

Skript zur Vorlesung

Höhere Mathematik II

für die Fachrichtungen

MACH, MIT, MWT, BIW, CIW, VT, GEOD

25. Auflage

PD Dr. Tilo Arens

PD Dr. Frank Hettlich

Prof. Dr. Andreas Kirsch

Institut für Angewandte und Numerische Mathematik
Karlsruher Institut für Technologie (KIT)

Literatur für diese Vorlesung (u.a.)

- T. Arens, F. Hettlich, C. Karpfinger, U. Kockelkorn, K. Lichtenegger, H. Stachel: Mathematik. Spektrum Akademischer Verlag, Heidelberg, 5. Auflage, 2022.
- K. Burg, H. Haf, F. Wille: Höhere Mathematik für Ingenieure. Band I-III. Teubner Verlag, Stuttgart.
- G. Merziger, T. Wirth: Repetitorium der höheren Mathematik. Binomi Verlag.

Formelsammlungen:

- T. Arens, F. Hettlich, Ch. Karpfinger, U. Kockelkorn, K. Lichtenegger, H. Stachel: Mathematik zum Mitnehmen. Springer Spektrum, Heidelberg.
- G. Merziger, G. Mühlbach, D. Wille, T. Wirth: Formeln und Hilfen zur Höheren Mathematik. Binomi Verlag.

Gedruckt auf 100% Recyclingpapier mit dem Gütesiegel „Blauer Engel“

Inhaltsverzeichnis

1	Vektorräume	3
1.1	Einführung	3
1.2	Das Gauß'sche Eliminationsverfahren	8
1.3	Linearkombinationen und lineare Unabhängigkeit	11
1.4	Euklidische Vektorräume	17
2	Lineare Abbildungen	25
2.1	Lineare Abbildungen und Matrizen	25
2.2	Theorie der linearen Gleichungssysteme	32
2.3	Determinanten	37
2.4	Eigenwerte und Eigenvektoren	44
3	Fouriertheorie	47
3.1	Trigonometrische Polynome	47
3.2	Fourierreihen	54
4	Differentialgleichungen	62
4.1	Homogene lineare Differentialgleichungen mit konstanten Koeffizienten	63
4.2	Der Satz von Picard-Lindelöf	68
4.3	Lineare Differentialgleichungssysteme	70
4.4	Bestimmung partikulärer Lösungen	78
4.5	Die Potenzreihenmethode	82
4.6	Numerische Lösungsverfahren	87
5	Die Laplacetransformation	92
5.1	Parameterintegrale	92
5.2	Definition und Grundlegende Eigenschaften der Laplace-Transformation	94
5.3	Anwendung bei Anfangswertaufgaben	98
5.4	Faltungsprodukt und Faltungssatz	105
5.5	Die Delta-Distribution	108

1 Vektorräume

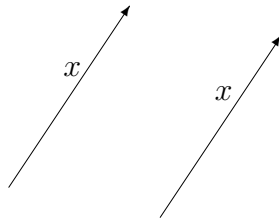
1.1 Einführung

Wir wollen die Grundbegriffe der Vektorrechnung wiederholen bzw. einführen. Dazu betrachten wir zunächst einige Beispiele.

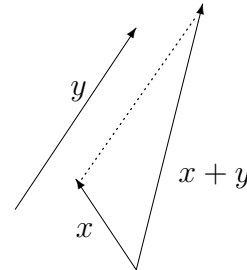
Beispiele:

- (a) In der Mechanik werden physikalische Größen wie Kräfte, Geschwindigkeiten oder Impulse als Pfeile veranschaulicht, die durch ihre **Richtung** und ihre **Länge** festgelegt sind. Man spricht von **Vektoren**. Da ein Vektor beliebig in der Ebene oder im Raum verschoben werden kann, spricht man auch von freien Vektoren.

Man kann zwei freie Vektoren **addieren**, indem man den Anfangspunkt des zweiten in den Endpunkt des ersten legt. Die Summe ist dann der Pfeil vom Anfangspunkt des ersten zum Endpunkt des zweiten. Man kann ferner einen freien Vektor **mit einer reellen Zahl multiplizieren**: Ist die Zahl positiv, ist das Ergebnis der um diese Zahl gestreckte Vektor. Ist die Zahl negativ, wird zusätzlich die Richtung umgekehrt.



Freier Vektor



Addition freier Vektoren

- (b) In der Geometrie kann mit einem Punkt in der Ebene oder im Raum dessen Ortsvektor assoziiert werden, der Verbindungsvektor des Ursprungs O mit diesem Punkt. Zwei Punkte A , B bzw. deren Ortsvektoren können addiert werden, indem man das Parallelogramm konstruiert, das die Eckpunkte A , O und B besitzt. Der vierte, O gegenüberliegende Punkt, ist dann die Summe von A und B . Auch die Multiplikation von A mit einer reellen Zahl λ ist möglich: Ist $\lambda > 0$, so ist das Ergebnis der Punkt auf dem durch O und A festgelegten Strahl mit dem λ -fachen des Abstands von O verglichen mit A . Ist $\lambda < 0$, so muss zusätzlich an O gespiegelt werden.
- (c) Zwei Polynome $p, q : \mathbb{C} \rightarrow \mathbb{C}$ vom Grad kleiner gleich N können im Sinne von Funktionen addiert und mit einer komplexen Zahl multipliziert werden. Ist

$$p(x) = \sum_{j=0}^N a_j x^j, \quad q(x) = \sum_{j=0}^N b_j x^j, \quad x \in \mathbb{C},$$

und $\lambda \in \mathbb{C}$, so ist

$$(p+q)(x) = \sum_{j=0}^N (a_j + b_j) x^j, \quad (\lambda p)(x) = \sum_{j=0}^N \lambda a_j x^j, \quad x \in \mathbb{C}.$$

Beide Funktionen $p+q$ bzw. λp sind wieder Polynome vom Grad kleiner gleich N .

In allen drei Fällen sind also auf einer Menge von Objekten zwei Operationen **Addition** und **Multiplikation mit einer Zahl** definiert. Durch Anwendung der Operationen erhält man jeweils wieder ein Objekt dieser Menge. Bevor wir die Rechenregeln identifizieren, die für diese Operationen gelten, wollen wir für uns wichtige Beispiele solcher Strukturen kennenlernen.

Definition 1.1 Sei $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$ und $n \in \mathbb{N}$ eine natürliche Zahl. (Wir denken vor allem an $n = 2$ oder $n = 3$ sowie $\mathbb{K} = \mathbb{R}$.)

Spaltenvektoren bestehen aus n übereinandergeschriebenen Zahlen,

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \quad \text{mit } x_1, \dots, x_n \in \mathbb{K}.$$

Die Zahlen $x_1, \dots, x_n$ heißen **Komponenten** oder **Koordinaten** des Vektors x . Die Menge aller Vektoren mit n Komponenten wird mit $\mathbb{K}^n$ bezeichnet (also $\mathbb{R}^n$ bzw. $\mathbb{C}^n$ für $\mathbb{K} = \mathbb{R}$ bzw. $\mathbb{K} = \mathbb{C}$).

Für $x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \in \mathbb{K}^n$ und $\lambda \in \mathbb{K}$ werden **Addition** und **skalare Multiplikation** definiert durch:

$$x + y = \begin{pmatrix} x_1 + y_1 \\ \vdots \\ x_n + y_n \end{pmatrix}, \quad \lambda x = \begin{pmatrix} \lambda x_1 \\ \vdots \\ \lambda x_n \end{pmatrix}.$$

Um Platz zu sparen, schreiben wir auch

$$x = (x_1, \dots, x_n)^\top \in \mathbb{K}^n \quad \text{statt} \quad x = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Bei der letzten Notation steht das hochgestellte $\top$ für *transponiert*. Im Zusammenhang mit der Matrizenrechnung wird dies später eine wichtige Operation.

Wir können nun nachrechnen, dass für $V = \mathbb{K}^n$ die folgenden Rechengesetze gelten:

- (V1) Mit $x, y \in V$ und $\lambda \in \mathbb{K}$ sind $x + y$ und λx wieder Elemente von V (Abgeschlossenheit bezüglich der Operationen)
- (V2) Für alle $x, y, z \in V$ gilt $x + (y + z) = (x + y) + z$ (Assoziativgesetz).
- (V3) Für alle $x, y \in V$ gilt $x + y = y + x$ (Kommutativgesetz).
- (V4) Es existiert genau ein Element $o \in V$ mit $x + o = x$ für alle $x \in V$. Das Element o heißt **neutrales Element** oder **Nullelement**.
- (V5) Zu jedem $x \in V$ existiert genau ein $x' \in V$ mit $x + x' = o$. Das Element x' wird mit $-x$ bezeichnet und heißt **Negatives** zu x . Es gilt also $x + (-x) = o$ für alle x .
- (V6) Für alle $x \in V$ und $\lambda, \mu \in \mathbb{K}$ gilt $(\lambda + \mu)x = \lambda x + \mu x$.
- (V7) Für alle $x, y \in V$ und $\lambda \in \mathbb{K}$ gilt $\lambda(x + y) = \lambda x + \lambda y$.

(V8) Für alle $x \in V$ und $\lambda, \mu \in \mathbb{K}$ gilt $(\lambda\mu)x = \lambda(\mu x)$

(V9) Für alle $x \in V$ gilt $1x = x$ (1 ist die Zahl Eins).

Dabei ist das Nullelement derjenige Spaltenvektor, dessen sämtliche Komponenten null sind. Das Negative von $x \in \mathbb{K}^n$ erhält man, indem man bei allen Komponenten das Vorzeichen ändert.

Sie können (und sollten!) sich davon überzeugen, dass auch bei den anderen einführenden Beispielen die Gesetze (V1)–(V9) erfüllt sind. Sehr unterschiedliche Mengen zeigen also ein und dieselbe Struktur. Hierfür führen wir einen eigenen Begriff ein.

Definition 1.2 Eine Menge V , auf der zusammen mit $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$ zwei algebraische Operationen **Addition** und **skalare Multiplikation** definiert sind, die den Gesetzen (V1)–(V9) genügen, heißt **Vektorraum** über $\mathbb{K}$ oder $\mathbb{K}$ -Vektorraum. Die Elemente von V heißen **Vektoren**.

Wenn wir einen Sachverhalt über einen Vektorraum nur unter Verwendung der Rechengesetze (V1)–(V9) nachweisen, so gilt dieser Sachverhalt automatisch für alle Vektorräume.

Die wesentliche Aussage, dass die Spaltenvektoren einen Vektorraum bilden, formulieren wir noch einmal als einen Satz.

Satz 1.3 Der $\mathbb{K}^n$ ist ein **Vektorraum** über $\mathbb{K}$, d.h. die Eigenschaften (V1)–(V9) gelten mit dem neutralen Element $0 = (0, \dots, 0)^\top$ und dem Negativen $-x = (-1)x$ zu $x \in \mathbb{K}^n$.

Bemerkung: Der Raum $\mathbb{C}^1 = \mathbb{C}$ ist einerseits ein $\mathbb{C}$ -Vektorraum. Andererseits können wir uns auch auf die Multiplikation nur mit reellen Zahlen einschränken, also $\mathbb{C}$ als $\mathbb{R}$ -Vektorraum betrachten. Dann entspricht $\mathbb{C}$ dem $\mathbb{R}^2$.

Für $n = 1$ wird $\mathbb{K}^1$ mit $\mathbb{K}$ identifiziert. In der Literatur wird häufig Fettdruck, oder eine spezielle Notation wie $\underline{x}$ oder $\vec{x}$ verwendet, um deutlich zu machen, dass es sich um einen Vektor handelt. Wir werden aber im Folgenden auf eine gesonderte Kennzeichnung verzichten; denn, ob mit oder ohne Kennzeichnung, es wird stets unerlässlich sein, sich im Zusammenhang klarzumachen, welche Elemente aus welcher Menge gerade durch einen bestimmten Buchstaben vertreten werden.

Veranschaulichung des $\mathbb{R}^2$ und $\mathbb{R}^3$: Für $n = 2$ oder $n = 3$ können wir $x = (x_1, x_2)^\top$ bzw. $x = (x_1, x_2, x_3)^\top$ als Punkte der Ebene bzw. des Raums auffassen, wenn wir ein festes Koordinatensystem zugrunde gelegt haben. Üblicherweise nehmen wir ein rechtwinkliges Koordinatensystem. Z.B. entspricht der Spaltenvektor $x = (1, 2)^\top$ dem Punkt X mit der x_1 -Koordinate 1 und der x_2 -Koordinate 2 der Ebene. Oft interpretieren wir $x = (x_1, x_2)^\top$ auch als Ortsvektor $\vec{OX}$ oder als freien Vektor. Ab jetzt identifizieren wir Vektoren mit Ortsvektoren und Punkten der Ebene, schreiben also z.B. $(3, 1)^\top$ und meinen damit den Spaltenvektor $(3, 1)^\top$, aber auch den entsprechenden Punkt oder Ortsvektor in einem Koordinatensystem.

Es gibt viele weitere Rechengesetze für Spaltenvektoren. Diese können direkt nachgewiesen werden, oder nur mit Hilfe der Axiome (V1)–(V9). Einige wichtige fassen wir im folgenden Satz zusammen:

Satz 1.4 (a) $0x = 0 \in \mathbb{K}^n$ für alle $x \in \mathbb{K}^n$ und $\lambda 0 = 0$ für alle $\lambda \in \mathbb{K}$.
(b) $\lambda x = 0 \in \mathbb{K}^n$ impliziert $\lambda = 0$ oder $x = 0$.
(c) $(-\lambda)x = -(\lambda x) = \lambda(-x)$ für alle $x \in \mathbb{K}^n$ und $\lambda \in \mathbb{K}$.

Beweis nur mit Hilfe der Axiome (V1)–(V9) (und damit gültig in allen Vektorräumen):

(a) Nach (V6) ist $0x + 0x = (0 + 0)x = 0x$. Addition des Negativen $-(0x)$ zu $0x$ liefert (benutzt wird auch (V2)) $0x + 0 = 0$, also wegen (V4) die Behauptung $0x = 0$. Genauso folgt: $\lambda 0 + \lambda 0 = \lambda(0 + 0) = \lambda 0$, also nach Subtraktion von $\lambda 0$ die Behauptung $\lambda 0 = 0$.

(b) Sei $\lambda x = 0$ und $\lambda \neq 0$. Dann ist $x = 1x = (\frac{1}{\lambda}\lambda)x = \frac{1}{\lambda}(\lambda x) = \frac{1}{\lambda}0 = 0$ wegen (a).

(c) $(-\lambda)x + \lambda x = (-\lambda + \lambda)x = 0x = 0$, also ist $(-\lambda)x = -(\lambda x)$. Speziell für $\lambda = 1$ ist $(-1)x = -(1x) = -x$ und $\lambda(-x) = \lambda(-1)x = (-\lambda)x$. $\square$

Geraden in der Ebene $\mathbb{R}^2$ oder im Raum $\mathbb{R}^3$ und Ebenen im $\mathbb{R}^3$ haben eine besondere Struktur, sie sind „nicht krumm“. Dies präzisieren wir:

Definition 1.5 (a) Sei $U \subseteq V$ eine nichtleere Teilmenge eines Vektorraums V über $\mathbb{K}$ mit den Eigenschaften:

Für alle $x, y \in U$ und $\lambda \in \mathbb{K}$ gilt $x + y \in U$ und $\lambda x \in U$.

(Man sagt auch dazu: U ist bzgl. Addition und skalarer Multiplikation **abgeschlossen**.) Dann heißt U ein **linearer Unterraum** oder **Teilraum** von V .

(b) Sei $z \in V$ ein fester Vektor und U ein linearer Unterraum. Dann heißt die Menge

$$z + U := \{z + x : x \in U\}$$

ein **affiner Unterraum**. Interessant ist dieser Begriff natürlich nur, wenn $z \notin U$ ist.

Wir erkennen sofort, wenn wir in obiger Definition $\lambda = 0$ bzw. $\lambda = -1$ setzen, dass für jeden Unterraum U gilt: $0 \in U$ und mit jedem $x \in U$ ist auch das Negative $-x \in U$. Die Eigenschaften (V1)–(V9) gelten auch, wenn wir uns auf Elemente von U einschränken.

Lemma 1.6 Ein linearer Unterraum ist selbst ein Vektorraum.

Für jeden affinen Unterraum $z + U$ gilt $z \in z + U$.

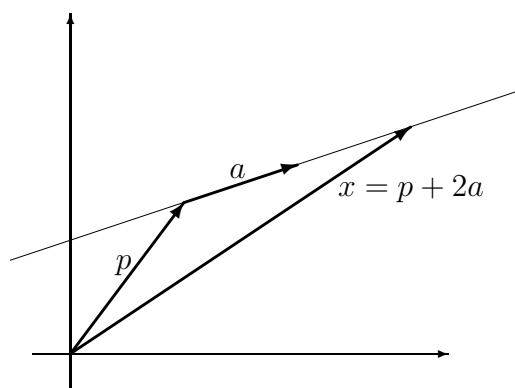
Triviale lineare Unterräume von $\mathbb{K}^n$ sind $U = \mathbb{K}^n$ und $U = \{0\}$. Dies sind die größt- bzw. kleinstmöglichen linearen Unterräume. Als Beispiele von affinen Unterräumen betrachten wir **Geraden** und **Ebenen** in Parameterform.

Definition 1.7 (Parameterform von Geraden im $\mathbb{R}^n$)

Sei $n \geq 2$, sowie $p \in \mathbb{R}^n$ und $a \in \mathbb{R}^n \setminus \{0\}$ gegeben. Die Gerade G durch den Punkt p in Richtung a ist gegeben durch

$$G = \{p + ta \in \mathbb{R}^n : t \in \mathbb{R}\}.$$

Da jedem Parameter $t \in \mathbb{R}$ der Punkt $x = p + ta \in G$ entspricht, spricht man auch von der **Parameterform** der Geraden.



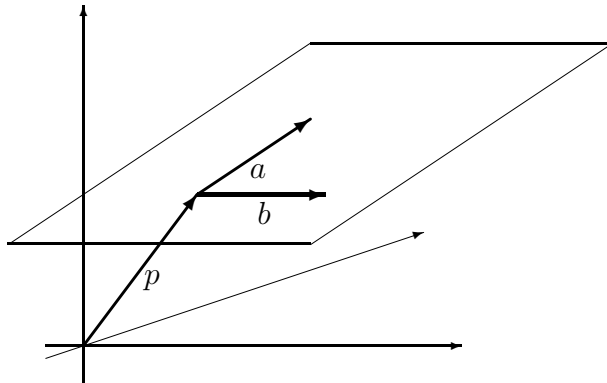
Die Menge $U = \{ta : t \in \mathbb{R}\}$ ist ein linearer Unterraum, denn $t_1a + t_2a = (t_1 + t_2)a$ und $t_1(t_2a) = (t_1t_2)a$ für alle $t_1, t_2 \in \mathbb{R}$. Daher ist $\{p + ta : t \in \mathbb{R}\} = p + U$ ein affiner Unterraum. Sind zwei Vektoren p und q des $\mathbb{R}^n$ gegeben, so gibt $q - p$ die Richtung der Geraden durch p und q an, also ist $x = p + t(q - p)$, $t \in \mathbb{R}$, die Parameterform der Geraden durch p und q . Für $t \in (0, 1)$ liegt $x = p + t(q - p)$ „zwischen“ p und q .

Definition 1.8 (*Parameterform von Ebenen im $\mathbb{R}^n$*)

Jetzt sei $n \geq 3$. Gegeben seien $p \in \mathbb{R}^n$, $a, b \in \mathbb{R}^n \setminus \{0\}$, für die $a \neq \lambda b$ für alle $\lambda \in \mathbb{R}$ gelte. Dann ist

$$E = \{p + ta + sb \in \mathbb{R}^n : s, t \in \mathbb{R}\},$$

die Ebene E durch den Punkt p , die mit den Richtungsvektoren a und b aufgespannt wird.



Die Menge $U = \{ta + sb : s, t \in \mathbb{R}\}$ ist ein linearer Unterraum, daher $\{p + ta + sb : s, t \in \mathbb{R}\} = p + U$ ein affiner Unterraum.

Beispiel: Bestimme die Ebene durch die Punkte $(1, 1, 1)^\top$, $(1, 2, 3)^\top$ und $(2, 4, 1)^\top$. Hier kann man nehmen:

$$p = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, a = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} - \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix}, b = \begin{pmatrix} 2 \\ 4 \\ 1 \end{pmatrix} - \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 3 \\ 0 \end{pmatrix}, \text{ also } x = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + t \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix} + s \begin{pmatrix} 1 \\ 3 \\ 0 \end{pmatrix}, \quad s, t \in \mathbb{R}, \text{ oder in Komponenten:}$$

$$\left. \begin{aligned} x_1 &= 1 + s \\ x_2 &= 1 + t + 3s \\ x_3 &= 1 + 2t \end{aligned} \right\} \quad t, s \in \mathbb{R}.$$

Die Ebene durch drei Punkte p , q und r ist also gegeben durch

$$\{x \in \mathbb{R}^3 : x = p + t(q - p) + s(r - p), \quad s, t \in \mathbb{R}\}.$$

Es muss vorausgesetzt werden, dass p , q und r nicht auf einer Geraden liegen, später führen wir hierfür die Sprechweise ein, dass die Vektoren $q - p$ und $r - p$ *linear unabhängig* sind.

Hält man in der Ebenengleichung $x = p + ta + sb$, $s, t \in \mathbb{R}$, einen der beiden Parameter (z.B. t) fest, so beschreibt der andere Parameter (hier s) eine Gerade in der Ebene (hier in Richtung b).

Wir erkennen, dass im Spezialfall $p = 0$ die Gerade $x = ta$, $t \in \mathbb{R}$, bzw. die Ebene $x = ta + sb$, $t, s \in \mathbb{R}$, durch den Nullpunkt gehen und damit lineare Unterräume sind. Für $p \neq 0$ ergeben sich entsprechend verschobene affine Unterräume. Wir betrachten zwei weitere Beispiele, die ebenfalls Geraden im $\mathbb{R}^3$ beschreiben.

Beispiel 1.9 Seien $U, V \subseteq \mathbb{R}^3$ definiert durch

$$U := \{x \in \mathbb{R}^3 : 2x_1 + x_2 - 4x_3 = 0, \quad x_1 + 5x_3 = 0\},$$

$$V := \{x \in \mathbb{R}^3 : 2x_1 + x_2 - 4x_3 = 1, \quad x_1 + 5x_3 = 2\}.$$

Dann ist U ein linearer Unterraum von $\mathbb{R}^3$ und V ein affiner Unterraum (weshalb?). Wie sehen diese Mengen anschaulich aus? Wir werden später zeigen, dass sie jeweils aus den Schnittpunkten von zwei Ebenen bestehen, also selbst Geraden beschreiben.

1.2 Das Gauß'sche Eliminationsverfahren

Im obigen Beispiel 1.9 wurden die Mengen beschrieben durch Gleichungen, die die Vektoren erfüllen. Genauer handelt es sich um *lineare Gleichungen*, d.h. Gleichungen der Form

$$\sum_{j=1}^n a_j x_j = b$$

mit gegebenen Zahlen (Koeffizienten) a_j , $j = 1, \dots, n$, b und den zu bestimmenden unbekannten Werten x_j , $j = 1, \dots, n$. Häufig müssen, wie im Beispiel, simultan einige solcher Gleichungen betrachtet werden. Dies führt auf **lineare Gleichungssysteme**:

Es ist ein Vektor $x = (x_1, \dots, x_n)^\top \in \mathbb{K}^n$ von reellen ($\mathbb{K} = \mathbb{R}$) oder komplexen ($\mathbb{K} = \mathbb{C}$) Zahlen gesucht, sodass m Gleichungen der Form

$$\begin{array}{ccccccccc} a_{11} x_1 & + & a_{12} x_2 & + & \cdots & + & a_{1n} x_n & = & b_1 \\ a_{21} x_1 & + & a_{22} x_2 & + & \cdots & + & a_{2n} x_n & = & b_2 \\ \vdots & & \vdots & & \vdots & & \vdots & & \vdots \\ a_{m1} x_1 & + & a_{m2} x_2 & + & \cdots & + & a_{mn} x_n & = & b_m. \end{array} \quad (1.1)$$

erfüllt sind mit gegebenen Koeffizienten $a_{ij} \in \mathbb{K}$, $i = 1, \dots, m$, $j = 1, \dots, n$ und rechten Seiten $b_i \in \mathbb{K}$, $i = 1, \dots, m$, die wir entsprechend der Gleichungen indizieren.

