heißt ja „senkrecht“, woher kommt also der Name?

Im Euklidischen Raum hat man ein Skalarprodukt

$$\langle \vec{x}, \vec{y} \rangle = |\vec{x}| |\vec{y}| \cos \alpha,$$

wobei α den Winkel zwischen den beiden Vektoren bezeichnet. Die beiden sind senkrecht zueinander genau dann, wenn der Winkel ein rechter ist, d.h. wenn $\langle \vec{x}, \vec{y} \rangle = 0$.

Für Funktionen kann man auch ein Skalarprodukt erklären, das aus zwei Funktionen einen reellen Zahl macht und dieselben Rechenregeln wie das Skalarprodukt im $\mathbb{R}^n$ erfüllt:

$$\langle f, g \rangle := \frac{2}{T} \int_0^T f(t)g(t)dt.$$

Dabei betrachten wir Funktionen auf dem Intervall $[0, T]$. Die Orthogonalitätsbedingungen besagen dann gerade, dass die Funktionen $\cos k\omega t$ und $\sin k\omega t$ paarweise orthogonal zueinander sind.

Die Analogie geht aber weiter. In der Vektorrechnung sind *Orthonormalbasen* sehr bequeme Basen. Bezüglich einer beliebigen Basis $\vec{e}_1, \dots, \vec{e}_n$ lässt sich jeder Vektor eindeutig darstellen als

$$\vec{x} = \sum_{k=1}^n x_k \vec{e}_k.$$

Ist die Basis orthonormal, d.h. sind die Basisvektoren senkrecht zueinander und von der Länge 1, gilt also

$$\langle \vec{e}_j, \vec{e}_k \rangle = \begin{cases} 1 & \text{für } j = k, \\ 0 & \text{für } j \neq k, \end{cases}$$

so sind die Koeffizienten gegeben durch

$$x_j = \langle \vec{x}, \vec{e}_j \rangle$$

und man hat

$$\vec{x} = \sum_{k=1}^n \langle \vec{x}, \vec{e}_k \rangle \vec{e}_k.$$

Der Beweis ist einfach:

$$\langle \vec{x}, \vec{e}_j \rangle = \left\langle \sum_{k=1}^n x_k \vec{e}_k, \vec{e}_j \right\rangle = \sum_{k=1}^n x_k \langle \vec{e}_k, \vec{e}_j \rangle.$$

Die Summanden sind aber alle 0, außer dem mit $k = j$, und der liefert $x_j \cdot 1$.

Sind $\vec{e}_1, \dots, \vec{e}_n$ zwar orthonormal, aber keine Basis, zum Beispiel zwei zueinander senkrechte Einheitsvektoren im 3-dimensionalen Raum, so ist $\sum_{k=1}^n x_k \vec{e}_k$ die Orthogonalprojektion von $\vec{x}$ auf den vom System $\vec{e}_1, \dots, \vec{e}_n$ aufgespannten Unterraum, also die beste Approximation von $\vec{x}$, die man mit diesen Vektoren finden kann.

Die Sin-Cos-Funktionen sind nun auch ein Orthonormalsystem, wenn man von dem kleinen Schönheitsfehler mit $\cos 0\omega T$ absieht²⁶. Die Rechnung im Beispiel 184 ist deshalb genau dieselbe wie im vorstehenden Beweis. In der Fourieranalyse stellt man Fourierpolynome bezüglich einer Orthonormalbasis dar und approximiert T -periodische Funktionen bezüglich des trigonometrischen Orthonomalsystems.

²⁶Diese Funktion hat nicht „Länge“ =1, sondern $=\sqrt{2}$