Wir haben etwa im Beispiel

$$\begin{array}{ccccccc} -x_1 & + & 2x_2 & - & 7x_3 & = & -3 \\ 3x_1 & + & 5x_2 & - & 6x_3 & = & -1 \end{array}$$

zwei Gleichungen in den drei Variablen $(x_1, x_2, x_3)^\top \in \mathbb{R}^3$ mit der rechten Seite $(b_1, b_2)^\top = (-3, -1)^\top \in \mathbb{R}^2$ und den Koeffizienten $a_{11} = -1, a_{12} = 2, \dots, a_{23} = -6$.

Falls $b_1 = \dots = b_m = 0$ sind, heißt das Gleichungssystem **homogen**. Ansonsten ist es ein **inhomogenes** Gleichungssystem.

Wegen des wichtigen Konzepts der *Linearisierung*, das wir bei einer Variablen bereits beim Ableiten kennengelernt haben und auch noch bei mehreren Variablen untersuchen werden, sind bei vielen Problemen solche Gleichungssysteme teilweise mit sehr großer Anzahl an Variablen zu lösen. Deswegen betrachten wir hier einen systematischen Lösungsweg, den **Gauß-Algorithmus**. Dieser Zugang wird auch in Programmpaketen verwendet, zumindest wenn die Anzahl der Variablen nicht zu groß ist. Wegen der Rechenzeit sind bei sehr großen Systemen andere Verfahren erforderlich, die Lösungen iterativ annähern.

Wir lösen Gleichungssysteme, indem wir Schritt für Schritt einfachere Systeme mit gleicher Lösungsmenge bestimmen. Ziel ist es *Variablen zu eliminieren*, d.h. Gleichungen so zu kombinieren, dass die Variable sich in einfacher Weise direkt aus den übrigen ergibt. Bei linearen Gleichungssystemen ist dies stets möglich.

Beispiel 1.10

$$\begin{array}{ccccccc} x_1 & + & x_2 & + & x_3 & - & x_4 & = & 1 \\ 4x_1 & + & 5x_2 & + & 5x_3 & + & 2x_4 & = & 0 \\ x_1 & - & 2x_2 & + & 2x_3 & - & 5x_4 & = & 5 \end{array}$$

Subtrahiere $4 \times$ (1. Zeile) von der 2. Zeile und die 1. Zeile von der 3. Zeile, dann erhalten wir:

$$\begin{array}{rrrrrr} x_1 & + & x_2 & + & x_3 & - & x_4 & = & 1 \\ & & x_2 & + & x_3 & + & 6x_4 & = & -4 \\ & - & 3x_2 & + & x_3 & - & 4x_4 & = & 4 \end{array}$$

Addiere $3 \times$ 2. zur 3. Zeile:

$$\begin{array}{rrrrrr} x_1 & + & x_2 & + & x_3 & - & x_4 & = & 1 \\ & & x_2 & + & x_3 & + & 6x_4 & = & -4 \\ & & & & 4x_3 & + & 14x_4 & = & -8 \end{array}$$

Genau diese Dreiecksgestalt des Gleichungssystems ist das Ziel des Vorgehens. Nun können wir durch Einsetzen von unten nach oben (**Rücksubstitution**) x_3, x_2, x_1 in Abhängigkeit von x_4 angeben. Wir erhalten $x_3 = -2 - \frac{7}{2}x_4$, $x_2 = -2 - \frac{5}{2}x_4$ und $x_1 = 5 + 7x_4$. Für jede Wahl von $x_4 \in \mathbb{R}$ ergibt sich eine reelle Lösung des Gleichungssystems.

Benutzt wurden die folgenden **elementaren Umformungen**:

- (a) eine Gleichung wird ersetzt durch ein von Null verschiedenes Vielfaches dieser Gleichung.
- (b) eine Gleichung wird ersetzt durch die Addition der Gleichung und einem Vielfachen einer anderen Gleichung.

Die Umformungen unter (a) und (b) sind offensichtlich Äquivalenzumformungen, d.h. sie ändern die Lösungsmenge des Gleichungssystems nicht.

Da bei allen Rechnungen $x_1, \dots, x_n$ unangetastet bleiben, bietet sich eine übersichtlichere Notation an. Wir stellen nur die Koeffizienten in einem Tableau zusammen:

$$\begin{array}{cccc|c} a_{11} & a_{12} & \cdots & a_{1n} & b_1 \\ a_{21} & a_{22} & \cdots & a_{2n} & b_2 \\ \vdots & \vdots & & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} & b_m \end{array} \quad (1.2)$$

welches das Gleichungssystem repräsentiert. Das **Gauß'sche Eliminationsverfahren** besteht nun aus folgenden Schritten:

- (A) Wähle ein sogenanntes **Pivotelement**, d.h. ein Element $a_{ij} \neq 0$ des Schemas (nicht aus der b -Spalte). Für Handrechnungen nehme möglichst eines aus einer Spalte mit vielen Nullen, und/oder eine 1 oder -1 , um die Rechnungen nicht unnötig zu erschweren.
- (B) Führe elementare Umformungen durch, um in der **Spalte** des aktuellen Pivotelements alle Einträge **leerzuräumen**, die nicht in einer Zeile mit einem Pivotelement stehen.
- (C) Wähle wieder ein Pivotelement $\neq 0$ aus einer Zeile, in der noch kein Pivotelement aufgetreten ist, und fahre fort mit (B).
- (D) Breche ab, falls kein Pivotelement in einer „unmarkierten“ Zeile mehr vorhanden ist.

Die Anzahl r der Pivotelemente heißt **Rang** des Gleichungssystems. Um die Dreiecksgestalt zu bekommen, müssten jetzt die Zeilen eventuell noch entsprechend umsortiert werden. Mit ein wenig Routine können wir uns diesen Schritt sparen.

Es ist immer $\boxed{r \leq m \text{ und } r \leq n}$, da verschiedene Pivotelemente in verschiedenen Zeilen und verschiedenen Spalten stehen.

Beispiel 1.11 (a) (vergleiche Beispiel 1.10 oben) Wir berechnen

$$\begin{array}{c|ccc|c} \boxed{1} & 1 & 1 & -1 & 1 \\ 4 & 5 & 5 & 2 & 0 \\ 1 & -2 & 2 & -5 & 5 \end{array} \longrightarrow \begin{array}{c|ccc|c} 1 & 1 & 1 & -1 & 1 \\ 0 & \boxed{1} & 1 & 6 & -4 \\ 0 & -3 & 1 & -4 & 4 \end{array} \longrightarrow \begin{array}{c|ccc|c} 1 & 1 & 1 & -1 & 1 \\ 0 & 1 & 1 & 6 & -4 \\ 0 & 0 & 4 & 14 & -8 \end{array}$$

Wir sind fertig, da in jeder Zeile ein Pivotelement steht. Die Spalten 1, 2 und 3 sind pivotisiert, daher nennen wir x_1 , x_2 und x_3 die *abhängigen* Variablen und x_4 die *unabhängige* oder *freie* Variable. Die Lösungsmenge $\mathcal{L}$ lässt sich ablesen,

$$\mathcal{L} = \left\{ \left(5 + 7x_4, -2 - \frac{5}{2}x_4, -2 - \frac{7}{2}x_4, x_4 \right)^\top \in \mathbb{R}^4 : x_4 \in \mathbb{R} \right\}$$

und der Rang ist 3.

(b) Gesucht ist die Lösung des linearen Gleichungssystems ($m = n$)

$$\begin{array}{cccccc} x_1 & + & 2x_2 & + & 3x_3 & + & x_4 & = & 1 \\ x_1 & + & 3x_2 & + & 3x_3 & + & 2x_4 & = & 0 \\ 2x_1 & + & 4x_2 & + & 3x_3 & + & 3x_4 & = & 0 \\ x_1 & + & x_2 & + & x_3 & + & x_4 & = & 0. \end{array}$$

Wir stellen das Schema auf und wählen das erste Element in der letzten Zeile als erstes Pivotelement.

$$\begin{array}{c|cccc|c} 1 & 2 & 3 & 1 & 1 \\ 1 & 3 & 3 & 2 & 0 \\ 2 & 4 & 3 & 3 & 0 \\ \boxed{1} & 1 & 1 & 1 & 0 \end{array} \longrightarrow \begin{array}{c|cccc|c} 0 & 1 & 2 & 0 & 1 \\ 0 & 2 & 2 & \boxed{1} & 0 \\ 0 & 2 & 1 & 1 & 0 \\ 1 & 1 & 1 & 1 & 0 \end{array} \longrightarrow \begin{array}{c|cccc|c} 0 & 1 & 2 & 0 & 1 \\ 0 & 2 & 2 & 1 & 0 \\ 0 & 0 & \boxed{-1} & 0 & 0 \\ 1 & 1 & 1 & 1 & 0 \end{array} \longrightarrow \begin{array}{c|cccc|c} 0 & \boxed{1} & 0 & 0 & 1 \\ 0 & 2 & 2 & 1 & 0 \\ 0 & 0 & -1 & 0 & 0 \\ 1 & 1 & 1 & 1 & 0 \end{array}$$

Also haben wir das äquivalente Gleichungssystem

$$\begin{array}{cccccc} x_2 & & & & & = & 1 \\ 2x_2 & + & 2x_3 & + & x_4 & = & 0 \\ & & - & x_3 & & = & 0 \\ x_1 & + & x_2 & + & x_3 & + & x_4 & = & 0 \end{array}$$

und erhalten nach Rücksubstitution die eindeutig bestimmte Lösung $x = (1, 1, 0, -2)^\top$. Der Rang ist hier 4. Beachten Sie zum Beispiel die Wahl des ersten Pivotelements. Es wurde die vierte Zeile gewählt; denn zum einen ist das Pivotelement gleich Eins (also treten nur einfache Multiplikationen im ersten Schritt auf), zweitens sind alle weiteren Elemente in der Zeile gleich Eins (also auch hier nur einfache Rechnungen beim Eliminieren) und drittens steht auf der rechten Seite eine Null (also sind keine weiteren Rechnungen auf der rechten Seite in diesem Schritt erforderlich)

(c) Ein weiteres Beispiel mit $m > n$:

$$\begin{array}{cccccc} x_1 & - & x_2 & + & 2x_3 & = & 2 \\ x_1 & + & 2x_2 & - & x_3 & = & 1 \\ 2x_1 & - & x_2 & + & x_3 & = & 0 \\ 3x_1 & + & 3x_2 & & & = & \alpha \end{array} \cdot \text{ Also haben wir } \begin{array}{c|ccc|c} 1 & -1 & 2 & 2 \\ 1 & 2 & -1 & 1 \\ 2 & -1 & \boxed{1} & 0 \\ 3 & 3 & 0 & \alpha \end{array}$$

$$\begin{array}{ccc} \begin{array}{ccc|c} -3 & \boxed{1} & 0 & 2 \\ 3 & 1 & 0 & 1 \\ 2 & -1 & \mathbf{1} & 0 \\ 3 & 3 & 0 & \alpha \end{array} & \longrightarrow & \begin{array}{ccc|c} -3 & \mathbf{1} & 0 & 2 \\ \boxed{6} & 0 & 0 & -1 \\ 2 & -1 & \mathbf{1} & 0 \\ 12 & 0 & 0 & \alpha - 6 \end{array} & \longrightarrow & \begin{array}{ccc|c} -3 & \mathbf{1} & 0 & 2 \\ \boxed{6} & 0 & 0 & -1 \\ 2 & -1 & \mathbf{1} & 0 \\ 0 & 0 & 0 & \alpha - 4 \end{array} \end{array}$$

Mit Rücksubstitution ergibt sich $x = (-1/6, 3/2, 11/6)^\top$ als eindeutig bestimmte Lösung, falls $\alpha = 4$ gilt. Ist $\alpha \neq 4$, so ist die letzte Gleichung, $0x_1 + 0x_2 + 0x_3 + 0x_4 = \alpha - 4$, nicht erfüllbar, d.h. das Gleichungssystem hat in diesem Fall keine Lösung. In jedem Fall ist der Rang 3.

(d) Analog lassen sich lineare Gleichungssysteme in $\mathbb{C}$ betrachten etwa

$$\begin{array}{l} x_1 + x_2 + ix_3 = 1 \\ x_1 - x_2 - 2x_3 = 0 \\ ix_1 + x_2 + x_3 = i \end{array} \quad \text{Also haben wir} \quad \begin{array}{ccc|c} 1 & 1 & i & 1 \\ \boxed{1} & -1 & -2 & 0 \\ i & 1 & 1 & i \end{array}$$

$$\longrightarrow \begin{array}{ccc|c} 0 & \boxed{2} & 2+i & 1 \\ 1 & -1 & -2 & 0 \\ 0 & 1+i & 1+2i & i \end{array} \longrightarrow \begin{array}{ccc|c} 0 & 2 & 2+i & 1 \\ 1 & -1 & -2 & 0 \\ 0 & 0 & \frac{1}{2} + \frac{1}{2}i & -\frac{1}{2} + \frac{1}{2}i \end{array}$$

Also ist das LGS äquivalent zu

$$\begin{array}{rcl} 2x_2 + (2+i)x_3 & = & 1 \\ x_1 + x_2 - 2x_3 & = & 0 \\ (1+i)x_3 & = & -1+i. \end{array}$$

und es folgt durch Auflösen der dritten, der ersten und danach der zweiten Gleichung die Lösung $x = (1+i, 1-i, i)^\top$.

1.3 Linearkombinationen und lineare Unabhängigkeit

Lineare Gleichungssysteme lassen sich auch anders lesen. Betrachten wir etwa das Beispiel 1.11(b), so lassen sich die vier Gleichungen auch durch Vektoren im $\mathbb{R}^4$ beschreiben,

$$x_1 \begin{pmatrix} 1 \\ 1 \\ 2 \\ 1 \end{pmatrix} + x_2 \begin{pmatrix} 2 \\ 3 \\ 4 \\ 1 \end{pmatrix} + x_3 \begin{pmatrix} 3 \\ 3 \\ 3 \\ 1 \end{pmatrix} + x_4 \begin{pmatrix} 1 \\ 2 \\ 3 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}.$$

Dies gilt allgemein. Fassen wir im Gleichungssystem

$$\sum_{j=1}^n a_{ij}x_j = b_i, \quad i = 1, \dots, m,$$

die Koeffizienten zu n Vektoren $a_{*j} = (a_{1j}, a_{2j}, \dots, a_{mj})^\top \in \mathbb{R}^m$ zusammen, so sind Faktoren x_j gesucht, mit denen wir die rechte Seite $b = (b_1, \dots, b_m)^\top$ als Kombination

$$\sum_{j=1}^n x_j a_{*j} = b$$

schreiben können.

Definition 1.12 (a) Es seien $v^{(1)}, \dots, v^{(m)} \in \mathbb{K}^n$ beliebige Vektoren in $\mathbb{K}^n$. Ein Ausdruck der Form

$$\sum_{j=1}^m \lambda_j v^{(j)} := \lambda_1 v^{(1)} + \lambda_2 v^{(2)} + \dots + \lambda_m v^{(m)}$$

mit **Koeffizienten** $\lambda_j \in \mathbb{K}$ heißt **Linearkombination** der Vektoren $v^{(1)}, \dots, v^{(m)}$.

(b) Für $A \subseteq \mathbb{K}^n$ definieren wir die Menge

$$\text{span } A := \left\{ \sum_{j=1}^m \lambda_j a^{(j)} : \lambda_j \in \mathbb{K}, a^{(j)} \in A, j = 1, \dots, m, m \in \mathbb{N} \right\}.$$

$\text{span } A$ ist ein linearer Unterraum von $\mathbb{K}^n$ und wird der von A **aufgespannte Unterraum** genannt.

Es ist zu beweisen, dass $\text{span } A$ tatsächlich ein linearer Unterraum ist. Dies ist aber der Fall, da die Summe von zwei Linearkombinationen und die skalare Multiplikation einer Linearkombination wieder Linearkombinationen sind.

Im allgemeinen darf A aus unendlich vielen Elementen bestehen. Besteht A aber nur aus endlich vielen Elementen, etwa $A = \{a^{(1)}, \dots, a^{(m)}\}$, so ist

$$\text{span } A = \left\{ \sum_{j=1}^m \lambda_j a^{(j)} : \lambda_j \in \mathbb{K} \right\}.$$

Beispiele 1.13 (a) Sei $a = (2, 1)^\top \in \mathbb{R}^2$. Dann ist

$$\text{span } \{a\} = \left\{ t \begin{pmatrix} 2 \\ 1 \end{pmatrix} : t \in \mathbb{R} \right\},$$

also eine Gerade im $\mathbb{R}^2$ durch die Punkte 0 und $(2, 1)^\top$.

(b) Es seien $a = (1, 1, 1)^\top$, $b = (2, 0, 1)^\top$ zwei Vektoren des $\mathbb{R}^3$. Dann beschreibt

$$\text{span } \{a, b\} = \{t(1, 1, 1)^\top + s(2, 0, 1)^\top : t, s \in \mathbb{R}\},$$

die Ebene im $\mathbb{R}^3$ durch die Punkte 0, $(1, 1, 1)^\top$ und $(2, 0, 1)^\top$.

Es werden sich endliche Mengen $A = \{v^{(1)}, v^{(2)}, \dots, v^{(m)}\} \subseteq \mathbb{K}^n$ von Vektoren als besonders wichtig erweisen, bei denen sich die Elemente von $\text{span } A$ auf nur genau eine Weise als Linearkombination der $v^{(j)}$, $j = 1, \dots, n$ darstellen lassen.

Definition 1.14 (Lineare Unabhängigkeit)

Eine Menge von endlich vielen Vektoren $\{v^{(1)}, v^{(2)}, \dots, v^{(m)}\} \subseteq \mathbb{K}^n$ heißt **linear unabhängig**, wenn sich das neutrale Element 0 **nur** als triviale Linearkombination, also als

$$0 = \sum_{j=1}^m \alpha_j v^{(j)} \quad \text{mit} \quad \alpha_1 = \alpha_2 = \dots = \alpha_m = 0,$$

schreiben lässt. Andernfalls heißt die Menge **linear abhängig**.

Man nennt auch häufig die Vektoren selbst linear unabhängig bzw. linear abhängig.
 Eine linear unabhängige Teilmenge des $\mathbb{K}^n$ hat in der Tat die gewünschte Eigenschaft, dass die Darstellungen der Vektoren als Linearkombinationen eindeutig sind.

Satz 1.15 Die Menge $A = \{v^{(1)}, v^{(2)}, \dots, v^{(m)}\} \subseteq \mathbb{K}^n$ ist genau dann linear unabhängig, wenn jedes $x \in \text{span } A$ eine eindeutige Darstellung als Linearkombination der $v^{(j)}$, $j = 1, \dots, m$, besitzt.

Beweis: Besitzt jedes Element aus $\text{span } A$ eine eindeutige Darstellung als Linearkombination der $v^{(j)}$, $j = 1, \dots, m$, so gilt dies insbesondere für den Nullvektor. Somit ist A linear unabhängig. Sei umgekehrt A linear unabhängig. Betrachte $x \in \text{span } A$ mit

$$x = \sum_{j=1}^m c_j v^{(j)} = \sum_{j=1}^m d_j v^{(j)}$$

für geeignete Zahlen $c_j, d_j \in \mathbb{K}$, $j = 1, \dots, m$. Dann ist

$$0 = \sum_{j=1}^m c_j v^{(j)} - \sum_{j=1}^m d_j v^{(j)} = \sum_{j=1}^m (c_j - d_j) v^{(j)}.$$

Da A linear unabhängig ist, folgt $c_j - d_j = 0$, d.h. $c_j = d_j$ für $j = 1, \dots, m$. Die beiden Darstellungen von x sind somit gleich. $\square$

Beispiel 1.16 (a) Seien $a, b \in \mathbb{R}^2$ mit $a = (1, 2)^\top$ und $b = (1, 1)^\top$ gegeben. Dann ist $\{a, b\}$ linear unabhängig, denn aus

$$\lambda \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ 1 \end{pmatrix} = 0$$

folgt mit den ersten Koordinaten $\lambda = -\mu$ und anschließend mit den zweiten $\lambda = \mu = 0$. Somit lässt sich jedes

$$z \in \text{span } \{a, b\} = \left\{ \lambda \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} 1 \\ 1 \end{pmatrix} : \lambda, \mu \in \mathbb{R} \right\}.$$

in eindeutiger Weise als Linearkombination von a und b schreiben.

Wir berechnen die Koeffizienten und zeigen dabei auch $\text{span } \{a, b\} = \mathbb{R}^2$: Sei $z = (z_1, z_2)^\top \in \mathbb{R}^2$ beliebig. Wir müssen zeigen, dass $\lambda, \mu \in \mathbb{R}$ existieren mit $\lambda a + \mu b = z$, d.h.

$$\begin{aligned} \lambda + \mu &= z_1, \\ 2\lambda + \mu &= z_2. \end{aligned}$$

Dies sind zwei lineare Gleichungen in den zwei Unbekannten λ und μ . Subtraktion der Gleichungen liefert $\lambda = z_2 - z_1$ und hieraus $\mu = 2z_1 - z_2$.

(b) Wir betrachten $c, d \in \mathbb{R}^2$ mit $c = (1, 2)^\top$ und $d = (2, 4)^\top$. Dann ist $\{c, d\}$ linear abhängig, denn es gilt zum Beispiel

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

Der Nullvektor lässt sich zum Beispiel durch

$$0 = 2 \begin{pmatrix} 1 \\ 2 \end{pmatrix} - \begin{pmatrix} 2 \\ 4 \end{pmatrix}$$

nichttrivial linear kombinieren.

(c) Es seien $e = (i, 1, 1)^\top \in \mathbb{C}^3$ und $f = (1, 1, i)^\top \in \mathbb{C}^3$. Dann ist $\{e, f\} \subseteq \mathbb{C}^3$ linear unabhängig; denn sei $\lambda e + \mu f = 0$, d.h.

$$\begin{aligned}\lambda i + \mu &= 0, \\ \lambda + \mu &= 0, \\ \lambda + i\mu &= 0.\end{aligned}$$

Subtraktion der ersten beiden Gleichungen liefert $(i-1)\lambda = 0$, also $\lambda = 0$ und hieraus $\mu = 0$. Also sind $\{e, f\}$ linear unabhängig.

Im Beispiel 1.16 (a) haben wir zwei Vektoren gefunden, so dass jeder Vektor des $\mathbb{R}^2$ eindeutig als Linearkombination dieser beiden dargestellt werden kann. Wir definieren allgemein:

Definition 1.17 Sei $U \subseteq \mathbb{K}^n$ ein linearer Unterraum. Eine Menge $B = \{b^{(1)}, b^{(2)}, \dots, b^{(m)}\} \subseteq U$ von endlich vielen Elementen heißt **Basis** von U , wenn sich jedes Element $x \in U$ auf **genau eine Weise** als Linearkombination der $b^{(j)}$ schreiben lässt, d.h. wenn es eindeutig bestimmte $\alpha_j \in \mathbb{K}$ gibt mit

$$x = \sum_{j=1}^m \alpha_j b^{(j)}.$$

Bemerkungen: (a) Es lässt sich zeigen, dass jeder Unterraum von $\mathbb{K}^n$ eine Basis besitzt – und zwar nicht nur eine, sondern viele Basen. Wie wir aber gleich noch sehen werden, haben je zwei Basen gleich viele Elemente.

(b) Bei Linearkombinationen von Vektoren nennt man allgemein die Zahlen α_j Koeffizienten. In Zusammenhang mit einer Basis spezifiziert man dies und spricht von den **Koordinaten** bezüglich der Basis B .

(c) Nach Satz 1.15 können wir auch äquivalent formulieren: Genau dann ist $B \subseteq U$ eine Basis von U , wenn $U = \text{span } B$ und B linear unabhängig ist.

Beispiele 1.18 (a) Im $\mathbb{R}^n$ seien die folgenden Vektoren $e^{(j)}$, $j = 1, \dots, n$, gegeben:

$$e^{(j)} := (0, \dots, 0, 1, 0, \dots, 0)^\top \in \mathbb{R}^n,$$

wobei die 1 an der j -ten Stelle steht. Man nennt $e^{(j)}$ den **j-ten Koordinateneinheitsvektor**. Jedes $x = (x_1, \dots, x_n)^\top \in \mathbb{R}^n$ hat die Form

$$x = \sum_{j=1}^n x_j e^{(j)},$$

und dies ist auch die einzige Möglichkeit, x als Linearkombination der $e^{(j)}$ zu schreiben. Daher bildet $\{e^{(j)} : j = 1, \dots, n\}$ eine Basis des $\mathbb{R}^n$. Bezüglich dieser Basis fallen also die Komponenten des Vektors x mit den Koeffizienten in der Linearkombination zusammen. Daher heißt diese Basis die **Koordinateneinheitsbasis** des $\mathbb{R}^n$.

(b) Sei $U = \{x \in \mathbb{C}^3 : 3ix_1 - 2x_2 + 5x_3 = 0\}$. Überlegen Sie sich, dass dann U ein linearer Unterraum des $\mathbb{C}^3$ ist. Sei $a = (2, 3i, 0)^\top \in U$, $b = (0, 5, 2)^\top \in U$. Dann bildet $\{a, b\}$ eine Basis von U ; denn zu jedem $x \in U$ ist $\alpha, \beta \in \mathbb{C}$ zu bestimmen mit $x = \alpha a + \beta b$, also das Gleichungssystem zu lösen:

$$\begin{aligned}2\alpha &= x_1, \\ 3i\alpha + 5\beta &= x_2, \\ 2\beta &= x_3.\end{aligned}$$

Dies sind 3 Gleichungen und nur 2 Unbekannte α und β . Trotzdem ist es lösbar: Aus der 1. und 3. Gleichung erhalten wir $\alpha = x_1/2$ und $\beta = x_3/2$, und die zweite Gleichung ist auch erfüllt wegen $3i x_1 - 2 x_2 + 5 x_3 = 0$ (nachprüfen!). Die Koeffizienten α und β sind auch *eindeutig* durch x bestimmt.

(c) Wir wollen eine Basis bestimmen von

$$U = \text{span} \left\{ \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix}, \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} \right\} \subseteq \mathbb{R}^3.$$

Dazu untersuchen wir zunächst die drei Vektoren auf lineare Unabhängigkeit. Es ist die Lösungsmenge des homogenen linearen Gleichungssystems

$$\lambda_1 \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix} + \lambda_2 \begin{pmatrix} 2 \\ 1 \\ 1 \end{pmatrix} + \lambda_3 \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} = 0$$

gesucht. Wir führen das Gauß'sche Eliminationsverfahren durch, wobei wir die Pivotelemente markieren:

$$\begin{pmatrix} 1 & 2 & 1 \\ 2 & 1 & -1 \\ -1 & 1 & 2 \end{pmatrix} \longrightarrow \begin{pmatrix} \boxed{1} & 2 & 1 \\ 0 & -3 & -3 \\ 0 & 3 & 3 \end{pmatrix} \longrightarrow \begin{pmatrix} \boxed{1} & 2 & 1 \\ 0 & \boxed{-3} & -3 \\ 0 & 0 & 0 \end{pmatrix}$$

Somit können wir $\lambda_3 \in \mathbb{R}$ frei wählen und erhalten $\lambda_2 = -\lambda_3$, $\lambda_1 = \lambda_3$. Es gibt nicht-triviale Lösungen, die Vektoren sind somit linear abhängig, bilden also keine Basis von U .

Um eine solche zu finden, müssen wir aber keine neue Rechnung durchführen. Die bereits durchgeführte Rechnung zeigt nämlich, dass die ersten beiden Vektoren linear unabhängig sind. Lässt man nämlich den dritten Vektor weg, so entspricht dies der Vorgabe $\lambda_3 = 0$. Damit folgt unmittelbar $\lambda_1 = \lambda_2 = 0$ als einzige Lösung des verbleibenden LGS.

Allgemeines Verfahren: Sind m Vektoren $v^{(1)}, \dots, v^{(m)} \in \mathbb{K}^n$ auf lineare Unabhängigkeit zu überprüfen, so wendet man auf das lineare Gleichungssystem

$$\lambda_1 v^{(1)} + \dots + \lambda_m v^{(m)} = 0$$

für die Unbekannten $\lambda_1, \dots, \lambda_m$ das Gaußsche Eliminationsverfahren an. Die Spalten, in denen Pivotelemente gewählt werden, entsprechen einer größtmöglichen linear unabhängigen Teilmenge von $\{v^{(1)}, \dots, v^{(m)}\}$ und bilden somit eine Basis von $\text{span}\{v^{(1)}, \dots, v^{(m)}\}$.

Wie schon erwähnt, ist die Anzahl der Basiselemente eine charakteristische Größe für einen linearen Unterraum. Kennen wir sie, können wir einfacher zeigen, dass eine Menge eine Basis ist.

Satz 1.19 *Es sei U ein linearer Unterraum von $\mathbb{K}^n$ und $B = \{b^{(1)}, b^{(2)}, \dots, b^{(m)}\} \subseteq U$ eine Basis von U . Dann gilt:*

- (a) *Jede Menge $V \subseteq U$ mit mindestens $m + 1$ Elementen ist linear abhängig.*
- (b) *Hat $V \subseteq U$ m Elemente und ist linear unabhängig, so ist V eine Basis von U .*
- (c) *Jede Basis von U hat genau m Elemente.*

Beweis: (a) Seien $v^{(1)}, \dots, v^{(m+1)} \in V$ paarweise verschiedene Elemente von V . Wir stellen den Nullvektor als Linearkombination der $v^{(j)}$ dar,

$$c_1 v^{(1)} + c_2 v^{(2)} + \dots + c_{m+1} v^{(m+1)} = 0.$$

Zu zeigen ist, dass mindestens ein c_j von null verschieden gewählt werden kann. Da B eine Basis ist, kann jedes $v^{(j)}$ als Linearkombination der $b^{(k)}$ dargestellt werden,

$$v^{(j)} = \sum_{k=1}^m \alpha_{j,k} b^{(k)}, \quad j = 1, \dots, m+1.$$

Somit sind c_j zu finden mit

$$0 = \sum_{j=1}^{m+1} c_j v^{(j)} = \sum_{j=1}^{m+1} \sum_{k=1}^m c_j \alpha_{j,k} b^{(k)} = \sum_{k=1}^m \left(\sum_{j=1}^{m+1} c_j \alpha_{j,k} \right) b^{(k)}.$$

Da B als Basis linear unabhängig ist, folgen die Bedingungen

$$\sum_{j=1}^{m+1} c_j \alpha_{j,k} = 0, \quad k = 1, \dots, m.$$

Dies ist ein lineares Gleichungssystem mit m Gleichungen für die $m+1$ Unbekannten c_j . Eine der Unbekannten c_j kann also frei und damit von null verschieden gewählt werden.

(b) Wir schreiben $V = \{v^{(1)}, \dots, v^{(m)}\}$ und wählen $x \in U$ beliebig. Es ist zu zeigen, dass x als Linearkombination der $v^{(j)}$ geschrieben werden kann.

Nach Teil (a) ist die Mengen $\{x, v^{(1)}, \dots, v^{(m)}\}$ linear abhängig. Somit existieren Zahlen $\alpha_0, \dots, \alpha_m$, die nicht alle gleich null sind, mit

$$0 = \alpha_0 x + \sum_{j=1}^m \alpha_j v^{(j)}.$$

Ist $\alpha_0 = 0$, so folgt aus der linearen Unabhängigkeit der $v^{(j)}$ auch $\alpha_1 = \dots = \alpha_m = 0$. Somit gilt $\alpha_0 \neq 0$, und wir haben

$$x = -\frac{1}{\alpha_0} \sum_{j=1}^m \alpha_j v^{(j)}.$$

(c) Es ist nur noch zu zeigen, dass keine Basis V von U mit weniger als m Elementen existiert. Wäre dies aber der Fall, so wäre nach Teil (a) die Menge B linear abhängig. Dies ist ein Widerspruch dazu, dass B Basis von U ist. $\square$

Da die Anzahl der Elemente jeder Basis eine charakteristische Eigenschaft jedes Unterraums des $\mathbb{K}^n$ ist, verdient sie einen besonderen Namen.

Definition 1.20 Ist $U \subseteq \mathbb{K}^n$ ein linearer Unterraum, so heißt die Anzahl $m \in \mathbb{N}$ der Elemente jeder Basis von U die **Dimension** von U , als Formel

$$\dim U = m.$$

Bemerkungen:

- (a) Mit Satz 1.19 können wir nun formulieren: Ist die Dimension m von U bereits bekannt, so bilden m Elemente von U genau dann eine Basis von U , wenn sie linear unabhängig sind.
- (b) Die Koordinateneinheitsbasis $\{e^{(1)}, \dots, e^{(n)}\}$ des $\mathbb{R}^n$ aus Beispiel 1.18 besitzt n Elemente. Damit ist $\dim \mathbb{R}^n = n$.
- (c) Bei $\{e^{(1)}, \dots, e^{(n)}\}$ handelt es sich auch um eine Basis des $\mathbb{C}^n$. Also ist auch $\dim \mathbb{C}^n = n$.

Beispiele 1.21 (a) Die Vektoren $e = (i, 1, 1)^\top \in \mathbb{C}^3$ und $f = (1, 1, i)^\top \in \mathbb{C}^3$ aus Beispiel 1.16 sind linear unabhängig, bilden aber keine Basis des $\mathbb{C}^3$. Dazu wären drei linear unabhängige Vektoren nötig, etwa noch $g = (0, 0, i)^\top$. Rechnen Sie selbst nach, dass diese drei Vektoren linear unabhängig sind.

- (b) Die Menge $V = \{v^{(1)}, v^{(2)}, v^{(3)}, v^{(4)}\} \subseteq \mathbb{R}^3$ mit

$$v^{(1)} = \begin{pmatrix} 1 \\ 3 \\ -2 \end{pmatrix}, \quad v^{(2)} = \begin{pmatrix} -2 \\ -6 \\ 4 \end{pmatrix}, \quad v^{(3)} = \begin{pmatrix} 2 \\ 0 \\ 1 \end{pmatrix}, \quad v^{(4)} = \begin{pmatrix} 1 \\ -3 \\ 3 \end{pmatrix},$$

ist linear abhängig, da V vier Elemente enthält, aber $\dim \mathbb{R}^3 = 3$ ist. Um $\dim(\text{span } V)$ zu bestimmen, betrachten wir trotzdem das homogene lineare Gleichungssystem

$$\sum_{j=1}^4 \lambda_j v^{(j)} = 0.$$

Beim Gauß'schen Eliminationsverfahren markieren wir wieder die Pivotelemente:

$$\begin{pmatrix} 1 & -2 & 2 & 1 \\ 3 & -6 & 0 & -3 \\ -2 & 4 & 1 & 3 \end{pmatrix} \longrightarrow \begin{pmatrix} \boxed{1} & -2 & 2 & 1 \\ 0 & 0 & -6 & -6 \\ 0 & 0 & 5 & 5 \end{pmatrix} \longrightarrow \begin{pmatrix} \boxed{1} & -2 & 2 & 1 \\ 0 & 0 & \boxed{-6} & -6 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Somit ist $\{v^{(1)}, v^{(3)}\}$ eine Basis von $\text{span } V$, und es gilt $\dim \text{span } V = 2$.

1.4 Euklidische Vektorräume

Bislang haben wir für Vektoren nur die algebraischen Operationen „Addition“ (und Subtraktion) und die „Streckung“ benutzt. Alle unsere Darstellungen von Vektoren in der Ebene und im Raum haben auch eine **Länge** und je zwei von ihnen schließen einen **Winkel** ein. Dies können wir allgemein für $x, y \in \mathbb{R}^n$ definieren (bzw. noch allgemeiner für $x, y \in \mathbb{C}^n$).

Definition 1.22 (a) Für $x = (x_1, \dots, x_n)^\top \in \mathbb{C}^n$, $y = (y_1, \dots, y_n)^\top \in \mathbb{C}^n$ definieren wir das **Skalarprodukt** oder **innere Produkt** durch

$$\langle x, y \rangle := \sum_{j=1}^n x_j \overline{y_j} = x_1 \overline{y_1} + \dots + x_n \overline{y_n},$$

wobei $\overline{y_j}$ die konjugiert komplexe Zahl zu $y_j \in \mathbb{C}$ ist.

(b) Für $x \in \mathbb{C}^n$ definieren wir die **Norm** (die **Länge**, oder den **Betrag**) durch

$$\|x\| := \sqrt{\langle x, x \rangle} = \sqrt{\sum_{j=1}^n |x_j|^2}.$$

Hier ist $|x_j|$ der Betrag der komplexen Zahl x_j .

Bemerkung: Beachte, dass sich die Definitionen für reelle Vektoren auf die folgenden Definitionen reduzieren.

(a) Wenn $x = (x_1, \dots, x_n)^\top \in \mathbb{R}^n$, $y = (y_1, \dots, y_n)^\top \in \mathbb{R}^n$ gilt, so verwendet man eine spezielle Schreibweise mit einem Multiplikationspunkt. Das Skalarprodukt ist gegeben durch

$$x \cdot y := \langle x, y \rangle = \sum_{j=1}^n x_j y_j = x_1 y_1 + \dots + x_n y_n$$

gegeben.

(b) Wenn $x \in \mathbb{R}^n$ ist, so ist durch

$$\|x\| := \sqrt{x \cdot x} = \sqrt{\sum_{j=1}^n x_j^2}$$

die **euklidische Norm** definiert.

Der **euklidische Abstand** von zwei Punkten $x, y \in \mathbb{R}^n$ ist die Norm des Differenzvektors bzw. Verbindungsvektors,

$$\|x - y\| = \sqrt{\sum_{j=1}^n (x_j - y_j)^2}.$$

Als erstes rechnen wir die **Binomischen Formeln** nach. Es seien $x, y \in \mathbb{C}^n$. Dann ist

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle = \sum_{j=1}^n (x_j + y_j) \overline{(x_j + y_j)} \\ &= \sum_{j=1}^n |x_j|^2 + \sum_{j=1}^n |y_j|^2 + \sum_{j=1}^n x_j \overline{y_j} + \overline{x_j} y_j \\ &= \|x\|^2 + \|y\|^2 + 2 \operatorname{Re} \langle x, y \rangle, \end{aligned}$$

und genauso noch zwei weitere Identitäten, also zusammengefasst:

$$\begin{aligned} \|x + y\|^2 &= \|x\|^2 + \|y\|^2 + 2 \operatorname{Re} \langle x, y \rangle, \\ \|x - y\|^2 &= \|x\|^2 + \|y\|^2 - 2 \operatorname{Re} \langle x, y \rangle, \\ \langle x + y, x - y \rangle &= \|x\|^2 - \|y\|^2 - 2i \operatorname{Im} \langle x, y \rangle. \end{aligned}$$

Weitere wichtige Eigenschaften fasst der folgende Satz zusammen:

Satz 1.23 (*Eigenschaften des Skalarprodukts und der Norm*)

Für beliebige (komplexe) Vektoren gilt:

- (a) $\langle x, y \rangle = \overline{\langle y, x \rangle}$ *hermitesch*
- (b) $\langle x + y, z \rangle = \langle x, z \rangle + \langle y, z \rangle$ *distributiv*
- (c) $\lambda \langle x, y \rangle = \langle \lambda x, y \rangle = \langle x, \overline{\lambda} y \rangle$ *assoziativ*
- (d) $x \neq 0 \iff \langle x, x \rangle > 0$ *positiv definit*
- (e) $\|\lambda x\| = |\lambda| \|x\|$ *homogen*
- (f) $|\langle x, y \rangle| \leq \|x\| \|y\|$ *Cauchy-Schwarzsche Ungleichung*
- Gilt $|\langle x, y \rangle| = \|x\| \|y\|$, so sind x und y linear abhängig.*
- (g) $\|x + y\| \leq \|x\| + \|y\|$ *erste Dreiecksungleichung*
- (h) $\|x - y\| \geq |\|x\| - \|y\||$ *zweite Dreiecksungleichung*

Bemerkungen: (a) Ausgeschrieben heißt die Cauchy-Schwarzsche Ungleichung (f) nach dem Quadrieren:

$$\left| \sum_{j=1}^n x_j \overline{y_j} \right|^2 \leq \sum_{j=1}^n |x_j|^2 \sum_{j=1}^n |y_j|^2.$$

(b) Beachte, dass im reellen Fall, $\mathbb{R}^n$, die erste Eigenschaft Symmetrie bedeutet, d.h. das Skalarprodukt ist dann kommutativ, $x \cdot y = y \cdot x$.

Beweis (Auszug): (a)–(e) sind einfach zu zeigen. Interessant ist Aussage (f), d.h. die Cauchy-Schwarzsche Ungleichung. Wir führen den Beweis nur für reelle Vektoren durch. Seien also $x, y \in \mathbb{R}^n$ festgehalten, $y \neq 0$. Für $t \in \mathbb{R}$ definieren wir die Funktion

$$f(t) := \|x - ty\|^2 = \|x\|^2 + t^2 \|y\|^2 - 2t(x \cdot y).$$

Dann ist f ein reelles quadratisches Polynom in t und wegen (d) ist $f(t) \geq 0$ für alle $t \in \mathbb{R}$. Quadratische Ergänzung liefert:

$$0 \leq f(t) = \left(t \|y\| - \frac{x \cdot y}{\|y\|} \right)^2 + \|x\|^2 - \frac{[x \cdot y]^2}{\|y\|^2}.$$

Für $t_0 := [x \cdot y] / \|y\|^2$ ist $(\dots) = 0$, also

$$0 \leq f(t_0) = \|x\|^2 - \frac{[x \cdot y]^2}{\|y\|^2}.$$

Nach Multiplikation mit $\|y\|^2$ ist dies genau die Cauchy-Schwarzsche Ungleichung. Ist $|x \cdot y| = \|x\| \|y\|$, so ist $f(t) = (\dots)^2$, also $f(t_0) = 0$, also

$$0 = f(t_0) = \|x - t_0 y\|^2, \quad \text{d.h.} \quad x = t_0 y,$$

und x, y sind linear abhängig.

(g) Mit der Binomischen Formel und der Cauchy-Schwarzschen Ungleichung rechnen wir einfach aus:

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2 \operatorname{Re}\langle x, y \rangle \leq \|x\|^2 + \|y\|^2 + 2 \|x\| \|y\| = (\|x\| + \|y\|)^2.$$

(h) Dies führen wir auf die erste Dreiecksungleichung zurück:

$$\|x\| = \|(x - y) + y\| \leq \|x - y\| + \|y\|, \quad \text{also } \|x - y\| \geq \|x\| - \|y\|.$$

Genauso zeigt man $\|y - x\| \geq \|y\| - \|x\|$ und daher $\|x - y\| \geq |\|x\| - \|y\||$. $\square$

Wie beim Begriff des Vektorraums können wir $\mathbb{R}^n$ als ein mögliches (aber wichtigstes) Beispiel für einen sogenannten euklidischen Vektorraum betrachten. Jeder Vektorraum V , der ein Skalarprodukt hat, also ein Produkt $V \times V \rightarrow \mathbb{K}$ mit den Eigenschaften (a)–(d) des letzten Satzes, heißt **Prä - Hilbertraum**. Ist $\mathbb{K} = \mathbb{R}$, so nennt man V auch einen **euklidischen Vektorraum**, falls $\mathbb{K} = \mathbb{C}$ einen **unitären Vektorraum**. In jedem euklidischen oder unitären Vektorraum kann man eine Norm über $\|x\| := \sqrt{\langle x, x \rangle}$ einführen. Den entstehenden Raum nennt man dann einen **normierten Vektorraum**, und es gelten die Eigenschaften (e)–(h) des obigen Satzes. Wir werden noch häufiger auf diese algebraischen Strukturen stoßen.

Aus der Cauchy-Schwarzschen Ungleichung folgt für beliebige Vektoren $x, y \in \mathbb{R}^n$, $x, y \neq 0$:

$$\left| \frac{x}{\|x\|} \cdot \frac{y}{\|y\|} \right| \leq 1.$$

Daher gibt es einen eindeutig bestimmten Winkel $\omega \in [0, \pi]$ mit

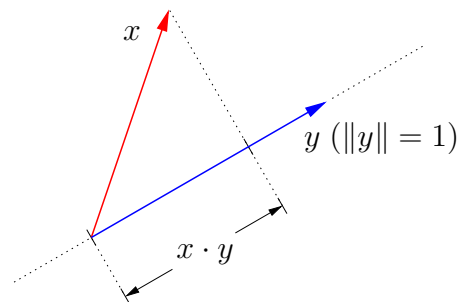
$$\frac{x}{\|x\|} \cdot \frac{y}{\|y\|} = \cos \omega.$$

Der Winkel ω heißt der von den Vektoren x und y **eingeschlossene Winkel**. Der Fall $x \cdot y = 0$ gehört zu $\omega = \pi/2$ (im Bogenmaß, d.h. 90°). In diesem Fall stehen also die beiden Vektoren x und y **senkrecht** aufeinander (**orthogonal**). Wir schreiben auch $x \perp y$.

Bemerkung: Wir haben den Winkel zwischen zwei Vektoren für beliebige Dimensionen n erklärt. Aus der Definition folgt sofort

$$x \cdot y = \|x\| \|y\| \cos \omega,$$

und dies ist die bekannte Formel aus der Schule. Die Formel wurde benutzt, um mit dem Längen- und Winkelbegriff das Skalarprodukt zu definieren. Wir haben hier den Spieß umgedreht und zuerst das Skalarprodukt definiert und damit Länge und Winkel. Eine wichtige anschauliche Interpretation des Skalarproduktes ergibt sich, wenn wir annehmen, dass einer der beiden Vektoren etwa y die Länge $\|y\| = 1$ hat. Dann ist das Skalarprodukt mit einem Vektor $x \in \mathbb{R}^n$ gerade die Länge der orthogonalen Projektion von x auf die Richtung y (s. Abbildung).



Das Skalarprodukt $x \cdot y$ mit $\|y\| = 1$ ist die Länge der orthogonalen Projektion von x auf y .

Beispiele 1.24

(a) Die beiden Vektoren $x, y \in \mathbb{R}^4$ seien gegeben durch

$$x = (1, 2, 3, 4)^\top, \quad y = (1, 0, 2, 0)^\top.$$

Dann ist $\|x\|^2 = 30$, $\|y\|^2 = 5$ und $x \cdot y = 7$. Der Winkel ω zwischen x und y wird aus $\cos \omega = 7/\sqrt{30 \cdot 5}$ berechnet und ist (im Bogenmaß) $\omega \approx 0.9624$ oder 55.1418 Grad.

(b) Für die Koordinateneinheitsvektoren $e^{(j)}$, $j = 1, \dots, n$, gilt $e^{(j)} \cdot e^{(k)} = 0$ für $j \neq k$ und $\|e^{(j)}\| = 1$ für alle j . Diese Vektoren stehen also alle senkrecht aufeinander und haben die Länge 1. Es gilt die Darstellung:

$$\text{Mit } x = (x_1, \dots, x_n)^\top, \quad \text{ist } x = \sum_{j=1}^n x_j e^{(j)} = \sum_{j=1}^n (x \cdot e^{(j)}) e^{(j)}.$$

Wir kommen jetzt zurück zu **Geraden** und **Ebenen** im $\mathbb{R}^n$. Wir hatten gesehen, dass diese in Parameterform angegeben werden können, etwa

$$G = \{x = p + tq : t \in \mathbb{R}\} \quad \text{bzw.} \quad E = \{x = p + sq + tr : s, t \in \mathbb{R}\}$$

mit $p \in \mathbb{R}^n$, $q \in \mathbb{R}^n \setminus \{0\}$ bzw. $p \in \mathbb{R}^n$, $\{q, r\} \subseteq \mathbb{R}^n$ linear unabhängig. Eine alternative Darstellung ist für Geraden im $\mathbb{R}^2$ und für Ebenen im $\mathbb{R}^3$ die **Normalform**.

Normalform einer Geraden im $\mathbb{R}^2$: Zu $q \in \mathbb{R}^2 \setminus \{0\}$ existiert $a \in \mathbb{R}^2 \setminus \{0\}$ mit $a \cdot q = 0$, denn dies ist ein lineares Gleichungssystem mit einer Gleichung für die zwei Unbekannten a_1, a_2 . Es hat also eine freie Variable. Mit $\rho = a \cdot p$, gilt dann für $x \in G$

$$a \cdot x = a \cdot (p + tq) = \rho + t a \cdot q = \rho.$$

Gilt umgekehrt für ein $x \in \mathbb{R}^2$ die Gleichung $a \cdot x = \rho$, so ist $a \cdot (x - p) = 0$. Das homogene lineare Gleichungssystem $a \cdot z = 0$ mit einer Gleichung und den Unbekannten z_1, z_2 hat die Lösungsmenge $\text{span}\{q\}$ mit einem $q \in \mathbb{R}^2 \setminus \{0\}$. Man hat also

$$x - p = z = tq, \quad t \in \mathbb{R}.$$

Diese Überlegung zeigt, dass man aus der Parameterform $x = p + tq$ stets die Normalform $a \cdot x = \rho$ ableiten kann, und umgekehrt. Es ist

$$G = \{x = p + tq : t \in \mathbb{R}\} = \{x \in \mathbb{R}^2 : a \cdot x = \rho\}$$

Den Vektor a , der bis auf die Länge eindeutig bestimmt ist, nennt man die **Normale** zu G .

Normalform einer Ebene im $\mathbb{R}^3$: Das Vorgehen ist ganz analog. Man findet zu den linear unabhängigen Richtungsvektoren $q, r \in \mathbb{R}^3$ ein $a \in \mathbb{R}^3 \setminus \{0\}$ mit $a \cdot q = a \cdot r = 0$. Umgekehrt findet man zu gegebenem a zwei linear unabhängige Vektoren q, r , die die Lösungsmenge von $a \cdot z = 0$ aufspannen. Damit hat man die zwei äquivalenten Darstellungen einer Ebene im $\mathbb{R}^3$,

$$E = \{x = p + sq + tr : s, t \in \mathbb{R}\} = \{x \in \mathbb{R}^3 : a \cdot x = \rho\}.$$

Auch hier heißt a die **Normale** zu E .

Auch im $\mathbb{R}^n$ können die entsprechenden Mengen betrachtet werden.

Satz 1.25 Sei $a \in \mathbb{R}^n$, $a \neq 0$. Die Menge $A_0 := \{x \in \mathbb{R}^n : a \cdot x = 0\}$ ist ein linearer Unterraum des $\mathbb{R}^n$. Für $\rho \in \mathbb{R}$ ist die Menge $A := \{x \in \mathbb{R}^n : a \cdot x = \rho\}$ ein affiner Unterraum.

Beweis: Es seien $x, y \in A_0$. Dann ist $a \cdot (x + y) = a \cdot x + a \cdot y = 0$ und $a \cdot (\lambda x) = \lambda a \cdot x = 0$. Also ist A_0 ein Unterraum. Wir setzen

$$\hat{x} = \frac{\rho}{\|a\|^2} a.$$

Dann ist $\hat{x} \cdot a = \rho$, also $\hat{x} \in A$. Nun zeigen wir $A = \hat{x} + A_0$.

(i) Ist $x \in A$, so ist $x = \hat{x} + (x - \hat{x})$ und $x - \hat{x} \in A_0$, denn $a \cdot (x - \hat{x}) = a \cdot x - a \cdot \hat{x} = \rho - \rho = 0$, d.h. $x \in \hat{x} + A_0$.

(ii) Ist umgekehrt $z \in A_0$, so ist $x = z + \hat{x} \in A$, denn $a \cdot (z + \hat{x}) = a \cdot z + a \cdot \hat{x} = \rho$. □

Das Zusammenspiel der beiden Darstellungsformen von Geraden und Ebenen ist bei der Lösung von Abstandsaufgaben besonders interessant. Wir definieren den **Abstand** eines Punktes $x \in \mathbb{R}^n$ von einem affinen Unterraum $A \subseteq \mathbb{R}^n$ also

$$d(x, A) = \min\{\|y - x\| : y \in A\}.$$

Ausführlich wollen wir $d(x, A)$ für den Fall bestimmen, dass A eine Gerade durch 0 im $\mathbb{R}^n$ ist, d.h. ein eindimensionaler linearer Unterraum des $\mathbb{R}^n$. Wir nehmen ferner zunächst an, dass der Richtungsvektor q normiert ist, also $A = \{tq : t \in \mathbb{R}\}$, und es gilt $\|q\| = 1$. Ist $x \in \mathbb{R}^n$ und $y = tq \in A$, so gilt mit den Eigenschaften des Skalarprodukts und einer quadratischen Ergänzung

$$\begin{aligned} \|x - y\|^2 &= \|x\|^2 - 2x \cdot y + \|y\|^2 = \|x\|^2 - 2tx \cdot q + t^2 \\ &= \|x\|^2 - (x \cdot q)^2 + (x \cdot q - t)^2. \end{aligned}$$

Hieraus erkennen wir, dass der Abstand $d(x, A)$ wohldefiniert ist, denn der Ausdruck nimmt für $t = x \cdot q$ tatsächlich sein Minimum an. Es ist

$$d(x, A) = \|x - \hat{x}\| \quad \text{mit} \quad \hat{x} = (x \cdot q)q.$$

Ferner ist

$$(x - \hat{x}) \cdot q = x \cdot q - (x \cdot q)q \cdot q = 0,$$

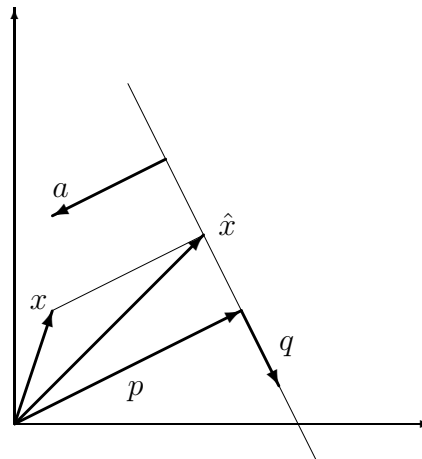
d.h. der Verbindungsvektor von x und $\hat{x}$ ist orthogonal zu A . Wir nennen $\hat{x}$ die **orthogonale Projektion** von x auf A , oder auch **Lotfußpunkt** zu x in A .

Allgemeiner erhält man für einen Richtungsvektor q mit $\|q\| \neq 1$ den Lotfußpunkt

$$\hat{x} = \frac{x \cdot q}{\|q\|^2} q.$$

Ist $A = \{y = p + tq : t \in \mathbb{R}\}$ eine beliebige Gerade, so verschiebt man zunächst x und A um $-p$, wendet dann die bisher hergeleiteten Formeln an und verschiebt schließlich wieder zurück mit p . Dies ergibt den Lotfußpunkt

$$\hat{x} = p + \frac{(x - p) \cdot q}{\|q\|^2} q.$$



Orthogonale Projektion bzw. Lotfußpunkt $\hat{x}$

Ist nun $A \subseteq \mathbb{R}^2$ in Normalform $a \cdot y = \rho$ mit Normale a gegeben, so ist $x - \hat{x} = d a$ mit einem $d \in \mathbb{R}$. Hierbei ist $|d| = d(x, A) / \|a\|$. Es folgt aus $\hat{x} \in A$ die Gleichung

$$a \cdot x - \rho = a \cdot x - a \cdot \hat{x} = a \cdot (x - \hat{x}) = a \cdot (d a) = d \|a\|^2.$$

Somit ergibt sich

$$d(x, A) = \frac{|a \cdot x - \rho|}{\|a\|}.$$

Das Vorgehen beim Bestimmen des Abstands eines Punktes von einer Ebene E ist theoretisch analog. Man kann hier zunächst den Fall betrachten, dass die Ebene 0 enthält und die Richtungsvektoren q, r orthogonal sind und die Norm 1 besitzen. Dann gilt für den Abstand von $x \in \mathbb{R}^n$ von $y = sq + tr \in E$

$$\begin{aligned} \|x - y\|^2 &= \|x\|^2 - 2x \cdot y + \|y\|^2 = \|x\|^2 - 2sx \cdot q - 2tx \cdot r + s^2 \|q\|^2 + 2st q \cdot r + t^2 \|r\|^2 \\ &= \|x\|^2 - 2sx \cdot q + s^2 - 2tx \cdot r + t^2 \\ &= \|x\|^2 - (x \cdot q)^2 - (x \cdot r)^2 + (x \cdot q - s)^2 + (x \cdot r - t)^2. \end{aligned}$$

Wieder sehen wir, dass der Abstand für genau eine Wahl der Parameter s, t minimal wird, nämlich für $s = x \cdot q, t = x \cdot r$. Mit $\hat{x} = (x \cdot q)q + (x \cdot r)r$ folgt die Orthogonalität von $x - \hat{x}$ zu E ,

$$(x - \hat{x}) \cdot q = x \cdot q - (x \cdot q)q \cdot q - (x \cdot r)r \cdot q = 0,$$

$$(x - \hat{x}) \cdot r = x \cdot r - (x \cdot q)q \cdot r - (x \cdot r)r \cdot r = 0.$$

Wir nennen $\hat{x}$ wieder die **orthogonale Projektion** von x auf E , oder **Lotfußpunkt**.

Es wäre nun auch wieder möglich, eine explizite Formel für der Lotfußpunkt für eine Ebene in allgemeiner Parameterform anzugeben, aber diese wird kompliziert und unhandlich. Praktisch einfacher ist es, die Orthogonalitätseigenschaft auszunutzen und direkt $\hat{x} = p + sq + tr$ mit $(x - \hat{x}) \cdot q = (x - \hat{x}) \cdot r = 0$ zu suchen. Dies ist ein lineares Gleichungssystem für die Parameter s, t . Der Abstand $d(x, E)$ ergibt sich dann als Abstand von x und $\hat{x}$.

Ist $E \subseteq \mathbb{R}^3$, so können wir E auch in der Normalform $a \cdot x = \rho$ betrachten. Dann ist mit derselben Rechnung wie bei Geraden im $\mathbb{R}^2$

$$d(x, A) = \frac{|a \cdot x - \rho|}{\|a\|}.$$

Auch in beliebigen Dimensionen führen analoge Überlegungen zu einer entsprechenden Abstandsformel. Damit erhalten wir den folgenden Satz.

Satz 1.26 Gilt für $a \in \mathbb{R}^n \setminus \{0\}$, $\rho \in \mathbb{R}$ zusätzlich, dass $\|a\| = 1$ und $\rho > 0$ ist, so sagt man, der affine Unterraum $A = \{y \in \mathbb{R}^n : a \cdot y - \rho = 0\}$ liegt in **Hesse'scher Normalform** vor. In diesem Fall gilt

$$d(x, A) = |a \cdot x - \rho| \quad \text{für alle } x \in \mathbb{R}^n.$$

Der Ausdruck im Betrag ist genau dann negativ, wenn x auf „derselben Seite“ von A wie 0 liegt.

Die Bestimmung von Abständen in vielen weiteren geometrischen Situationen lässt sich ähnlich durchführen.

- (a) Bestimmung des Abstands zwischen **zwei Geraden** $x = p + ta$, $t \in \mathbb{R}$, und $y = q + sb$, $s \in \mathbb{R}$:

Hier müssen analog zwei Parameter $\hat{s} \in \mathbb{R}$ und $\hat{t} \in \mathbb{R}$ gesucht werden, so dass $(\hat{x} - \hat{y}) \cdot a = 0$ und $(\hat{x} - \hat{y}) \cdot b = 0$ ist. Dies ergibt zwei Gleichungen in den Unbekannten $\hat{t}$ und $\hat{s}$. Das Gleichungssystem ist genau dann eindeutig lösbar, wenn a und b linear unabhängig sind, d.h. die Geraden nicht parallel sind.

Falls sich die beiden Geraden schneiden, so ist der **Schnittwinkel** ω zwischen ihnen gegeben durch

$$\cos \omega = \frac{a \cdot b}{\|a\| \|b\|}.$$

- (b) Genauso können wir den Abstand zwischen einer Geraden und einer Ebene bestimmen. Wir gehen darauf nicht näher ein, sondern verweisen auf die Übungen.

2 Lineare Abbildungen

2.1 Lineare Abbildungen und Matrizen

Die Struktur eines Vektorraums wird bestimmt durch die beiden Operationen *Addition* und *skalare Multiplikation*. Betrachtet man nun Abbildungen zwischen zwei Vektorräumen, so sind solche Abbildungen besonders interessant, die sich mit diesen Operationen „gut vertragen“.

Definition 2.1 Eine Abbildung $\mathcal{A} : V \rightarrow W$ von einem Vektorraum V in einen Vektorraum W mit den Eigenschaften

$$\mathcal{A}(x + y) = \mathcal{A}(x) + \mathcal{A}(y) \quad \text{und} \quad \mathcal{A}(\lambda x) = \lambda \mathcal{A}(x)$$

für alle $x, y \in V$ und $\lambda \in \mathbb{K}$ heißt **linear**.

Bei linearen Abbildungen ist es dasselbe, ob man erst addiert und dann abbildet, oder erst abbildet und dann addiert.

Wir haben diesen Begriff auch schon bei Differentialgleichungen verwendet, dabei handelt es sich um dieselben Bedingungen, aber im entsprechenden Vektorraum von Funktionen, d.h. die Funktionen spielen dort die Rolle der Vektoren und die linke Seite der Differentialgleichung lässt sich als Abbildung (Operator) auffassen, die Funktionen auf Funktionen abbildet. Dieser Operator entspricht der Abbildung $\mathcal{A}$ in der Definition 2.1.

Beispiele: (a) Im Vektorraum $V = \mathbb{R}$ sind die linearen Abbildungen $\mathcal{A} : V \rightarrow V$ genau diejenigen mit $\mathcal{A}(x) = ax$, $x \in V$, für ein $a \in \mathbb{R}$. Allgemeiner nennt man Polynome ersten Grades, $\mathcal{A}(x) = ax + b$, *lineare Funktionen*. Sie haben aber streng genommen nicht die Linearitätseigenschaft sondern unterscheiden sich von einer linearen Abbildung durch eine Konstante. Präziser müsste man sie *affin linear* nennen.

(b) Die Abbildung $x = (x_1, \dots, x_n)^\top \mapsto x_1 = x \cdot e^{(1)}$ ist eine lineare Abbildung vom $\mathbb{R}^n$ nach $\mathbb{R}$ und heißt Projektion auf die erste Komponente. Hierbei bezeichne $e^{(1)}$ den ersten Koordinateneinheitsvektor.

(c) Die Abbildung $\mathcal{A}(x) = (-x_2, x_1)^\top$ ist linear und beschreibt eine Drehung um den Ursprung um 90° gegen den Uhrzeigersinn.

Wir konzentrieren uns in diesem Kapitel auf lineare Abbildungen auf den endlich dimensionalen Räumen $\mathbb{K}^n$. Zu einem Vektor $x \in \mathbb{K}^n$ betrachte man die Darstellung in der Koordinateneinheitsbasis,

$$x = \sum_{j=1}^n x_j e^{(j)},$$

so gilt für die Anwendung einer linearen Abbildung $\mathcal{A} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ auf $x \in \mathbb{K}^n$:

$$\mathcal{A}(x) = \sum_{j=1}^n x_j \mathcal{A}(e^{(j)}) = \begin{pmatrix} | & & | \\ \mathcal{A}(e^{(1)}) & \cdots & \mathcal{A}(e^{(n)}) \\ | & & | \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

Sind umgekehrt $v^{(1)}, \dots, v^{(n)} \in \mathbb{K}^m$, so wird durch

$$\mathcal{A}(x) = \sum_{j=1}^n x_j v^{(j)}, \quad x = (x_1, \dots, x_n)^\top \in \mathbb{K}^n,$$

eine lineare Abbildung definiert, die $\mathcal{A}(e^{(j)}) = v^{(j)}$, $j = 1, \dots, n$, erfüllt. Also ist eine lineare Abbildung vollständig durch die Bilder $\mathcal{A}(e^{(j)}) = (a_{1j}, \dots, a_{mj})^\top \in \mathbb{K}^m$ der Einheitsvektoren bestimmt. Schreibt man diese spaltenweise nebeneinander, so entsteht ein rechteckiges Zahlenschema.

Definition 2.2 Unter einer reellen oder komplexen **(m, n)-Matrix** A verstehen wir ein rechteckiges Schema von Zahlen $a_{ij} \in \mathbb{K}$, $i = 1, \dots, m$, $j = 1, \dots, n$, wobei wieder $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$ ist. Wir schreiben:

$$A = \begin{pmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & & & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{pmatrix} = (a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}}.$$

Die Zahlen a_{ij} heißen **Elemente** oder **Koeffizienten** der Matrix. Wir schreiben $A \in \mathbb{K}^{m \times n}$ für (m, n) -Matrizen mit Koeffizienten in $\mathbb{K}$. Zwei Matrizen heißen gleich, wenn die Größen n und m und alle ihre Elemente übereinstimmen.

Die Überlegung vor dieser Definition zeigt, dass jede lineare Abbildung genau einer Matrix entspricht, deren Spalten die Bilder der Standardbasisvektoren sind. Man nennt diese Matrix die **Abbildungsmatrix** oder **Darstellungsmatrix** der linearen Abbildung. Ist $\mathcal{A} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ linear mit Darstellungsmatrix $A = (a_{ij}) \in \mathbb{K}^{m \times n}$, so gilt für alle $x \in (x_1, \dots, x_n)^\top \in \mathbb{K}^n$ mit $y = \mathcal{A}(x)$:

$$y = \begin{pmatrix} y_1 \\ \vdots \\ y_m \end{pmatrix} = \sum_{j=1}^n \begin{pmatrix} a_{1j} \\ \vdots \\ a_{mj} \end{pmatrix} x_j, \quad \text{und somit} \quad y_i = \sum_{j=1}^n a_{ij} x_j \quad i = 1, \dots, m.$$

Spezielle Matrizen liefern auch spezielle lineare Abbildungen.

Spezielle Beispiele:

(a) $O = (0)_{\substack{i=1,\dots,m \\ j=1,\dots,n}}$ heißt **Nullmatrix**, und entspricht der linearen Abbildung $\mathcal{A}(x) = 0$ für alle $x \in \mathbb{K}^n$

(b) Die Einheitsmatrix

$$I_n = \begin{pmatrix} 1 & 0 & \cdots & 0 & 0 \\ 0 & 1 & & & \vdots \\ 0 & 0 & \ddots & & \vdots \\ \vdots & \vdots & & 1 & 0 \\ 0 & 0 & \cdots & 0 & 1 \end{pmatrix} = (\delta_{ij})_{\substack{i=1,\dots,n \\ j=1,\dots,n}} \in \mathbb{R}^{n \times n}$$

ist der Identitätsabbildung $\mathcal{A}(x) = x$ zuzuordnen. Sie ist nur für $m = n$ definiert. Das verwendete Kroneckersymbol ist δ_{ij} erklärt durch

$$\delta_{ij} := \begin{cases} 1, & \text{falls } i = j, \\ 0, & \text{falls } i \neq j. \end{cases}$$

(c) Die $(m, 1)$ -Matrix $A = \begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix}$ können wir offensichtlich mit dem Spaltenvektor $a = \begin{pmatrix} a_{11} \\ \vdots \\ a_{m1} \end{pmatrix}$

identifizieren, d.h. es ist $\mathbb{K}^{m \times 1} = \mathbb{K}^m$. Die zugehörige lineare Abbildung $\mathcal{A} : \mathbb{K} \rightarrow \mathbb{K}^m$ ist einfach die Streckung des Vektors a um das Argument $x \in \mathbb{K}$, d.h. $\mathcal{A}(x) = x a$.

(d) Ein Zeilenvektor, d.h. eine $(1, n)$ -Matrix $(a_{11}, \dots, a_{1n})$ ist der linearen Abbildung

$$\mathcal{A}(x) = \sum_{j=1}^n a_{1j} x_j$$

zugeordnet, also im reellen Fall dem Skalarprodukt $\mathcal{A}(x) = b \cdot x$ mit dem Vektor $b = (a_{11}, \dots, a_{1n})^\top \in \mathbb{R}^n$.

Wir wollen nun Rechenoperationen wie Addition oder Multiplikation für Matrizen definieren. Da Matrizen so eng mit linearen Abbildungen verknüpft sind, sollen solche Operationen zu den entsprechenden Operationen für lineare Abbildungen kompatibel sein. Lineare Abbildungen können wie skalarwertige Funktionen addiert oder mit einer Zahl multipliziert werden, um neue lineare Abbildungen zu erhalten. Sind $\mathcal{A}, \mathcal{B} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ linear und $\lambda \in \mathbb{K}$, so sind

$$(\mathcal{A} + \mathcal{B})(x) = \mathcal{A}(x) + \mathcal{B}(x), \quad (\lambda \mathcal{A})(x) = \lambda \mathcal{A}(x), \quad x \in \mathbb{K}^n,$$

ebenfalls linear.

Wir führen entsprechende Operationen direkt für die Darstellungsmatrizen ein. Sind $A, B \in \mathbb{K}^{m \times n}$, so gehören hierzu lineare Abbildungen $\mathcal{A}, \mathcal{B} : \mathbb{K}^n \rightarrow \mathbb{K}^m$. Man definiert die **Addition** von A und B nun so, dass $A + B$ die Darstellungsmatrix von $\mathcal{A} + \mathcal{B}$ ist. Analog macht man es für die **skalare Multiplikation**. Es ist leicht nachzurechnen, dass dies auf eine komponentenweise Addition bzw. Multiplikation herausläuft: Sind $A = (a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}}$ und $B = (b_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}}$ $\in \mathbb{K}^{m \times n}$, dann setzen wir

$$A + B := (a_{ij} + b_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}} \quad \text{und} \quad \lambda A = (\lambda a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}}.$$

Eine weitere wichtige Operation ist die **Verkettung (Komposition)** von Abbildungen, d.h. wir verbinden $\mathcal{A} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ und $\mathcal{B} : \mathbb{K}^p \rightarrow \mathbb{K}^n$ zu einer wiederum linearen(!) Abbildung $\mathcal{C} : \mathbb{K}^p \rightarrow \mathbb{K}^m$, die durch $\mathcal{C}(x) = \mathcal{A}(\mathcal{B}(x))$ definiert ist. Für die Komposition von Abbildungen nutzen wir weiterhin die Notation $\mathcal{C} = \mathcal{A} \circ \mathcal{B}$. Die Abbildungsmatrix von $\mathcal{C}$ ergibt sich nun aus den Matrizen A zu $\mathcal{A}$ und B zu $\mathcal{B}$ durch das **Matrixprodukt**.

Definition 2.3 (*Matrixprodukt*)

Seien $A = (a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,p}} \in \mathbb{K}^{m \times p}$, $B = (b_{ij})_{\substack{i=1,\dots,p \\ j=1,\dots,n}} \in \mathbb{K}^{p \times n}$ zwei Matrizen, wobei die Spaltenzahl von A mit der Zeilenzahl von B übereinstimmt. Dann heißt

$$AB = (c_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}} \in \mathbb{K}^{m \times n},$$

definiert durch

$$c_{ij} := \sum_{k=1}^p a_{ik} b_{kj} \quad i = 1, \dots, m, \quad j = 1, \dots, n,$$

das **Matrixprodukt** oder **Matrizenprodukt** von A und B .

Achtung: Die Größen der zu multiplizierenden Matrizen müssen zueinander passen.

Wir begründen kurz, dass das Matrixprodukt tatsächlich die oben angesprochene Eigenschaft besitzt, kompatibel zur Verkettung linearer Abbildungen zu sein. Wir wählen $x \in \mathbb{K}^n$ und setzen $y = \mathcal{B}(x) \in \mathbb{K}^p$ und $z = \mathcal{C}(x) = \mathcal{A}(y) \in \mathbb{K}^m$. Dann ist $y = Bx$, $z = Ay$ und somit

$$y_k = \sum_{j=1}^n b_{kj} x_j, \quad k = 1, \dots, p, \quad z_i = \sum_{k=1}^p a_{ik} y_k, \quad i = 1, \dots, m.$$

Es folgt

$$z_i = \sum_{k=1}^p a_{ik} \sum_{j=1}^n b_{kj} x_j = \sum_{j=1}^n \left(\sum_{k=1}^p a_{ik} b_{kj} \right) x_j = \sum_{j=1}^n c_{ij} x_j .$$

Die Darstellungsmatrix C von $\mathcal{C} = \mathcal{A} \circ \mathcal{B}$ ist also tatsächlich das Matrixprodukt AB .

Beispiele:

$$\underbrace{\begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}}_{\in \mathbb{R}^{3 \times 1}} \underbrace{\begin{pmatrix} -1 & 2 \end{pmatrix}}_{\in \mathbb{R}^{1 \times 2}} = \underbrace{\begin{pmatrix} -1 & 2 \\ -2 & 4 \\ -3 & 6 \end{pmatrix}}_{\in \mathbb{R}^{3 \times 2}}$$

$$\underbrace{\begin{pmatrix} 1 & 2 \\ 0 & 3 \\ 4 & 1 \end{pmatrix}}_{\in \mathbb{R}^{3 \times 2}} \underbrace{\begin{pmatrix} 1 & 3 & 4 & 1 \\ 2 & 0 & 1 & 2 \end{pmatrix}}_{\in \mathbb{R}^{2 \times 4}} = \underbrace{\begin{pmatrix} 5 & 3 & 6 & 5 \\ 6 & 0 & 3 & 6 \\ 6 & 12 & 17 & 6 \end{pmatrix}}_{\in \mathbb{R}^{3 \times 4}}$$

Die uns schon bekannten Spaltenvektoren, also die Elemente von $\mathbb{K}^n$, finden wir unter den Matrizen als Elemente des $\mathbb{K}^{n \times 1}$ wieder. Diese beiden Mengen werden identifiziert. Wichtig ist hierbei, dass für $x \in \mathbb{K}^n$ das Bild $\mathcal{A}(x)$ unter einer linearen Abbildung $\mathcal{A} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ durch das Matrixprodukt der Abbildungsmatrix $A \in \mathbb{K}^{m \times n}$ von $\mathcal{A}$ mit x gegeben ist:

$$\mathcal{A}(x) = Ax, \quad x \in \mathbb{K}^n .$$

Umgekehrt ist für jede Matrix $A \in \mathbb{K}^{m \times n}$ eine lineare Abbildung $\mathcal{A} : \mathbb{K}^n \rightarrow \mathbb{K}^m$ durch $\mathcal{A}(x) = Ax$, $x \in \mathbb{K}^n$ eindeutig festgelegt.

Auch ein lineares Gleichungssystem mit m Gleichungen und n Unbekannten, ausgeschrieben:

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

bzw.

$$\sum_{j=1}^n a_{ij} x_j = b_i, \quad i = 1, \dots, m ,$$

kann jetzt als $Ax = b$ mit $A = (a_{ij}) \in \mathbb{K}^{m \times n}$ und $b = (b_i) \in \mathbb{K}^m$ geschrieben werden. Die Lösung eines LGS zu finden, bedeutet also nichts weiter, als dass wir Urbilder zum Vektor b (etwa Nullstellen, $b = 0$) unter der zugehörigen linearen Abbildung $\mathcal{A}$ suchen.

Für das Matrixprodukt gelten folgende Rechenregeln, die wir durch Nachrechnen leicht zeigen können:

Satz 2.4 (Rechenregeln)

Für Matrizen A, B, C der „richtigen“ Dimensionen und $\lambda \in \mathbb{K}$ gilt:

- | | | |
|-----|---|---|
| (a) | $\lambda(AB) = (\lambda A)B = A(\lambda B)$ | 1. Assoziativgesetz, |
| (b) | $A(BC) = (AB)C$ | 2. Assoziativgesetz, |
| (c) | $A(B + C) = AB + AC$ | 1. Distributivgesetz, |
| (d) | $(A + B)C = AC + BC$ | 2. Distributivgesetz, |
| (e) | $A0 = 0, \quad 0A = 0$ | (0 bezeichne verschiedene Nullmatrizen!), |
| (f) | $I_m A = A, \quad A I_n = A$ | |

Achtung: Im Allgemeinen ist selbst für quadratische Matrizen $AB \neq BA$. Die Matrizenmultiplikation (und somit auch die Komposition von linearen Abbildungen) ist also nicht kommutativ. Darüber hinaus folgt aus $AB = O$ **nicht**, dass $A = O$ oder $B = O$ ist!

Beispiel: Mit

$$A = \begin{pmatrix} 1 & 0 \\ 1 & 0 \end{pmatrix} \text{ und } B = \begin{pmatrix} 0 & 0 \\ 1 & 1 \end{pmatrix} \text{ ist } AB = O \text{ und } BA = \begin{pmatrix} 0 & 0 \\ 2 & 0 \end{pmatrix}.$$

Also ist das Produkt nicht kommutativ und besitzt „Nullteiler“, d.h. das Produkt kann O sein, ohne dass einer der beiden Faktoren O ist.

Für Vektoren haben wir schon $(\dots)^\top$ kennengelernt, um aus einem Zeilenvektor einen Spaltenvektor zu machen und umgekehrt. Diese Notation führen wir nun allgemein für Matrizen ein.

Definition 2.5 (*Transponierte und adjungierte Matrix*)

Sei $A = (a_{ij})_{\substack{i=1,\dots,m \\ j=1,\dots,n}} \in \mathbb{K}^{m \times n}$.

(a) Die Matrix $A^\top := (\tilde{a}_{ij})_{\substack{i=1,\dots,n \\ j=1,\dots,m}} \in \mathbb{K}^{n \times m}$ mit $\tilde{a}_{ij} = a_{ji}$ heißt **transponierte Matrix** zu **A**.
(Die Zeilen und Spalten werden vertauscht)

(b) Die Matrix $A^* := (\tilde{a}_{ij})_{\substack{i=1,\dots,n \\ j=1,\dots,m}} \in \mathbb{K}^{n \times m}$ mit $\tilde{a}_{ij} = \overline{a_{ji}}$ heißt **adjungierte Matrix** zu **A**. Hier bezeichnen wir wieder mit $\overline{a_{ji}}$ die zu $a_{ji} \in \mathbb{C}$ komplex konjugierte Zahl. (Die Koeffizienten von A^* sind die komplex konjugierten Einträge von $A^\top$).

Für reelle Matrizen stimmen beide Begriffe überein.

Beispiel: Sei

$$A = \begin{pmatrix} 1+i & 3 \\ 2 & 1-3i \\ 0 & 4 \end{pmatrix} \in \mathbb{C}^{3 \times 2}.$$

Dann ist

$$A^\top = \begin{pmatrix} 1+i & 2 & 0 \\ 3 & 1-3i & 4 \end{pmatrix} \in \mathbb{C}^{2 \times 3} \quad \text{und} \quad A^* = \begin{pmatrix} 1-i & 2 & 0 \\ 3 & 1+3i & 4 \end{pmatrix} \in \mathbb{C}^{2 \times 3}.$$

Für das Rechnen mit der transponierten oder adjungierten Matrix gelten die folgenden **Regeln**:

$$(A^\top)^\top = A, \quad (A+B)^\top = A^\top + B^\top \quad (\lambda A)^\top = \lambda A^\top \quad \text{und}$$

$$(A^*)^* = A, \quad (A+B)^* = A^* + B^* \quad (\lambda A)^* = \overline{\lambda} A^*$$

für jedes $\lambda \in \mathbb{K}$.

Interessant ist nun das Zusammenspiel dieser Operationen mit den verschiedenen Formen der Multiplikation von Matrizen und Vektoren, die wir kennengelernt haben.

Satz 2.6 Es seien $A \in \mathbb{K}^{m \times n}$, $B \in \mathbb{K}^{n \times p}$ Matrizen.

(a) Für die Transposition bzw. Adjunktion des Matrixprodukts gilt

$$(AB)^\top = B^\top A^\top, \quad (AB)^* = B^* A^*.$$

(b) Sind $x, y \in \mathbb{K}^n$, so gilt

$$\langle x, y \rangle = x^\top \bar{y} = y^* x$$

(c) Sind $x \in \mathbb{K}^n$, $y \in \mathbb{K}^m$, so gilt

$$\langle Ax, y \rangle = \langle x, A^* y \rangle.$$

Beweis: (a) und (b) können leicht nachgerechnet werden. Dann folgt (c) aus (b), denn

$$\langle Ax, y \rangle = y^*(Ax) = (y^*A)x = (A^*y)^*x = \langle x, A^*y \rangle.$$

□

Ein weiterer Spezialfall ist das **dyadische Produkt**, $ba^\top$, von $a \in \mathbb{K}^n$ und $b \in \mathbb{K}^m$. Ausgeschrieben ergibt sich

$$ba^\top = \begin{pmatrix} b_1 \\ \vdots \\ b_m \end{pmatrix} (a_1, \dots, a_n) = \begin{pmatrix} b_1 a_1 & b_1 a_2 & \cdots & b_1 a_n \\ \vdots & & & \vdots \\ b_m a_1 & \cdots & \cdots & b_m a_n \end{pmatrix} \in \mathbb{K}^{m \times n}.$$

Die Definition impliziert sofort folgende Rechenregel mit dyadischen Produkten:

$$A(ba^\top) = (Ab)a^\top \quad \text{für } A \in \mathbb{K}^{m \times n}, b \in \mathbb{K}^n, a \in \mathbb{K}^p,$$

$$\underbrace{(ab^\top)}_{\in \mathbb{K}^{n \times m}} \underbrace{(cd^\top)}_{\in \mathbb{K}^{m \times p}} = \underbrace{(b^\top c)}_{\in \mathbb{K}} \underbrace{(ad^\top)}_{\in \mathbb{K}^{n \times p}} \quad \text{für } a \in \mathbb{K}^n, b, c \in \mathbb{K}^m, d \in \mathbb{K}^p$$

und analog für adjungierte Elemente.

Wir betrachten zum Abschluss dieses Abschnitts noch einige wichtige lineare Abbildungen der Ebene und des Raumes und ihre Abbildungsmatrizen. Durch diese Beispiele werden wir auf eine spezielle Klasse von linearen Abbildungen geführt.

Beispiele 2.7 (a) Mit $\alpha \in [0, 2\pi]$ sei

$$R = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \in \mathbb{R}^{2 \times 2}.$$

Die Matrix R angewandt auf die Koordinateneinheitsvektoren $e^{(1)}$ und $e^{(2)}$ ergibt $Re^{(1)} = \begin{pmatrix} \cos \alpha \\ \sin \alpha \end{pmatrix}$ und $Re^{(2)} = \begin{pmatrix} -\sin \alpha \\ \cos \alpha \end{pmatrix}$. Diese beiden Vektoren werden also unter der Matrix R um den Winkel α entgegen dem Uhrzeigersinn gedreht. Da die zugehörige Abbildung $\mathcal{R}$ linear ist, und $\{e^{(1)}, e^{(2)}\}$ eine Basis des $\mathbb{R}^2$ ist, beschreibt Rx für jeden Vektor $x \in \mathbb{R}^2$ seine **Drehung** (oder Rotation) um den Ursprung und den Winkel α .

(b) Sei wieder $\alpha \in [0, 2\pi]$. Die $(3, 3)$ -Matrizen

$$R_3 = \begin{pmatrix} \cos \alpha & -\sin \alpha & 0 \\ \sin \alpha & \cos \alpha & 0 \\ 0 & 0 & 1 \end{pmatrix}, \quad R_2 = \begin{pmatrix} \cos \alpha & 0 & \sin \alpha \\ 0 & 1 & 0 \\ -\sin \alpha & 0 & \cos \alpha \end{pmatrix},$$

$$R_1 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \alpha & -\sin \alpha \\ 0 & \sin \alpha & \cos \alpha \end{pmatrix} \in \mathbb{R}^{3 \times 3}$$

beschreiben Rotationen im $\mathbb{R}^3$ gegen den Uhrzeigersinn um die x_3 -Achse, bzw. x_2 -Achse oder x_1 -Achse.

(c) Es sei $n \in \mathbb{N}$, $n \geq 2$, beliebig und $a \in \mathbb{R}^n$ mit $\|a\|^2 = a \cdot a = a^\top a = 1$ ein Einheitsvektor. Setze

$$P := I_n - a a^\top \in \mathbb{R}^{n \times n},$$

wobei I_n die n -dimensionale Einheitsmatrix ist. Es ist Px die **orthogonale Projektion** auf den Unterraum $E := \{x \in \mathbb{R}^n : a^\top x = 0\}$, d.h.

$$Px \in E \quad \text{und} \quad (x - Px) \cdot y = 0 \text{ für alle } y \in E.$$

Es ist $Px \in E$ wegen $a^\top Px = a^\top x - (a^\top a) a^\top x = 0$ und $x - Px = a(a^\top x) \in \text{span}\{a\}$.

Für $n = 2$ ist Px also die orthogonale Projektion auf die Gerade $\{x \in \mathbb{R}^2 : a^\top x = 0\}$, für $n = 3$ auf die Ebene $\{x \in \mathbb{R}^3 : a^\top x = 0\}$.

Wir betrachten noch den Fall des affinen (d.h. verschobenen) Unterraums. Sei wieder $a \in \mathbb{R}^n$ ein Einheitsvektor und $\rho \in \mathbb{R}$. Für $x \in \mathbb{R}^n$ ist die orthogonale Projektion p_x auf

$$E := \{x \in \mathbb{R}^n : a^\top x = \rho\} = \hat{x} + \{x \in \mathbb{R}^n : a^\top x = 0\}$$

gegeben durch

$$p_x = Px + (a^\top \hat{x}) a = x + (a^\top (\hat{x} - x)) a.$$

(d) Sei $n \in \mathbb{N}$, $n \geq 2$, beliebig und $a \in \mathbb{R}^n$ ein Einheitsvektor, d.h. $\|a\|_2 = 1$. Setze

$$S := I_n - 2 a a^\top \in \mathbb{R}^{n \times n}.$$

Dann beschreibt Sx die **Spiegelung** von x an dem Unterraum $E = \{x \in \mathbb{R}^n : a^\top x = 0\}$. Spiegelung bedeutet folgendes: Ist $P = I_n - a a^\top$ der orthogonale Projektor von Teil (c), so liegen x , Px und Sx auf derselben Geraden mit Richtung a und $\|Sx - Px\| = \|x - Px\|$. Man veranschauliche sich dies für $n = 2$ und $n = 3$! □

Rotationen und Spiegelungen sind Beispiele für orthogonale Matrizen.

Definition 2.8 (*orthogonale, unitäre Matrix*)

- (a) Eine reelle, quadratische Matrix $A \in \mathbb{R}^{n \times n}$ heißt **orthogonal**, falls $A^\top A = A A^\top = I_n$.
- (b) Eine komplexe, quadratische Matrix $A \in \mathbb{C}^{n \times n}$ heißt **unitär**, falls $A^* A = A A^* = I_n$.

Interpretiert man Matrizen als Abbildungen, so lassen orthogonale Matrizen Längen und Winkel **invariant**, d.h.

$$(Ax) \cdot (Ay) = (Ax)^\top (Ay) = x^\top \underbrace{A^\top A}_{=I_n} y = x^\top y = x \cdot y.$$

Man rechne selbst nach, dass die Rotationen R_j aus Teil (b) des letzten Beispiels orthogonale Matrizen sind. Dies gilt auch für die Spiegelung S von Teil (d), denn

$$\begin{aligned} S^\top S &= (I_n - 2aa^\top)^\top (I_n - 2aa^\top) = (I_n - 2aa^\top)(I_n - 2aa^\top) \\ &= I_n - 4aa^\top + 4a \underbrace{a^\top a}_{=1} = I_n. \end{aligned}$$

2.2 Theorie der linearen Gleichungssysteme

Wir haben gesehen, dass sich lineare Abbildungen durch entsprechende Matrizen beschreiben lassen. Das Lösen eines linearen Gleichungssystems von der Form $Ax = b$ kann somit interpretiert werden, als die Suche nach der Lösung einer Gleichung $\mathcal{A}(x) = b$, also einem Urbild x zu b unter der zugehörigen linearen Abbildung $\mathcal{A}$.

In diesem Abschnitt fragen wir uns zunächst, unter welchen Bedingungen es eine solche Lösung gibt. Sei also ein lineares Gleichungssystem $Ax = b$ gegeben. Wenn wir den j -ten Spaltenvektor der Matrix A mit $a_{*j} := (a_{1j}, \dots, a_{mj})^\top \in \mathbb{K}^m$ bezeichnen, können wir das lineare Gleichungssystem auch durch

$$\sum_{j=1}^n x_j a_{*j} = b$$

beschreiben. Also ist das Gleichungssystem lösbar genau dann, wenn b Element des durch die Spalten von A aufgespannten Unterraums ist, d.h.

$$b \in B = \text{span} \{a_{*j} : j = 1, \dots, n\} = \{y \in \mathbb{K}^m : \text{es gibt } x \in \mathbb{K}^n \text{ sodass } y = Ax\}.$$

Diese Menge B nennen wir das **Bild** der Matrix A . Weiter setzen wir

$$\mathcal{L} := \{x \in \mathbb{K}^n : Ax = b\} \quad \text{und} \quad \mathcal{L}_0 := \{x \in \mathbb{K}^n : Ax = 0\}.$$

$\mathcal{L}_0$ wird der **Kern** von A genannt. In den Beispielen 1.9 und 1.18, aber auch in Satz 1.25 sind uns schon solche Mengen begegnet. Der Kern einer Matrix spielt eine entscheidende Rolle in der allgemeinen Lösungstheorie. Dies besagt der folgende Satz.

Satz 2.9 *Der Kern $\mathcal{L}_0$ von A ist ein linearer Unterraum des $\mathbb{K}^n$.*

Die Lösungsmenge $\mathcal{L}$ ist entweder leer, oder $\mathcal{L}$ ist ein affiner Unterraum des $\mathbb{K}^n$ und

$$\mathcal{L} = \hat{x} + \mathcal{L}_0 = \{x \in \mathbb{K}^n : x = \hat{x} + x_0, x_0 \in \mathcal{L}_0\}$$

für jedes $\hat{x} \in \mathcal{L}$.

Dies bedeutet, dass sich alle Lösungen des inhomogenen Gleichungssystems als Summe aus irgend-einer **partikulären** Lösung $\hat{x}$ des inhomogenen Gleichungssystems und Lösungen des homogenen Gleichungssystems schreiben lassen.

Beweis: Sei $\mathcal{L} \neq \emptyset$. Dann gibt es $\hat{x} \in \mathcal{L}$. Ist nun $y \in \mathcal{L}$ ein weiteres Element in der Lösungsmenge $\mathcal{L}$, so folgt, dass $A(y - \hat{x}) = b - b = 0$. Also ist $y - \hat{x} \in \mathcal{L}_0$ bzw. $y = \hat{x} + (y - \hat{x}) \in \hat{x} + \mathcal{L}_0$. Andererseits, wenn $y \in \hat{x} + \mathcal{L}_0$ ist, also $v \in \mathcal{L}_0$ existiert mit $y = \hat{x} + v$, dann gilt $Ay = A(\hat{x} + v) = A\hat{x} + Av = b$. Damit ist $y \in \mathcal{L}$. Insgesamt erhalten wir, dass die beiden Mengen gleich sind. $\square$

Mit diesem Satz können wir nun beim Lösen von linearen Gleichungssystemen $Ax = b$ genau drei Fälle unterscheiden

1. **Fall:** Es gibt überhaupt **keine Lösung** $\hat{x}$ zu $Ax = b$.
2. **Fall:** Es existiert $\hat{x} \in \mathbb{K}^n$ mit $A\hat{x} = b$ und $\mathcal{L}_0 = \{0\}$. Dann gibt es **genau eine Lösung** des Gleichungssystems nämlich $\hat{x}$.
3. **Fall:** Es existiert $\hat{x} \in \mathbb{K}^n$ mit $A\hat{x} = b$ und $\mathcal{L}_0 \neq \{0\}$. Dann besitzt das Gleichungssystem **unendlich viele Lösung** nämlich den affinen Unterraum $\mathcal{L} = \hat{x} + \mathcal{L}_0$.

Welcher der drei Fälle zutrifft, lässt sich zumindest theoretisch stets mit dem Gaußschen Eliminationsverfahren ermitteln. Mit Hilfe des Verfahrens lassen sich wichtige allgemeine Aussagen zeigen. Es sei dazu an den Begriff des *Rangs eines linearen Gleichungssystems* aus dem ersten Kapitel erinnert: Dieser bezeichnet die Anzahl der Pivotelemente, die im Gauß'schen Eliminationsverfahren gewählt werden können (siehe Seite 9).

Satz 2.10 Betrachte ein lineares Gleichungssystem $Ax = b$ mit $A \in \mathbb{K}^{m \times n}$, $b \in \mathbb{K}^m$.

- (a) Ist $m = n$, gibt es also ebenso viele Gleichungen wie Unbekannte, so ist das inhomogene Gleichungssystem für jede rechte Seite $b \in \mathbb{K}^m$ genau dann eindeutig lösbar, wenn das homogene Gleichungssystem nur die Lösung $x = 0$ besitzt.
- (b) Ist r der Rang des linearen Gleichungssystems, so gilt $r + \dim \mathcal{L}_0 = n$.
- (c) Ist $r < n$, so gibt es **nichttriviale** (d.h. von $0 \in \mathbb{K}^n$ verschiedene) Lösungen des homogenen Gleichungssystems, d.h. $\mathcal{L}_0 \neq \{0\}$ und $\dim \mathcal{L}_0 \geq 1$. Da $r \leq \min m, n$, ist dies insbesondere der Fall, wenn weniger Gleichungen als Freiheitsgrade gegeben sind. Hat das inhomogene System eine Lösung, so gibt es unendlich viele Lösungen.

Beweis: Wir skizzieren den Beweis zu Teil (a).

Wenn $Ax = b$ für jedes $b \in \mathbb{K}^m$ genau eine Lösung besitzt, so trifft dies auch für $b = 0$ zu. Da $x = 0$ Lösung des homogenen Systems ist, folgt, dass $x = 0$ die einzige Lösung ist, also $\mathcal{L}_0 = \{0\}$. Dies beweist eine Richtung der Aussage.

Andererseits, wenn $\mathcal{L}_0 = \{0\}$ ist, so kann in einem Endtableau, dass wir durch das oben beschriebene Gaußsche Eliminationsverfahren gewinnen, keine Zeile verschwinden. Da sonst wegen $m = n$ in einer Spalte kein Pivotelement auftritt und wir mit beliebiger Wahl der zugehörigen Variable stets das homogene System lösen könnten im Widerspruch zu $\mathcal{L}_0 = \{0\}$. Also erreichen wir mit dem Algorithmus nach eventuellem Umbenennen der Variablen, Umsortieren der Zeilen und Dividieren der Zeilen durch die Pivotelemente zu beliebiger rechter Seite $b \in \mathbb{K}^m$ ein Endtableau der Form

$$\begin{array}{cccc|c} 1 & * & \dots & * & \tilde{b}_1 \\ 0 & 1 & \ddots & \vdots & \vdots \\ \vdots & \ddots & \ddots & * & \\ 0 & \dots & 0 & 1 & \tilde{b}_n, \end{array}$$

von dem wir die Lösung des Systems durch Rücksubstitution ablesen können. $\square$

Beispiel 2.11

Betrachte das Gleichungssystem

$$\begin{aligned} x_1 &+ 2x_3 = 1 \\ 4x_1 + 2x_2 + 3x_3 &= 0 \\ 2x_1 + 2x_2 + \alpha x_3 &= \gamma \end{aligned}$$

mit $\alpha, \gamma \in \mathbb{R}$. Mit dem Gaußschen Eliminationsverfahren erhalten wir

$$\begin{array}{ccc|c} 1 & 0 & 2 & 1 \\ 4 & 2 & 3 & 0 \\ 2 & 2 & \alpha & \gamma \end{array} \longrightarrow \begin{array}{ccc|c} \boxed{1} & 0 & 2 & 1 \\ 4 & 2 & 3 & 0 \\ -2 & 0 & \alpha - 3 & \gamma \end{array} \longrightarrow \begin{array}{ccc|c} 1 & 0 & 2 & 1 \\ 4 & 2 & 3 & 0 \\ 0 & 0 & \boxed{\alpha + 1} & \gamma + 2 \end{array}$$

Schreiben wir das reduzierte lineare Gleichungssystem ausführlich hin,

$$\begin{aligned} x_1 &+ 2x_3 &= 1 \\ 4x_1 + 2x_2 + 3x_3 &= 0 \\ (\alpha + 1)x_3 &= \gamma + 2, \end{aligned}$$

so sehen wir sofort die drei möglichen Fälle. Ist $\alpha + 1 \neq 0$, so ist das Gleichungssystem eindeutig lösbar (für beliebige Werte $\gamma \in \mathbb{R}$), also der zweite Fall. Wenn $\alpha = -1$ und $\gamma \neq -2$ ist, gibt es keine Lösung. Das ist oben als erster Fall beschrieben. Wenn aber $\alpha = -1$ und $\gamma = -2$ gilt, so gibt es unendlich viele Lösungen (3. Fall) und wir können die Lösungsmenge eine Gerade, durch Rücksubstitution mit $x_3 = t$ ablesen:

$$\mathcal{L} = \left\{ \underbrace{\begin{pmatrix} 1 \\ -2 \\ 0 \end{pmatrix}}_{\hat{x}} + t \underbrace{\begin{pmatrix} -2 \\ \frac{5}{2} \\ 1 \end{pmatrix}}_{\in \mathcal{L}_0}, \quad t \in \mathbb{R} \right\}$$

Wir interpretieren nun noch diese Theorie der linearen Gleichungssysteme und linearen Abbildungen im Fall quadratischer Matrizen. Dazu definieren wir die zu einer Matrix A *inverse Matrix*, also die Matrix, die der Umkehrabbildung zu $\mathcal{A}$ zugeordnet werden kann, wenn die lineare Abbildung $\mathcal{A}$ umkehrbar ist.

Definition 2.12 Eine quadratische Matrix $A \in \mathbb{K}^{n \times n}$ heißt **regulär, nicht singulär oder invertierbar**, wenn es eine Matrix $X \in \mathbb{K}^{n \times n}$ gibt mit $A X = X A = I_n$. Falls es so eine Matrix X gibt, so ist X eindeutig. Sie heißt die zu A **inverse Matrix** und wird mit A^{-1} bezeichnet. Es gilt also

$$A A^{-1} = A^{-1} A = I_n.$$

Andernfalls heißt die Matrix A **singulär**.

Die Eindeutigkeit ergibt sich wie folgt: Gäbe es zwei Matrizen X und Y mit $A X = X A = I_n$ und $A Y = Y A = I_n$, so wäre $Y = Y I_n = Y A X = I_n X = X$.

Bemerkung: Bei quadratischen Matrizen lässt sich zeigen, dass eine Matrix genau dann invertierbar ist, wenn es eine Matrix X gibt mit $X A = I_n$, also ohne die in unserer Definition geforderte zweite Identität $A X = I_n$. Wir verzichten auf einen Beweis.

Die Lösungstheorie der linearen Gleichungssysteme liefert eine ganze Reihe von äquivalenten Bedingungen, um die Regularität einer Matrix zu beschreiben. Die Wichtigsten führen wir hier auf:

Satz 2.13 Sei $A \in \mathbb{K}^{n \times n}$. Dann sind äquivalent:

- (a) $A \in \mathbb{K}^{n \times n}$ ist regulär.
- (b) $A^\top$ ist regulär.
- (c) A^* ist regulär.
- (d) Das Gleichungssystem $Ax = b$ hat für jedes $b \in \mathbb{K}^n$ genau eine Lösung $x \in \mathbb{K}^n$.
- (e) Das homogene Gleichungssystem $Ax = 0$ besitzt nur die Lösung $x = 0$.
- (f) Der Rang des Gleichungssystems bzw. der Matrix ist n .

Beweis: (a) $\Leftrightarrow$ (b): Ist A regulär, so ist

$$(A^{-1})^\top A^\top = (AA^{-1})^\top = I_n \quad \text{und} \quad A^\top (A^{-1})^\top = (A^{-1}A)^\top = I_n.$$

Genauso zeigen wir umgekehrt aus der Regularität von $A^\top$ die von A .

(a) $\Leftrightarrow$ (c): Dies lässt sich genauso zeigen.

(a) $\Rightarrow$ (d): Ist A regulär, so ist $x = A^{-1}b$ Lösung des Gleichungssystems, denn $Ax = AA^{-1}b = I_nb = b$. Außerdem ergibt sich für $x, y \in \mathbb{K}^n$ mit $Ax = Ay = b$, dass $A(x-y) = Ax - Ay = 0$ ist und mit der Inversen folgt $x - y = A^{-1}0 = 0$, d.h. es gibt genau eine Lösung des Gleichungssystems.

(d) $\Rightarrow$ (a): Besitzt das Gleichungssystem für jede rechte Seite eine eindeutige Lösung, so können wir für jedes $j = 1, \dots, n$ Spaltenvektoren $x^{(j)} \in \mathbb{K}^n$ finden mit $Ax^{(j)} = e^{(j)}$. Hier ist wieder $e^{(j)}$ der j -te Koordinateneinheitsvektor im $\mathbb{K}^n$. Fassen wir die Vektoren $x^{(j)}$ zu der Matrix $X \in \mathbb{K}^{n \times n}$ zusammen, so bedeutet diese Gleichung genau $AX = I_n$. Wir haben $XA = I_n$ zu zeigen. Zunächst zeigen wir, dass auch das Gleichungssystem $A^*x = z$ für jedes $z \in \mathbb{K}^n$ eine Lösung besitzt. Nach Satz 2.10(a) genügt es zu zeigen, dass das homogene System $A^*x = 0$ nur die Lösung $x = 0$ besitzt. Sei $x \in \mathbb{K}^n$ eine Lösung von $A^*x = 0$ und sei $u \in \mathbb{K}^n$ Lösung von $Au = x$. Dann ist

$$0 = (A^*x)^*u = x^*Au = x^*x = \sum_{j=1}^n |x_j|^2.$$

Hieraus folgt $x = 0$. Daher hat $A^*x = z$ für jedes $z \in \mathbb{K}^n$ eine Lösung. Wir können wieder für jedes $j = 1, \dots, n$ Spaltenvektoren $y^{(j)} \in \mathbb{K}^n$ finden mit $A^*y^{(j)} = e^{(j)}$. Fassen wir die Vektoren $y^{(j)}$ zu der Matrix $Y \in \mathbb{K}^{n \times n}$ zusammen, so bedeutet diese Gleichung $A^*Y = I_n$. Damit ist $Y^*A = I_n$ und

$$Y^* = Y^*I_n = Y^*(AX) = (Y^*A)X = I_nX = X,$$

und daher auch $XA = I_n$.

(d) $\Leftrightarrow$ (e): siehe Satz 2.10

(e) $\Leftrightarrow$ (f): Dies folgt aus der Durchführung des Gaußschen Eliminationsverfahrens. □

Fazit:

Ist A quadratisch und regulär, so wird das Gleichungssystem $Ax = y$ durch $x = A^{-1}y$ gelöst.

Also ist die inverse Matrix A^{-1} die Abbildungsmatrix derjenigen linearen Abbildung, die die Umkehrabbildung zu $\mathcal{A}(x) = Ax$ ist.

Für das Rechnen mit inversen Matrizen gelten die folgenden **Regeln**:

$$(A^{-1})^{-1} = A, \quad (AB)^{-1} = B^{-1}A^{-1}, \quad (A^\top)^{-1} = (A^{-1})^\top,$$

aber $(A + B)^{-1} \neq A^{-1} + B^{-1}$.

Beispiel:

Für $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathbb{K}^{2 \times 2}$ ist $A^{-1} = \frac{1}{ad-bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$, falls $ad-bc \neq 0$. Dies lässt sich direkt nachrechnen.

Die inverse Matrix A^{-1} zu A erhalten wir, indem wir die Gleichungssysteme $AX = I_n$ lösen, d.h. die linearen Gleichungssysteme für die n rechten Seiten $e^{(j)} \in \mathbb{K}^n$, $j = 1, \dots, n$. Dazu nutzen wir wieder das Gauß-Verfahren mit einer Matrix auf der rechten Seite, in deren Spalten die rechten Seiten zusammengestellt sind. Beim letzten Tableau vertauschen wir die Zeilen so, dass links vom vertikalen Strich die Einheitsmatrix steht. Rechts vom Strich können wir so die inverse Matrix ablesen.

Beispiel 2.14

Bestimme die inverse Matrix zu $A = \begin{pmatrix} 1 & 2 & 1 \\ 0 & 2 & 1 \\ 2 & 6 & 2 \end{pmatrix}$. Wir haben das folgende Gleichungssystem zu lösen:

$$\begin{array}{ccc|ccc} \boxed{1} & 2 & 1 & 1 & 0 & 0 \\ 0 & 2 & 1 & 0 & 1 & 0 \\ 2 & 6 & 2 & 0 & 0 & 1 \end{array} \longrightarrow \begin{array}{ccc|ccc} 1 & 2 & 1 & 1 & 0 & 0 \\ 0 & 2 & \boxed{1} & 0 & 1 & 0 \\ 0 & 2 & 0 & -2 & 0 & 1 \end{array}$$

$$\longrightarrow \begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 2 & 1 & 0 & 1 & 0 \\ 0 & \boxed{2} & 0 & -2 & 0 & 1 \end{array} \longrightarrow \begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & 2 & 1 & -1 \\ 0 & 2 & 0 & -2 & 0 & 1 \end{array}$$

Dies beendet das eigentliche Gauß-Verfahren. Jetzt dividieren wir die letzte Gleichung noch durch 2 und sortieren die Zeilen (d.h. die Gleichungen) so, dass links die Einheitsmatrix steht. Dies liefert

$$\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & -1 & 0 \\ 0 & 1 & 0 & -1 & 0 & \frac{1}{2} \\ 0 & 0 & 1 & 2 & 1 & -1 \end{array} \quad \text{und daher ist} \quad A^{-1} = \begin{pmatrix} 1 & -1 & 0 \\ -1 & 0 & \frac{1}{2} \\ 2 & 1 & -1 \end{pmatrix}.$$

□

Einen Aspekt in Hinblick auf lineare Abbildungen haben wir bisher noch nicht angesprochen. Eine Abbildung $\mathcal{A} : \mathbb{R}^n \rightarrow \mathbb{R}^m$ ist unabhängig von der Koordinatendarstellung der Vektoren in $\mathbb{R}^n$. Aber die zugehörige Matrix hängt vom Koordinatensystem ab. Also sieht die Matrix zu einer linearen Abbildung bezüglich einer anderen Basis völlig anders aus.

Zunächst beobachten wir, dass Basiswechsel im $\mathbb{R}^n$ (bzw. $\mathbb{C}^n$) auch durch Matrizen gegeben sind. Betrachten wir zwei Basen zum $\mathbb{R}^n$, $B = \{b^{(1)}, \dots, b^{(n)}\}$ und $C = \{c^{(1)}, \dots, c^{(n)}\}$. Da B eine Basis ist, gibt es eindeutig bestimmte Zahlen $t_{ij} \in \mathbb{R}$ mit

$$c^{(j)} = \sum_{i=1}^n t_{ij} b^{(i)}.$$

Wir erhalten eine Matrix ${}_B T_C = (t_{ij})_{i,j=1,\dots,n} \in \mathbb{R}^{n \times n}$ mit den Einträgen $t_{ij} \in \mathbb{R}$. Ist nun $x \in \mathbb{R}^n$ ein Koordinatenvektor bezüglich der Basis C , d.h. $x = (x_1, x_2, \dots, x_n)^\top$ beschreibt den Vektor

$$v = x_1 c^{(1)} + x_2 c^{(2)} + \dots + x_n c^{(n)},$$

so erhalten wir die Koordinaten $y \in \mathbb{R}^n$ desselben Vektors v bezüglich der anderen Basis B durch die Matrixmultiplikation $y = {}_B T_C x$. Das bedeutet, es gilt

$$v = y_1 b^{(1)} + y_2 b^{(2)} + \dots + y_n b^{(n)}.$$

Die Matrix ${}_B T_C$ heißt **Basistransformationsmatrix**. Die Darstellungen sind eindeutig. Daher ist ${}_B T_C$ invertierbar und beschreibt den Koordinatenwechsel vollständig.

Betrachten wir nun eine lineare Abbildung $\mathcal{A} : \mathbb{R}^n \rightarrow \mathbb{R}^n$, die bezüglich der Basis C durch eine Matrix $A \in \mathbb{R}^{n \times n}$ gegeben ist. Dann gilt für einen Vektor v mit Koordinaten $x \in \mathbb{R}^n$ bezüglich der Basis C , dass das Bild $\mathcal{A}(v)$ die Koordinaten $Ax = b$ bezüglich der Basis C hat. Nun wollen wir dieselbe Abbildung bezüglich der Basis B angeben. Wir bezeichnen mit $y = {}_B T_C x$ und $\hat{b} = {}_B T_C b$ die Koordinaten von v beziehungsweise des Bilds $\mathcal{A}(v)$ bezüglich B . Dann gilt

$$\hat{b} = {}_B T_C b = {}_B T_C Ax = {}_B T_C A {}_B T_C^{-1} y.$$

Also ist durch das Matrizenprodukt ${}_B T_C A {}_B T_C^{-1} = \hat{A}$ die Matrix $\hat{A}$ zu $\mathcal{A}$ gegeben, wenn wir Koordinaten in Bezug auf die Basis B betrachten.

Definition 2.15 (Ähnliche Matrizen)

Zwei Matrizen $A, \hat{A} \in \mathbb{R}^{n \times n}$ (bzw. $\mathbb{C}^{n \times n}$) heißen zueinander **ähnlich**, wenn es eine invertierbare $n \times n$ Matrix T gibt mit

$$\hat{A} = T A T^{-1}.$$

In den nächsten beiden Abschnitten suchen wir nun Kennzahlen bei Matrizen, die sich unter Ähnlichkeitstransformationen nicht ändern. Das bedeutet, dass diese Kennzahlen der lineare Abbildung zuzuordnen sind, unabhängig vom betrachteten Koordinatensystem.

2.3 Determinanten

Sei $n \in \mathbb{N}$ und $A \in \mathbb{R}^{n \times n}$ eine **reelle quadratische** Matrix. Dann werden durch die zugehörige lineare Abbildung $x \mapsto Ax$ die Koordinateneinheitsvektoren $e^{(j)} \in \mathbb{R}^n$ in die Spaltenvektoren $Ae^{(j)} = a_{*j} \in \mathbb{R}^n$ von A abgebildet. Die $\{e^{(j)} : j = 1, \dots, n\}$ spannen den „Einheitswürfel“ im $\mathbb{R}^n$ auf mit Volumen 1. Die Bilder $\{a_{*j} : j = 1, \dots, n\}$ spannen ein Parallelepiped auf. Welches Volumen hat dies?

Wir beantworten diese Frage zunächst für die Fälle $n = 2$ und $n = 3$:

Spezialfall $n = 2$: Die Fläche des von zwei Vektoren $a, b \in \mathbb{R}^2$ aufgespannten Parallelogramms ist

$$\left| b \cdot \frac{a^\perp}{\|a\|} \right| \|a\| = |a_1 b_2 - a_2 b_1|,$$

wenn $a^\perp = (a_2, -a_1)^\top$ den zu a senkrechten Vektor bezeichnet (Skizze!). Wir definieren die Abbildung

$$\det(a, b) = a_1 b_2 - a_2 b_1 \quad \text{für } a, b \in \mathbb{R}^2,$$

die je zwei Vektoren im $\mathbb{R}^2$ bis auf das Vorzeichen die Fläche des von ihnen aufgespannten Parallelogramms zuordnet.

Spezialfall $n = 3$: Um das Volumen des von drei Vektoren $a, b, c \in \mathbb{R}^3$ aufgespannten Spats zu bestimmen, betrachten wir zunächst die Fläche des von a und b aufgespannten Parallelogramms. Dazu definieren wir das Vektorprodukt im $\mathbb{R}^3$ (bzw. $\mathbb{C}^3$).

Definition 2.16 Sei $\mathbb{K} = \mathbb{R}$ oder $\mathbb{K} = \mathbb{C}$. Für $a, b \in \mathbb{K}^3$ ist das **Vektorprodukt** (oder Kreuzprodukt) $a \times b \in \mathbb{K}^3$ erklärt durch

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

Einige Eigenschaften lassen sich direkt durch Nachrechnen nachweisen.

Satz 2.17 Für $a, b, c \in \mathbb{K}^3$ und $\lambda \in \mathbb{K}$ gilt:

- (a) $a \times a = 0$,
- (b) $a \times b = -(b \times a)$,
- (c) $(a + b) \times c = a \times c + b \times c$, und $a \times (b + c) = a \times b + a \times c$,
- (d) $(\lambda a) \times b = \lambda(a \times b) = a \times (\lambda b)$,
- (e) $(a \times b) \cdot c = a \cdot (b \times c)$,
- (f) $a \times (b \times c) = (a \cdot c)b - (a \cdot b)c$.

Bemerkung: Für komplexe Vektoren bedeutet der Malpunkt in den letzten beiden Fällen, dass die Formel für das reelle Skalarprodukt verwendet wird, also ohne komplexe Konjugation.

Der folgende Satz gibt nun die gesuchte geometrische Interpretation des Vektorprodukts.

Satz 2.18 Es seien $a, b \in \mathbb{R}^3$. Dann gilt:

- (i) Es steht $a \times b$ senkrecht auf a und auf b , d.h. $(a \times b) \cdot a = 0$ und $(a \times b) \cdot b = 0$. Damit ist also die Richtung des Vektorprodukts festgelegt.
- (ii) Es ist $\|a \times b\| = \|a\| \|b\| |\sin \omega|$, wobei ω der von a und b eingeschlossene Winkel ist. Damit ist $\|a \times b\|$ die Fläche des von a und b aufgespannten Parallelogramms.

Beweis: (i) Es ist $(a \times b) \cdot b = a \cdot (b \times b) = a \cdot 0 = 0$ und $(a \times b) \cdot a = -(b \times a) \cdot a = 0$.

(ii) Mit $c = a \times b$ ist $\|a \times b\|^2 = (a \times b) \cdot c = a \cdot (b \times c)$ und $b \times c = b \times (a \times b) = \|b\|^2 a - (a \cdot b)b$, also

$$\begin{aligned} \|a \times b\|^2 &= a \cdot (b \times c) = \|a\|^2 \|b\|^2 - |a \cdot b|^2 \\ &= \|a\|^2 \|b\|^2 [1 - \cos^2 \omega] = \|a\|^2 \|b\|^2 \sin^2 \omega. \end{aligned}$$

Dies beendet den Beweis. □

Bemerkungen:

- (a) Durch die angegebenen Eigenschaften ist das Vektorprodukt bis auf ein Vorzeichen eindeutig charakterisiert. Der Vektor $a \times b$ zeigt in die Richtung, so dass die drei Vektoren $a, b, a \times b$ ein **Rechtssystem** bilden. Dies bedeutet: Man spreize die rechte Hand, sodass der Daumen in Richtung a weist, der Zeigefinger in Richtung b , und der Mittelfinger rechtwinklig zu Daumen und Zeigefinger zeigt. Dann weist der Mittelfinger in Richtung von $a \times b$.
- (b) Multiplikation der ersten beiden Koordinateneinheitsvektoren ergibt:

$$e^{(1)} \times e^{(2)} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} \times \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix} = e^{(3)}.$$

Analog ist $e^{(2)} \times e^{(3)} = e^{(1)}$ und $e^{(3)} \times e^{(1)} = e^{(2)}$.

- (c) Die Bedingung $a \times b = 0$ ist gleichbedeutend mit der **linearen Abhängigkeit** von a und b . Ist nämlich z.B. $a = \lambda b$, so folgt $a \times b = \lambda b \times b = 0$. Ist umgekehrt $a \times b = 0$, so folgt aus (ii), dass $\omega = 0$ oder π gilt. Dies ist aber genau dann der Fall, wenn a und b linear abhängig sind.

Die Projektion von c auf die Normale zur von a und b aufgespannten Ebene liefert die Höhe des Spats, so dass wir insgesamt das Volumen eines durch $a, b, c \in \mathbb{R}^3$ aufgespannten Spats durch

$$\|a \times b\| \|c\| |\cos \omega| = |(a \times b) \cdot c|$$

erhalten, wobei ω der Winkel zwischen $a \times b$ und c ist (Skizze!). Analog zum zweidimensionalen Fall definieren wir:

$$\begin{aligned} \det(a, b, c) &:= (a \times b) \cdot c = a_2 b_3 c_1 - a_3 b_2 c_1 + a_3 b_1 c_2 - a_1 b_3 c_2 + a_1 b_2 c_3 - a_2 b_1 c_3 \\ &= a_1(b_2 c_3 - b_3 c_2) - a_2(b_1 c_3 - b_3 c_1) + a_3(b_1 c_2 - b_2 c_1) \\ &= a_1(b_2 c_3 - b_3 c_2) - b_1(a_2 c_3 - a_3 c_2) + c_1(a_2 b_3 - a_3 b_2) \end{aligned} \quad (2.3)$$

für $a, b, c \in \mathbb{R}^3$. Also etwa

$$\det \begin{pmatrix} 3 & 1 & 2 \\ 0 & 2 & 3 \\ 1 & 1 & 1 \end{pmatrix} = 0 \cdot 1 \cdot 2 - 1 \cdot 2 \cdot 2 + 1 \cdot 1 \cdot 3 - 3 \cdot 1 \cdot 3 + 3 \cdot 2 \cdot 1 - 0 \cdot 1 \cdot 1 = -4.$$

Die vorletzte Identität nutzen wir für eine allgemeine Definition der Determinante einer $n \times n$ Matrix.

Definition 2.19 (Determinante)

Ist $A \in \mathbb{K}^{n \times n}$ eine quadratische Matrix, so ist $\det(A) = a_{11}$ im Fall $n = 1$ und für $n \geq 2$ wird rekursiv definiert

$$\det A = \sum_{i=1}^n (-1)^{i+1} a_{i1} \det A^{(i)},$$

wobei $A^{(i)}$ die $(n-1) \times (n-1)$ -Matrix bezeichnet, die durch Streichen der 1. Spalte und der i -ten Zeile entsteht.

Bemerkung: Für die Determinante einer Matrix wird auch abkürzend $|A| = \det A$ geschrieben.

Beispiel 2.20 (a) Wir berechnen

$$\begin{vmatrix} 3 & 2 & 0 \\ 1 & 0 & 2 \\ 2 & 1 & 1 \end{vmatrix} = 3 \cdot \begin{vmatrix} 0 & 2 \\ 1 & 1 \end{vmatrix} - 1 \cdot \begin{vmatrix} 2 & 0 \\ 1 & 1 \end{vmatrix} + 2 \cdot \begin{vmatrix} 2 & 0 \\ 0 & 2 \end{vmatrix} = 3(0 - 2) - (2 - 0) + 2(4 - 0) = 0.$$

(b) Für die Einheitsmatrix $I \in \mathbb{R}^{n \times n}$ gilt

$$\det(I_n) = 1.$$

Wir bezeichnen im Folgenden die j -te Spalte einer Matrix A mit a_{*j} und die i -te Zeile mit a_{i*} , $i, j = 1, \dots, n$ und machen zunächst zwei wichtige Beobachtungen.

Lemma 2.21 (a) Die Determinante einer Matrix ist linear bezüglich der ersten Spalte, d.h. mit $a_{*1} = a + b \in \mathbb{R}^n$ gilt

$$\det(a + b, a_{*2}, \dots, a_{*n}) = \det(a, a_{*2}, \dots, a_{*n}) + \det(b, a_{*2}, \dots, a_{*n})$$

und mit $\lambda \in \mathbb{R}$ ist

$$\det(\lambda a_{*1}, a_{*2}, \dots, a_{*n}) = \lambda \det(a_{*1}, a_{*2}, \dots, a_{*n}).$$

(b) Ist die Matrix $\tilde{A}$ gegeben durch Vertauschen der ersten und zweiten Spalte, so gilt

$$\det(\tilde{A}) = -\det(A).$$

Beweis: Die erste Beobachtung ergibt sich offensichtlich direkt aus der Definition. Für die Determinante nach Vertauschen der beiden ersten Spalten rechnen wir nach:

$$\begin{aligned} \det(\tilde{A}) &= \sum_{i=1}^n (-1)^{i+1} a_{i2} \left| \tilde{A}^{(i)} \right| \\ &= \sum_{i=1}^n (-1)^{i+1} a_{i2} \left[\sum_{l=1}^{i-1} (-1)^{l+1} a_{l1} |B^{(i),(l)}| + \sum_{l=i+1}^n (-1)^l a_{l1} |B^{(i),(l)}| \right] \\ &= \sum_{l=1}^n \left[\sum_{i=l+1}^n (-1)^{i+1} (-1)^{l+1} a_{l1} a_{i2} |B^{(i),(l)}| + \sum_{i=0}^{l-1} (-1)^{i+1} (-1)^l a_{l1} a_{i2} |B^{(i),(l)}| \right] \\ &= - \sum_{l=1}^n (-1)^{l+1} a_{l1} \left[\sum_{i=1}^{l-1} (-1)^{i+1} a_{i2} |B^{(i),(l)}| + \sum_{i=l+1}^n (-1)^i a_{i2} |B^{(i),(l)}| \right] \\ &= - \sum_{l=1}^n (-1)^{l+1} a_{l1} |A^{(l)}| = -\det(A). \end{aligned}$$

Bei diesen Rechnungen haben wir mit $B^{(i),(l)}$ die $(n-2) \times (n-2)$ Matrizen bezeichnet, die sich nach Streichung der ersten beiden Spalten und der i -ten und j -ten Zeile aus A ergeben. Außerdem haben wir die Konvention $\sum_{l=1}^0 = 0 = \sum_{l=n+1}^n$ genutzt. $\square$

Die rekursive Definition der Determinante reduziert die Berechnung einer $n \times n$ Determinante auf n Berechnungen von $(n-1) \times (n-1)$ Determinanten. Mit dem Lemma(b) zur Vertauschung von Spalten wird deutlich, dass wir dabei nicht notwendig die erste Spalte betrachten müssen. Dies führt uns auf die allgemeine Entwicklungsregel zur Berechnung von Determinanten. Mit der letzten Gleichung in (2.3) sehen wir, dass durch entsprechendes Umsortieren der Summanden auch Zeilen anstelle von Spalten zur Entwicklung genutzt werden können. Auf einen vollständigen Beweis verzichten wir hier.

Satz 2.22 (*Entwicklungsregel*)

Sei $k \in \{1, \dots, n\}$ eine feste Zeilennummer, also $a_{k1}, \dots, a_{kn}$ die Elemente der k -ten Zeile. Für jedes $j \in \{1, \dots, n\}$ entsteht aus A eine $(n-1) \times (n-1)$ -Matrix $A^{(kj)}$ durch Streichen der k -ten Zeile und j -ten Spalte. Dann gilt:

$$\det A = \sum_{j=1}^n (-1)^{j+k} a_{kj} \det A^{(kj)}.$$

Diese Methode heißt „Entwicklung der Determinante nach der k -ten Zeile“. Genauso kann nach der j -ten Spalte entwickelt werden (wie im Fall der Definition mit $j = 1$).

Zur Illustration der Aussage dienen die folgenden beiden Beispiele.

Beispiele 2.23 (a) Sei

$$A = \begin{pmatrix} -1 & 0 & 3 & 4 \\ 2 & 1 & 0 & 3 \\ 0 & 1 & -2 & 1 \\ 1 & 2 & 0 & 3 \end{pmatrix}.$$

Wir entwickeln nach der dritten Spalte, da dort 2 Nullen auftreten:

$$\begin{aligned} |A| &= 3 \cdot \begin{vmatrix} 2 & 1 & 3 \\ 0 & 1 & 1 \\ 1 & 2 & 3 \end{vmatrix} - 2 \cdot \begin{vmatrix} -1 & 0 & 4 \\ 2 & 1 & 3 \\ 1 & 2 & 3 \end{vmatrix} \\ &= 3[2(3-2) + 1(1-3)] - 2[-1(3-6) + 4(4-1)] = -30. \end{aligned}$$

Hier haben wir die beiden dreireihigen Determinanten nach der 1. Spalte bzw. 1. Zeile entwickelt.

(b) Sei

$$A = \begin{pmatrix} 3 & 0 & -2 & 1 & 3 \\ 4 & 2 & 1 & 3 & 1 \\ 1 & 0 & 0 & -2 & 2 \\ 0 & 2 & 0 & 0 & 4 \\ 0 & 3 & 0 & 0 & 0 \end{pmatrix}.$$

Bei dieser Matrix gibt es jede Menge Nullen. Wir entwickeln zuerst nach der 5. Zeile, dann nach

der 4. und dann nach der 3. Zeile:

$$\begin{aligned}
 |A| &= -3 \cdot \begin{vmatrix} 3 & -2 & 1 & 3 \\ 4 & 1 & 3 & 1 \\ 1 & 0 & -2 & 2 \\ 0 & 0 & 0 & 4 \end{vmatrix} = (-3) \cdot 4 \cdot \begin{vmatrix} 3 & -2 & 1 \\ 4 & 1 & 3 \\ 1 & 0 & -2 \end{vmatrix} \\
 &= -12 [(-6 - 1) - 2(3 + 8)] = 12 \cdot 29 = 348.
 \end{aligned}$$

□

Bemerkung: Man beachte, dass es bei der Berechnung von Determinanten sinnvoll ist, sich eine Spalte bzw. Zeile auszuwählen, die viele Nullen enthält, um den Rechenaufwand gering zu halten.

Um die Berechnung von Determinanten noch effizienter durchführen zu können, listen wir wesentliche Eigenschaften von Determinanten auf ohne die Beweise genau auszuführen.

Satz 2.24 (*Eigenschaften*)

- (a) $\det A^\top = \det A$ und $\det A^* = \overline{\det A}$ für jede Matrix $A \in \mathbb{K}^{n \times n}$,
- (b) **Linearität** bzgl. jeder Spalte und jeder Zeile (vgl. Lemma 2.21 (a)),
- (c) **Antikommutativität:** Bei Vertauschung von zwei Zeilen oder Spalten vertauscht sich das Vorzeichen (vgl. Lemma 2.21, (b)).
- (d) **Reihen- oder Spaltenakkumulation:** Der Wert der Determinante ändert sich nicht, wenn man zu einer Zeile (bzw. Spalte) ein Vielfaches einer anderen Zeile (bzw. Spalte) addiert.
- (e) **Multiplikationsformel:** $\det(AB) = \det A \det B$ für alle Matrizen $A, B \in \mathbb{K}^{n \times n}$.
- (f) $\det A \neq 0$ genau dann, wenn die Matrix A regulär ist. In diesem Fall gilt

$$\det(A^{-1}) = \frac{1}{\det A}.$$

Bemerkung: Mit (e) und (f) sehen wir, dass sich die Determinante bei Koordinatenwechsel, also bei einer Ähnlichkeitstransformation, nicht ändert, denn für eine invertierbare Matrix T gilt

$$\det(TAT^{-1}) = \det(T) \det(A) \det(T^{-1}) = \frac{\det(T) \det(A)}{\det(T)} = \det(A).$$

Die Eigenschaft (d) bietet die Möglichkeit, zur Berechnung von Determinanten das Gauß-Verfahren zu verwenden. Dieses eignet sich besonders für Matrizen mit keinen oder wenigen Nullen. Man beachte dabei nur die folgenden **Einschränkungen**:

- (a) Vertausche keine Zeilen! Sonst würde sich das Vorzeichen ändern.
- (b) Ersetze keine Zeile durch ein Vielfaches von ihr! Sonst würde die Determinante auch multipliziert bzw. dividiert werden müssen.

Beispiele 2.25

(a) Wir betrachten zunächst die $(4, 4)$ -Matrix

$$A = \begin{pmatrix} 1 & 2 & 3 & 1 \\ 1 & 3 & 3 & 2 \\ 2 & 4 & 3 & 3 \\ 1 & 1 & 1 & 1 \end{pmatrix}.$$

Mit Gauß Schritten, bei denen sich die Determinante jeweils nicht ändert, erhalten wir die Tableaus

$$\begin{array}{cccc} \boxed{1} & 2 & 3 & 1 \\ 1 & 3 & 3 & 2 \\ 2 & 4 & 3 & 3 \\ 1 & 1 & 1 & 1 \end{array} \longrightarrow \begin{array}{cccc} \mathbf{1} & 2 & 3 & 1 \\ 0 & \boxed{1} & 0 & 1 \\ 0 & 0 & -3 & 1 \\ 0 & -1 & -2 & 0 \end{array} \longrightarrow \begin{array}{cccc} \mathbf{1} & 0 & 3 & -1 \\ 0 & \mathbf{1} & 0 & 1 \\ 0 & 0 & -3 & \boxed{1} \\ 0 & 0 & -2 & 1 \end{array}$$

$$\begin{array}{cccc} \mathbf{1} & 0 & 0 & 0 \\ 0 & \mathbf{1} & 3 & 0 \\ 0 & 0 & -3 & \mathbf{1} \\ 0 & 0 & \boxed{1} & 0 \end{array} \longrightarrow \begin{array}{cccc} \mathbf{1} & 0 & 0 & 0 \\ 0 & \mathbf{1} & 0 & 0 \\ 0 & 0 & 0 & \mathbf{1} \\ 0 & 0 & \mathbf{1} & 0 \end{array}$$

Diese Determinante können wir jetzt leicht entwickeln und erhalten $\det A = -1$. Wir erkennen, dass wir eigentlich zu viel leergeräumt haben. Betrachten wir das zweite Tableau. Entwickeln wir nach der ersten Spalte, so können wir nur mit dem reduzierten Tableau

$$\begin{array}{ccc} \mathbf{1} & 0 & 1 \\ 0 & -3 & 1 \\ -1 & -2 & 0 \end{array}$$

weitermachen. Wir müssen also die Elemente in der markierten Zeile nicht mehr leerräumen.

(b) Bei dem folgenden Beispiel muss man etwas mehr aufpassen: Sei

$$A = \begin{pmatrix} 1 & -2 & 9 & -3 \\ 3 & 4 & 1 & 3 \\ 6 & 8 & -2 & 1 \\ 2 & 1 & 3 & 10 \end{pmatrix}.$$

Dann erhalten wir

$$\begin{array}{cccc} \boxed{1} & -2 & 9 & -3 \\ 3 & 4 & 1 & 3 \\ 6 & 8 & -2 & 1 \\ 2 & 1 & 3 & 10 \end{array} \longrightarrow \begin{array}{cccc} \mathbf{1} & -2 & 9 & -3 \\ 0 & 10 & -26 & 12 \\ 0 & 20 & -56 & 19 \\ 0 & \boxed{5} & -15 & 16 \end{array}$$

Jetzt können wir nach der ersten Spalte entwickeln, d.h. wir brauchen in den folgenden Tableaus die erste Zeile nicht mehr leerräumen:

$$\begin{array}{cccc} \mathbf{1} & -2 & 9 & -3 \\ 0 & 0 & \boxed{4} & -20 \\ 0 & 0 & 4 & -45 \\ 0 & \mathbf{5} & -15 & 16 \end{array} \longrightarrow \begin{array}{cccc} \mathbf{1} & -2 & 9 & -3 \\ 0 & 0 & \mathbf{4} & -20 \\ 0 & 0 & 0 & -25 \\ 0 & \mathbf{5} & -15 & 16 \end{array}$$

Dies entwickeln wir nach der ersten Spalte und es folgt

$$\det A = 1 \cdot \begin{vmatrix} 0 & 4 & -20 \\ 0 & 0 & -25 \\ 5 & -15 & 16 \end{vmatrix} = 5 \cdot \begin{vmatrix} 4 & -20 \\ 0 & -25 \end{vmatrix} = 5 \cdot 4 \cdot (-25) = -500.$$

□

2.4 Eigenwerte und Eigenvektoren

Abschließend wollen wir noch weitere wichtige Kennzahlen bei linearen Abbildungen einführen.

Definition 2.26 (*Eigenwert, Eigenvektor*)

Sei $A \in \mathbb{C}^{n \times n}$ eine quadratische Matrix. Eine Zahl $\lambda \in \mathbb{C}$ heißt **Eigenwert** mit zugehörigem **Eigenvektor** $v \in \mathbb{C}^n$, wenn $v \neq 0$ ist und $Av = \lambda v$ gilt.

Bemerkung: Wir können uns reelle Eigenwerte und Eigenvektoren veranschaulichen. Wenn wir die zugehörige lineare Abbildung betrachten, so besteht die Abbildung für Vektoren in dem durch den Eigenvektor aufgespannten Unterraum nur in einer Streckung bzw. Stauchung (und gegebenenfalls Punktspiegelung bei negativen Eigenwerten) um den Faktor, der durch den Eigenwert gegeben ist.

Schreiben wir $Av = \lambda v$ um in die Form $(A - \lambda I_n)v = 0$, wobei I_n die Einheitsmatrix ist, so erkennen wir, dass v ein quadratisches homogenes lineares Gleichungssystem lösen muss. Es gibt genau dann nichttriviale Lösungen, wenn die Determinante von $A - \lambda I_n$ verschwindet. Wir müssen also solche $\lambda \in \mathbb{C}$ bestimmen mit $\det(A - \lambda I_n) = 0$.

Definition 2.27 (*Charakteristisches Polynom*)

Die Funktion $p(\lambda) = \det(A - \lambda I_n)$ heißt **charakteristisches Polynom** der Matrix A . Die Nullstellen des charakteristischen Polynoms sind die Eigenwerte der Matrix A .

Entwickeln wir die Determinante so erhalten wir wirklich ein Polynom vom Grad n . Indem wir die Nullstellen dieses Polynoms berechnen, können wir die Eigenwerte der Matrix bestimmen. Eigenvektoren erhalten wir dann aus dem entsprechenden linearen Gleichungssystem.

Beispiel 2.28 Sei

$$A = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 3 & -1 \\ 1 & -1 & 3 \end{pmatrix}$$

gegeben. Wir bestimmen die Nullstellen des charakteristischen Polynoms,

$$0 = \det(A - \lambda I_3) = \det \begin{pmatrix} 1 - \lambda & 0 & 0 \\ 0 & 3 - \lambda & -1 \\ 1 & -1 & 3 - \lambda \end{pmatrix} = (1 - \lambda)[(3 - \lambda)^2 - 1].$$

Also sind $\lambda_1 = 1$, $\lambda_2 = 2$ und $\lambda_3 = 4$ Eigenwerte. Einen Eigenvektor v etwa zu $\lambda = 4$ erhalten wir aus dem linearen Gleichungssystem $(A - 4I_3)v = 0$. Wir berechnen mit dem Gauß-Verfahren

$$\begin{array}{ccc|c} -3 & 0 & 0 & 0 \\ 0 & -1 & -1 & 0 \\ 1 & -1 & -1 & 0 \end{array} \longrightarrow \begin{array}{ccc|c} -3 & 0 & 0 & 0 \\ 0 & -1 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{array}.$$

Also ist etwa $v = (0, 1, -1)^\top$ ein Eigenvektor zum Eigenwert $\lambda = 4$. □

Bemerkung: Mit Hilfe des charakteristischen Polynoms sehen wir nun auch, dass wir die Eigenwerte der durch die Matrix beschriebenen linearen Abbildung zuordnen können. Denn die Eigenwerte einer Matrix ändern sich nicht unter Ähnlichkeitstransformationen, da

$$\det(TAT^{-1} - \lambda I_n) = \det(TAT^{-1} - \lambda TT^{-1}) = \det(T(A - \lambda I_n)T^{-1}) = \det(A - \lambda I_n)$$

gilt. Für Eigenvektoren v zu einem Eigenwert λ gilt wegen $\lambda v = Av$ die Identität

$$\lambda Tv = TAv = TAT^{-1}Tv = (TAT^{-1})Tv,$$

d.h. der Vektor Tv ist Eigenvektor zum Eigenwert λ der Matrix TAT^{-1} . Wir haben mit den Eigenwerten und der Determinante Kennzahlen zu einer linearen Abbildung kennengelernt, die sich bei Koordinatenwechsel nicht verändern.

Wir können analog Eigenvektoren zu den anderen Eigenwerten bestimmen und würden dann eventuell eine wichtige Beobachtung machen, die wir allgemein in folgendem Satz festhalten.

Satz 2.29 Seien $A \in \mathbb{C}^{n \times n}$ und $\lambda_1, \dots, \lambda_k \in \mathbb{C}$ paarweise verschiedene Eigenwerte von A mit zugehörigen Eigenvektoren $v^{(1)}, \dots, v^{(k)} \in \mathbb{C}^n$. Dann sind die Vektoren $\{v^{(1)}, \dots, v^{(k)}\}$ linear unabhängig.

Beweis: Dies zeigen wir mit vollständiger Induktion nach k (vergleiche auch Satz 4.14). Für $k = 1$ ist dies klar, da $v^1 \neq 0$. Es seien nun $\{v^{(1)}, \dots, v^{(k)}\}$ linear unabhängig und es sei $\sum_{j=1}^{k+1} \rho_j v^{(j)} = 0$. Zu zeigen ist, dass alle $\rho_j = 0$ sind. Multiplikation dieser Gleichung mit A bzw. mit λ_{k+1} liefert

$$\sum_{j=1}^{k+1} \rho_j \lambda_j v^{(j)} = \sum_{j=1}^{k+1} \rho_j A v^{(j)} = A \left(\sum_{j=1}^{k+1} \rho_j v^{(j)} \right) = 0 \quad \text{bzw.} \quad \sum_{j=1}^{k+1} \rho_j \lambda_{k+1} v^{(j)} = 0.$$

Die Differenz beider Gleichungen ergibt $\sum_{j=1}^k \rho_j [\lambda_j - \lambda_{k+1}] v^{(j)} = 0$. Beachte, dass die Summation jetzt nur noch bis k geht, da das letzte Glied sich weghebt. Aus der Induktionsvoraussetzung folgt, dass alle Koeffizienten Null sind, d.h. $\rho_j [\lambda_j - \lambda_{k+1}] = 0$ für $j = 1, \dots, k$. Da alle λ_j verschieden sind, folgt $\rho_j = 0$ für $j = 1, \dots, k$. Dann folgt schließlich auch $\rho_{k+1} = 0$, womit die Behauptung bewiesen ist. $\square$

Eine sehr spezielle Matrix ist eine Matrix $A \in \mathbb{R}^{n \times n}$ mit Eigenwerten $\lambda_1, \dots, \lambda_n$ und linear unabhängigen Eigenvektoren $v^{(1)}, \dots, v^{(n)}$, zum Beispiel wenn alle Eigenwerte verschieden sind (s. Satz 2.29). Dann ist durch die Eigenvektoren eine Basis gegeben. Wählen wir als weitere Basis die kanonische Orthogonalbasis mit den Einheitsvektoren $e^{(j)}$ und die zugehörige Basistransformationsmatrix T mit

$$Tv^{(j)} = e^{(j)},$$

die die Koordinaten $v^{(j)}$ zur üblichen Basis auf die Koordinaten zur Basis $\{v^{(1)}, \dots, v^{(n)}\}$ abbildet. Dann gilt

$$TAT^{-1}e^{(j)} = TAv^{(j)} = \lambda Tv^{(j)} = \lambda e^{(j)}.$$

Daraus ist ersichtlich, dass

$$TAT^{-1} = \begin{pmatrix} \lambda_1 & 0 & 0 \\ 0 & \ddots & 0 \\ 0 & 0 & \lambda_n \end{pmatrix}$$

gilt. Die Ähnlichkeitstransformation führt auf eine Diagonalmatrix mit den Eigenwerten auf der Diagonalen. Matrizen mit dieser Eigenschaft heißen **diagonalisierbar**.

In Hinblick auf die zugehörige lineare Abbildung ist Diagonalisierbarkeit eine sehr angenehme Eigenschaft. Geometrisch bedeutet es, dass sich die lineare Abbildung in diesem Koordinatensystem vollständig zerlegen lässt, als Streckungen/Stauchungen um Faktoren λ_j in Richtung der Basisvektoren $v^{(j)}$. Aus diesem Grund begegnen uns Eigenwerte und Eigenvektoren häufiger in Anwendungen etwa bei elastischen Phänomenen als *Hauptspannungen* bzw. *Hauptspannungsrichtungen*.

3 Fouriertheorie

Es wurde bereits angedeutet, dass wir etwa die Menge der stetigen Funktionen über einem Intervall oder die Menge der Polynome über $\mathbb{C}$ als Vektorraum betrachten können. In diesem Kapitel werden wir sehen, wie Funktionen sogar als Vektoren in einem euklidischen Vektorraum angesehen werden können. Dies liefert Darstellungen bzw. Approximationen für Funktionen mit sehr weitreichenden Konsequenzen in den Anwendungen etwa in der digitalen Signalverarbeitung.

3.1 Trigonometrische Polynome

Ist eine Funktion f in eine Potenzreihe entwickelbar, so kann sie durch eine Linearkombination von Polynomen, also durch sehr einfache Funktionen, approximiert werden. Zu beachten ist aber:

- (i) Die Klasse der Funktionen, die sich als Potenzreihe beschreiben lässt, ist klein. So muss f beliebig oft differenzierbar sein. Aber selbst das reicht noch nicht aus, um f als Potenzreihe darstellen zu können, wie Beispiel 5.36 im ersten Semester zeigt.
- (ii) Die Approximation ist im Allgemeinen nur in einer Umgebung des Entwicklungspunktes „gut“. Dies haben wir an den Plots zu Potenz- und Taylorreihen gesehen.

Es gibt viele verschiedene Methoden, eine möglichst große Klasse von Funktionen auf ihrem Definitionsbereich durch einfache Funktionen anzunähern. Historisch ein Meilenstein in der Entwicklung der Naturwissenschaften sind die „Fourierreihen“. In der Fouriertheorie werden Funktionen auf einem beschränkten Intervall $I \subset \mathbb{R}$ durch trigonometrische Polynome approximiert. Es ist im Folgenden nützlich, die komplexe Notation zu wählen, also $f : I \rightarrow \mathbb{C}$ zu betrachten.

Definition 3.1 Ein (komplexes) **trigonometrisches Polynom** vom Grad n ist eine Funktion der Form

$$p(x) = \sum_{k=-n}^n c_k e^{ikx} = \sum_{k=-n}^n c_k [e^{ix}]^k, \quad x \in \mathbb{R},$$

mit Koeffizienten $c_k \in \mathbb{C}$.

Beachte, dass der Name „Polynom“ gerechtfertigt ist, indem wir die Summe in positive und negative Indizes aufteilen und $z = e^{ix}$ bzw. $z = e^{-ix}$ setzen.

Bemerkung: Wegen $x \in \mathbb{R}$ erhalten wir mit der Eulerschen Formel

$$\begin{aligned} p(x) &= \sum_{k=-n}^n c_k e^{ikx} = \sum_{k=-n}^n c_k (\cos kx + i \sin kx) \\ &= c_0 + \sum_{k=1}^n (c_k + c_{-k}) \cos kx + i(c_k - c_{-k}) \sin kx. \end{aligned}$$

Also können wir $p(x)$ auch darstellen durch

$$p(x) = a_0 + \sum_{k=1}^n a_k \cos kx + b_k \sin kx$$

mit $a_0 = c_0$, $a_k = c_k + c_{-k}$ und $b_k = i(c_k - c_{-k})$. Beachte, dass $a_0, a_k, b_k \in \mathbb{R}$ für $k = 1, \dots, n$ impliziert, dass $p : I \rightarrow \mathbb{R}$ reellwertig ist. In der komplexen Darstellung ist dies äquivalent zu der Bedingung

$$c_{-k} = \overline{c_k}.$$

Anhand der zweiten Darstellung sehen wir, dass die trigonometrischen Polynome auf der reellen Achse periodisch sind mit Periode $L = 2\pi$.

Definition 3.2 Die Funktion $f : \mathbb{R} \rightarrow \mathbb{C}$ heißt **periodisch** mit der Periode $L > 0$, wenn $f(x + L) = f(x)$ gilt für alle $x \in \mathbb{R}$.

Da sich periodische Funktionen mit Periode L stets durch $\tilde{f}(x) := f(\frac{Lx}{2\pi})$ zu einer 2π -periodische Funktion transformieren lassen, können wir im Folgenden ohne Einschränkung 2π -periodische trigonometrische Polynome betrachten.

Wir erinnern uns an das erste Kapitel, in dem wir uns mit Vektorräumen, Unterräumen und Skalarprodukten beschäftigt haben. Der Raum

$$\mathcal{T}_n := \left\{ f : [-\pi, \pi] \rightarrow \mathbb{C} : f(x) = \sum_{|k| \leq n} c_k e^{ikx} : c_k \in \mathbb{C} \right\}$$

ist ein Vektorraum über $\mathbb{C}$, wobei die „Vektoren“ hier Funktionen sind. Die Vektorraum-Eigenschaften sehen wir, indem wir direkt die Bedingungen (V1)-(V9) in Definition 1.2 für die Addition von Funktionen $f, g \in \mathcal{T}_n$, d.h. $f + g \in \mathcal{T}_n$ ist gegeben durch $(f + g)(x) = f(x) + g(x)$, $x \in [-\pi, \pi]$, und die skalare Multiplikation $\lambda f \in \mathcal{T}_n$ mit $(\lambda f)(x) = \lambda f(x)$ für $x \in [-\pi, \pi]$ nachprüfen.

Offensichtlich ist $\mathcal{T}_n$ ein Unterraum von $\mathcal{T}_N$ mit $N > n$. Weiterhin bilden mit dieser Addition und skalaren Multiplikation etwa die stetigen Funktionen einen Vektorraum. Man bezeichnet diesen mit $C(I) = \{f : I \rightarrow \mathbb{C} : f \text{ ist stetig}\}$. Bei der Definition des Integrals haben wir auch bereits gezeigt, dass die Summe und ein Vielfaches integrierbarer Funktionen auch integrierbar ist und die Eigenschaften (V1)-(V9) gelten, d.h. die Menge der integrierbaren Funktionen $L(I)$ ist auch ein Vektorraum. Die Menge $\mathcal{T}_n$ kann somit auch als linearer Unterraum von $C([-\pi, \pi])$ oder als linearer Unterraum von $L((-\pi, \pi))$ aufgefasst werden.

Unser Ziel ist es nun, eine möglichst gute Näherung einer Funktion (aus einem der oben genannten Räume) durch ein trigonometrisches Polynom $p \in \mathcal{T}_n$ zu finden. Im $\mathbb{R}^3$ hatten wir durch Konstruktion des Lotfußpunktes eines Punkts auf eine Ebene (Unterraum) die Approximation mit dem kürzesten Abstand angeben können. Dazu haben wir das euklidische Skalarprodukt und die daraus sich ergebende euklidische Norm verwendet, also weitergehende Eigenschaften als nur die eines Vektorraums. Analog zum $\mathbb{R}^n$ bzw. $\mathbb{C}^n$ lässt sich auch über dem Vektorraum $\mathcal{T}_n$ ein Skalarprodukt definieren.

Lemma 3.3 Durch

$$\langle f, g \rangle := \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx$$

ist ein Skalarprodukt auf $\mathcal{T}_n$ für $n \in \mathbb{N}$ definiert.

Bemerkung: Das hier auftretende Integral über die komplexwertige Funktion $h = f\bar{g}$ wird berechnet, indem man den Integranden in Real- und Imaginärteil aufteilt:

$$\int_{-\pi}^{\pi} h(x) dx = \int_{-\pi}^{\pi} \operatorname{Re} h(x) dx + i \int_{-\pi}^{\pi} \operatorname{Im} h(x) dx$$

Beweis: Wir müssen die Eigenschaften (a)-(d) im Satz 1.23 verifizieren. Offensichtlich gilt

$$\langle f, g \rangle = \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx = \overline{\int_{-\pi}^{\pi} \overline{f(x)} g(x) dx} = \overline{\langle g, f \rangle}.$$

Die Distributivität, d.h. $\langle (f+g), h \rangle = \langle f, h \rangle + \langle g, h \rangle$ und die Assoziativität, $\lambda \langle f, g \rangle = \langle \lambda f, g \rangle = \langle f, \bar{\lambda} g \rangle$ ergeben sich direkt aus der Linearität des Integrals.

Mit Satz 6.8 (e) des ersten Semesters folgt aus

$$\langle f, f \rangle = \int_{-\pi}^{\pi} |f(x)|^2 dx = 0,$$

dass $|f(x)|^2 = 0$ für fast alle $x \in [-\pi, \pi]$ ist. Da $f \in \mathcal{T}_n$ stetig ist, folgt, dass $f = 0$ auf $[-\pi, \pi]$. Also gilt auch die Definitheit $f = 0 \Leftrightarrow \langle f, f \rangle = 0$ für Funktionen in $\mathcal{T}_n$. $\square$

Somit ist $\mathcal{T}_n$ ein euklidischer Vektorraum. Durch

$$\|f\|_2 = \sqrt{\langle f, f \rangle} = \sqrt{\int_{-\pi}^{\pi} |f(x)|^2 dx}$$

ist auf $\mathcal{T}_n$ eine Norm gegeben mit den in Satz 1.23, (e)-(h) aufgelisteten Eigenschaften. Durch

$$\|f - g\|_2 = \sqrt{\int_{-\pi}^{\pi} |f(x) - g(x)|^2 dx}$$

ist ein „Abstand“, das **quadratische Mittel**, zwischen Funktionen definiert.

Bemerkung: Die Festlegung auf das Intervall $[-\pi, \pi]$ ist vorteilhaft, aber nicht wesentlich. Es kann auch ein beliebiges anderes Intervall der Länge 2π , zum Beispiel $[0, 2\pi]$, verwendet werden. Denn ist $f : \mathbb{R} \rightarrow \mathbb{C}$ periodisch mit Periode 2π und $f|_{[-\pi, \pi]}$ integrierbar, so ist f auf jedem Intervall der Länge 2π integrierbar, und es gilt

$$\int_a^{a+2\pi} f(t) dt = \int_{-\pi}^{\pi} f(t) dt \quad \text{für jedes } a \in \mathbb{R}.$$

Das sieht man folgendermaßen: Wir spalten das Integrationsgebiet auf:

$$\int_a^{a+2\pi} f(t) dt = \int_a^{-\pi} f(t) dt + \int_{-\pi}^{\pi} f(t) dt + \int_{\pi}^{a+2\pi} f(t) dt.$$

Hierbei ist es unerheblich, ob z.B. $a > -\pi$ oder $a < -\pi$ ist. Die Substitution $t = s + 2\pi$ im letzten Integral liefert

$$\int_a^{-\pi} f(t) dt + \int_{\pi}^{a+2\pi} f(t) dt = \int_a^{-\pi} f(t) dt + \int_{-\pi}^a f(s) ds = 0.$$

Mit dem Skalarprodukt zeigt sich nun der wesentliche Vorteil, den die Wahl der Basis in der Fouriertheorie beinhaltet.

Satz 3.4 *Die Funktionen*

$$\varphi_k(x) := \frac{1}{\sqrt{2\pi}} e^{ikx}, \quad |k| \leq n,$$

bilden eine Orthonormalbasis des $\mathcal{T}_n$ bezüglich des Skalarproduktes $\langle \cdot, \cdot \rangle$. Insbesondere hat der Vektorraum $\mathcal{T}_n$ die Dimension $2n + 1$.

Bemerkung: Wir sprechen von einem **Orthogonalsystem** $W = \{v^{(1)}, \dots, v^{(m)}\}$ in einem euklidischen Vektorraum V , wenn die Vektoren paarweise orthogonal sind, d.h. $(v^{(i)}, v^{(j)}) = 0$ für $i \neq j$. Ist W sogar eine Basis von V und haben die Vektoren alle die Norm $\|v^{(i)}\| = 1$, $i = 1, \dots, m$, so heißt W **Orthonormalbasis**.

Beweis: Es gilt

$$\langle \varphi_k, \varphi_m \rangle = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{i(k-m)x} dx = \begin{cases} 1 & \text{für } k = m \\ \frac{1}{2\pi i(k-m)} e^{i(k-m)x} \Big|_{-\pi}^{\pi} = 0 & \text{für } k \neq m. \end{cases}$$

Die φ_k haben also die Norm 1 und sind paarweise orthogonal. Damit ergibt sich insbesondere, dass die Funktionen linear unabhängig sind; denn wir erhalten aus $0 = \sum_{k=-n}^n c_k \varphi_k = 0$ für die Koeffizienten

$$0 = \left\langle \sum_{k=-n}^n c_k \varphi_k, \varphi_l \right\rangle = \sum_{k=-n}^n c_k \langle \varphi_k, \varphi_l \rangle = c_l$$

für $l = -n, \dots, n$. Also ist $\{\varphi_k : k = -n, \dots, n\}$ eine Orthonormalbasis zu $\mathcal{T}_n$. $\square$

Im Unterraum $\mathcal{T}_n$ der trigonometrischen Polynome bis zum Grad n sind durch φ_k Basisvektoren gegeben analog zu den Einheitsvektoren im $\mathbb{R}^n$ und wir können die Koeffizienten c_k als „Koordinaten“ bezüglich dieser Basis verstehen.

Zu einer allgemeinen Funktion $f : [-\pi, \pi] \rightarrow \mathbb{C}$ suchen wir nun das trigonometrische Polynom $p_n \in \mathcal{T}_n$ mit dem kürzesten Abstand $\|f - p_n\|_2$. Um diesen Abstand formulieren zu können, muss nur vorausgesetzt werden, dass das Integral $\|f - p_n\|_2^2 = \int_{-\pi}^{\pi} |f(x) - p_n(x)|^2 dx$ existiert. Dies ist gewährleistet, wenn wir voraussetzen, dass $|f|^2$ Lebesgue-integrierbar ist, d.h. wir betrachten Funktionen f mit

$$f \in L^2(-\pi, \pi) := \left\{ f \in L(-\pi, \pi) : \int_{-\pi}^{\pi} |f(x)|^2 dx < \infty \right\}.$$

Die Menge $L^2(-\pi, \pi)$ ist ein Vektorraum; denn es gilt die Cauchy-Schwarzsche Ungleichung,

$$|\langle f, g \rangle| = \left| \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx \right| \leq \left(\int_{-\pi}^{\pi} |f(x)|^2 dx \right)^{\frac{1}{2}} \left(\int_{-\pi}^{\pi} |g(x)|^2 dx \right)^{\frac{1}{2}} = \|f\|_2 \|g\|_2,$$

(s. auch Satz 1.23). Damit existiert mit $f, g \in L^2(-\pi, \pi)$ auch

$$\begin{aligned} \int_{-\pi}^{\pi} |f(x) + g(x)|^2 dx &= \int_{-\pi}^{\pi} |f(x)|^2 + 2\operatorname{Re}(f(x)\overline{g(x)}) + |g(x)|^2 dx \\ &\leq \|f\|_2^2 + 2\|f\|_2\|g\|_2 + \|g\|_2^2 \leq \infty. \end{aligned}$$

Da auch $\lambda f \in L^2(-\pi, \pi)$ für $f \in L^2(-\pi, \pi)$ gilt und die Eigenschaften (V1)-(V9) erfüllt sind, ist $L^2(-\pi, \pi)$ ein Vektorraum. Weiterhin ist wegen der Cauchy-Schwarz Ungleichung durch

$$\langle f, g \rangle = \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx$$

ein Skalarprodukt zu $f, g \in L^2(-\pi, \pi)$ gegeben, so dass der Raum $L^2(-\pi, \pi)$ ein euklidischer Vektorraum ist.

Nun betrachten wir den Abstand einer Funktion $f \in L^2(-\pi, \pi)$ zu einem Polynom $q = \sum_{|k| \leq n} q_k \varphi_k \in \mathcal{T}_n$. Es gilt unter Ausnutzung der Orthogonalität

$$\begin{aligned} \|f - q\|_2^2 &= \langle f - q, f - q \rangle = \langle f, f \rangle - 2 \operatorname{Re} \langle f, q \rangle + \langle q, q \rangle \\ &= \langle f, f \rangle - 2 \operatorname{Re} \left\langle f, \sum_{|k| \leq n} q_k \varphi_k \right\rangle + \left\langle \sum_{|k| \leq n} q_k \varphi_k, \sum_{|k| \leq n} q_k \varphi_k \right\rangle \\ &= \langle f, f \rangle - 2 \operatorname{Re} \sum_{|k| \leq n} \overline{q_k} \langle f, \varphi_k \rangle + \sum_{|k| \leq n} |q_k|^2 \\ &= \|f\|_2^2 - \sum_{|k| \leq n} |\langle f, \varphi_k \rangle|^2 + \sum_{|k| \leq n} |q_k - \langle f, \varphi_k \rangle|^2. \end{aligned} \quad (3.4)$$

Wir sehen, dass der Abstand am kleinsten ist unter allen $q \in \mathcal{T}_n$, wenn die Koeffizienten so gewählt werden, dass

$$q_k = \langle q, \varphi_k \rangle = \langle f, \varphi_k \rangle$$

gilt. Bezeichnen wir dieses spezielle trigonometrische Polynom mit $p_n \in \mathcal{T}_n$, so besagt die Bedingung, dass

$$\langle f - p_n, \varphi_k \rangle = 0 \quad \text{für alle } |k| \leq n$$

oder kurz $(f - p_n) \perp \mathcal{T}_n$ gilt. Also ist p_n der „Lotfußpunkt“ $p_n \in \mathcal{T}_n$ der Funktion $f \in L^2(-\pi, \pi)$ in $\mathcal{T}_n$. Das trigonometrische Polynom p_n hat den kürzesten Abstand zu f unter allen trigonometrischen Polynomen der Ordnung kleiner oder gleich n , d.h. $\|f - p_n\|_2 \leq \|f - q\|_2$ für alle $q \in \mathcal{T}_n$ bzw.

$$\int_{-\pi}^{\pi} |f(x) - p_n(x)|^2 dx \leq \int_{-\pi}^{\pi} |f(x) - q(x)|^2 dx \quad \text{für alle } q \in \mathcal{T}_n.$$

Weiter folgt für die Koeffizienten c_k von $p_n = \sum_{|k| \leq n} c_k \varphi_k$ aus der Bedingung, dass $c_k = \langle p_n, \varphi_k \rangle = \langle f, \varphi_k \rangle$ ist. Also ist

$$p_n(x) = \sum_{|k| \leq n} \langle f, \varphi_k \rangle \varphi_k(x) = \frac{1}{2\pi} \sum_{k=-n}^n \left(\int_{-\pi}^{\pi} f(t) e^{-ikt} dt \right) e^{ikx}.$$

Definition 3.5 Für eine komplexwertige Funktion $f \in L^2(-\pi, \pi)$ und $n \in \mathbb{N}$ heißt

$$p_n(x) := \sum_{k=-n}^n c_k e^{ikx}, \quad x \in \mathbb{R},$$

$$\text{mit} \quad c_k := \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t) e^{-ikt} dt, \quad k \in \{-n, \dots, n\},$$

das (komplexe) **Fourierpolynom** n -ten Grades zu f mit den **Fourierkoeffizienten** c_k .

Mit der Euler'schen Formel können wir die Fourierpolynome auch umschreiben zu

$$p_n(x) = a_0 + \sum_{k=1}^n a_k \cos kx + b_k \sin kx,$$

wobei für die Koeffizienten

$$a_0 = c_0 = \frac{1}{2\pi} \int_{-\pi}^{+\pi} f(t) dt,$$

$$a_k = c_k + c_{-k} = \frac{1}{2\pi} \int_{-\pi}^{+\pi} f(t) (e^{-ikt} + e^{ikt}) dt = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \cos kt dt,$$

$$b_k = i(c_k - c_{-k}) = i \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t) (e^{-ikt} - e^{ikt}) dt = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \sin kt dt.$$

folgt.

Ist f **reellwertig**, so bieten sich diese reellen trigonometrischen Polynome an; denn wegen $c_{-k} = \overline{c_k}$ (s. Bemerkung nach Definition 3.1) ergibt sich

$$a_k = c_k + c_{-k} = 2 \operatorname{Re} c_k \in \mathbb{R} \quad \text{und} \quad b_k = i(c_k - c_{-k}) = -2 \operatorname{Im} c_k \in \mathbb{R}.$$

Für die praktische Berechnung von Fourierpolynomen im Reellen sind somit die folgenden Formulierungen nützlich, wobei bei den symmetrischen Fällen die Substitution $x \mapsto -x$ ausgenutzt wird.

Korollar 3.6

(a) Ist $f \in L^2(-\pi, \pi)$ eine reellwertige Funktion, so ist das Fourierpolynom $p_n(x) = a_0 + \sum_{k=1}^n a_k \cos kx + b_k \sin kx$ gegeben durch die reellen Koeffizienten

$$a_0 = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(t) dt, \quad a_k = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \cos kt dt, \quad b_k = \frac{1}{\pi} \int_{-\pi}^{\pi} f(t) \sin kt dt, \quad \text{für } k \in \mathbb{N}.$$

(b) Ist f darüber hinaus eine **gerade** Funktion, d.h. $f(x) = f(-x)$ für alle x , so ist $b_k = 0$ und

$$a_0 = \frac{1}{\pi} \int_0^{\pi} f(t) dt, \quad a_k = \frac{2}{\pi} \int_0^{\pi} f(t) \cos kt dt, \quad \text{für } k \in \mathbb{N}.$$

(c) Ist f **ungerade**, d.h. $f(x) = -f(-x)$ für alle x , so gilt $a_k = 0$ und

$$b_k = \frac{2}{\pi} \int_0^{\pi} f(t) \sin kt dt, \quad \text{für } k \in \mathbb{N}.$$

Beispiele 3.7

(a) Sei $f(x) = x$ für $|x| < \pi$. Dann ist f eine ungerade Funktion, also $a_k = 0$ für alle k und

$$\begin{aligned} b_k &= \frac{2}{\pi} \int_0^{\pi} t \sin kt dt = -\frac{2}{\pi k} t \cos kt \Big|_0^{\pi} + \frac{2}{\pi k} \int_0^{\pi} \cos kt dt \\ &= -\frac{2(-1)^k}{k} + \frac{2}{\pi k^2} \sin kt \Big|_0^{\pi} = \frac{2(-1)^{k+1}}{k}, \end{aligned}$$

also ist

$$p_n(x) = 2 \sum_{k=1}^n \frac{(-1)^{k+1}}{k} \sin(kx).$$

das n -te Fourierpolynom zu f .

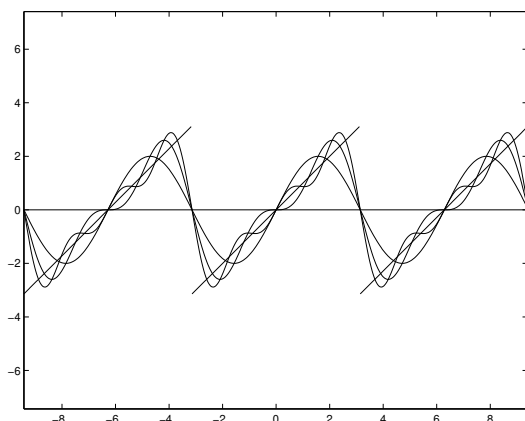
(b) Sei jetzt $f(x) = \operatorname{sign} x$, $|x| < \pi$. Dann ist f ebenfalls ungerade, also $a_k = 0$ für alle k und

$$b_k = \frac{2}{\pi} \int_0^{\pi} \sin kt dt = -\frac{2}{\pi k} \cos kt \Big|_0^{\pi} = \frac{2}{\pi k} [1 - (-1)^k] = \begin{cases} \frac{4}{\pi k}, & k \text{ ungerade,} \\ 0, & k \text{ gerade,} \end{cases}$$

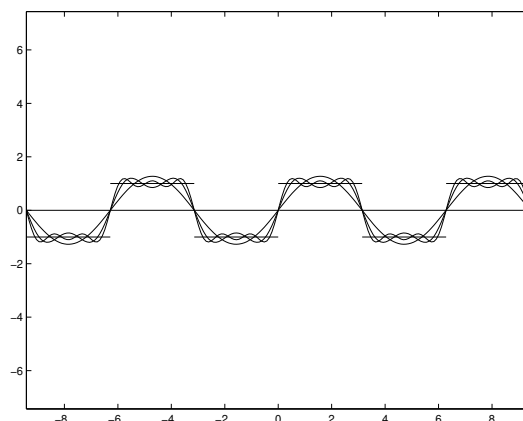
also

$$p_{2n+1}(x) = \frac{4}{\pi} \sum_{k=0}^n \frac{1}{2k+1} \sin((2k+1)x).$$

Für beide Beispiele sind die Funktionen und die drei ersten Fourierpolynome geplottet.



$f(x) = x, |x| \leq \pi$, und p_1, p_2, p_3



$f(x) = \text{sign}(x), |x| \leq \pi$, und p_1, p_3, p_5

Anscheinend approximiert das Fourierpolynom außerhalb des Intervalls $[-\pi, \pi]$ die 2π -**periodische Fortsetzung** der Funktion f , d.h. die Funktion $\tilde{f} : \mathbb{R} \rightarrow \mathbb{C}$ mit

$$\tilde{f}(x) := f(x - 2j\pi), \quad \text{für } -\pi + 2j\pi \leq x < \pi + 2j\pi, \quad j \in \mathbb{Z}.$$

Beachte, dass wir hier in den Stellen $\pi + 2j\pi$ den Funktionswert $f(-\pi)$ gewählt haben. Genauso hätten wir auch $f(\pi)$ verwenden können. Beide Fortsetzungen unterscheiden sich höchstens an den Nahtstellen (also auf einer Nullmenge!).

3.2 Fourierreihen

Nun interessiert das Konvergenzverhalten, wenn der Grad des Fourierpolynoms gegen unendlich geht. Wir beginnen mit einer allgemeinen Abschätzung, die insbesondere zeigt, dass die Fourierkoeffizienten für $k \rightarrow \pm\infty$ gegen 0 konvergieren.

Satz 3.8 (Besselsche Ungleichung)

Sind (c_k) die Fourierkoeffizienten einer Funktion $f \in L^2(-\pi, \pi)$ so gilt

$$\sum_{k=-n}^n |c_k|^2 \leq \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx \quad \text{für alle } n \in \mathbb{N}.$$

Da die Summe auf der linken Seite nur nichtnegative Summanden hat, konvergiert die Reihe $\left(\sum_{k=-\infty}^{\infty} |c_k|^2 \right)$. Insbesondere konvergieren die Fourierkoeffizienten c_k gegen 0 für $k \rightarrow \pm\infty$.

Beweis: Mit den Beziehungen $|u - v|^2 = |u|^2 + |v|^2 - 2 \operatorname{Re}(u \bar{v})$ und $|u|^2 = u \bar{u}$ für komplexe Zahlen $u, v \in \mathbb{C}$ folgt aus der Rechnung in (3.4)

$$0 \leq \int_{-\pi}^{+\pi} \left| f(x) - \sum_{|k| \leq n} c_k e^{ikx} \right|^2 dx = \int_{-\pi}^{\pi} |f(x)|^2 dx - 2\pi \sum_{|k| \leq n} |c_k|^2.$$

Dies beweist die Besselsche Ungleichung. □

Die Aussage, dass für $f \in L^2(-\pi, \pi)$

$$c_k = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-ikx} dx \rightarrow 0, \quad k \rightarrow \pm\infty$$

bzw. reell ausgedrückt

$$\int_{-\pi}^{\pi} f(x) \cos kx dx \rightarrow 0, \quad \text{und} \quad \int_{-\pi}^{\pi} f(x) \sin kx dx \rightarrow 0, \quad k \rightarrow \infty,$$

gilt, wird Riemann-Lebesgue-Lemma genannt.

Wir wollen jetzt zeigen, dass für $n \rightarrow \infty$ die Fouriersumme von f gegen die Funktion f konvergiert. Wir müssen hier aufpassen mit der Art der Konvergenz. (Aus dem ersten Semester kennen wir ja verschiedene Formen: punktweise und gleichmäßige Konvergenz, Konvergenz im quadratischen Mittel). Wir beginnen mit stückweise differenzierbaren Funktionen und der punktweisen Konvergenz.

Definition 3.9 Sei I ein beschränktes Intervall mit Endpunkten $a < b$. Eine Funktion $f : I \rightarrow \mathbb{C}$ heißt **stückweise stetig**, wenn es eine Zerlegung

$$a = x_0 < x_1 < \dots < x_n = b \quad \text{von} \quad [a, b]$$

gibt, so dass alle Einschränkungen $f|_{(x_{j-1}, x_j)}$ für $j = 1, \dots, n$ stetig und stetig fortsetzbar auf $[x_{j-1}, x_j]$ sind.

Mit $PC[a, b]$ sei der Raum der auf $[a, b]$ stückweise stetigen Funktion bezeichnet („PC“ soll an „piecewise continuous“ erinnern).

Wir benötigen noch die folgende Bezeichnung: Für $f \in PC[-\pi, \pi]$ definieren wir die links- und rechtsseitigen Grenzwerte durch

$$f_{\pm}(x) := \lim_{\substack{y \rightarrow x \\ y \gtrless x}} f(y) \quad \text{für alle inneren Punkte} \quad x \in (-\pi, \pi).$$

Ist f am Punkt x stetig, so ist natürlich $f_+(x) = f_-(x)$. Für die Randpunkte $\pm\pi$ sind $f_+(-\pi)$ und $f_-(-\pi)$ analog erklärt. Wir setzen noch $f_-(-\pi) := f_-(\pi)$ und $f_+(\pi) := f_+(-\pi)$.

In den folgenden Sätzen treten die Mittelwerte $\frac{1}{2}[f_+(x) + f_-(x)]$ auf. Diese Mittelwerte stimmen natürlich mit $f(x)$ überein, wenn f in x stetig ist. An den Randpunkten $\pm\pi$ ist aber im allgemeinen $f(-\pi) \neq f(\pi)$.

Satz 3.10 Sei $f \in PC[-\pi, \pi]$. Für alle $n \in \mathbb{N} \cup \{0\}$ und $x \in [-\pi, \pi]$ ist

$$\frac{1}{2}[f_+(x) + f_-(x)] = \sum_{|k| \leq n} c_k e^{ikx} + R_n(x),$$

wobei c_k die Fourierkoeffizienten von f sind und

$$R_n(x) = \frac{1}{2\pi} \left[\int_{x-\pi}^x \{f_-(x) - f(t)\} \frac{\sin[(2n+1)\frac{x-t}{2}]}{\sin \frac{x-t}{2}} dt + \int_x^{x+\pi} \{f_+(x) - f(t)\} \frac{\sin[(2n+1)\frac{x-t}{2}]}{\sin \frac{x-t}{2}} dt \right].$$

Beweis: Wir setzen die Funktion $f|_{[-\pi, \pi]}$ auf $\mathbb{R}$ zu einer 2π -periodischen Funktion fort. Wir definieren die Funktion

$$A_n(s) = \operatorname{Re} \sum_{k=-n}^n e^{iks} = 1 + 2 \operatorname{Re} \sum_{k=1}^n e^{iks}, \quad s \in \mathbb{R},$$

und rechnen mit der geometrischen Summenformel aus (Übung):

$$A_n(s) = \frac{\sin[(2n+1)\frac{s}{2}]}{\sin \frac{s}{2}}, \quad s \in \mathbb{R}.$$

A_n ist also die Funktion, die in R_n auftritt. Wir beweisen jetzt die Behauptung mit vollständiger Induktion nach n .

Sei zuerst $n = 0$. Dann ist wegen $c_0 = \frac{1}{2\pi} \int_{x-\pi}^{x+\pi} f(t) dt$:

$$\begin{aligned} \frac{1}{2}[f_+(x) + f_-(x)] - c_0 &= \frac{1}{2\pi} \left[\pi f_+(x) + \pi f_-(x) - \int_{x-\pi}^{x+\pi} f(t) dt \right] \\ &= \frac{1}{2\pi} \left[\int_{x-\pi}^x \{f_-(x) - f(t)\} dt + \int_x^{x+\pi} \{f_+(x) - f(t)\} dt \right] = R_0(x). \end{aligned}$$

Sei nun die Behauptung für n richtig. Dann gilt für $n + 1$

$$\begin{aligned}
& \frac{1}{2} [f_+(x) + f_-(x)] - \sum_{k=-(n+1)}^{n+1} c_k e^{ikx} \\
&= R_n(x) - c_{-(n+1)} e^{-i(n+1)x} - c_{n+1} e^{i(n+1)x} \\
&= \frac{1}{2\pi} \left\{ \int_{x-\pi}^x \{f_-(x) - f(t)\} A_n(x-t) dt + \int_x^{x+\pi} \{f_+(x) - f(t)\} A_n(x-t) dt \right. \\
&\quad \left. - \int_{x-\pi}^{x+\pi} f(t) \underbrace{\left[e^{-i(n+1)(x-t)} + e^{i(n+1)(x-t)} \right]}_{=2 \operatorname{Re} \exp[i(n+1)(x-t)]} dt \right\}.
\end{aligned}$$

Wegen

$$\operatorname{Re} \int_{x-\pi}^x e^{i(n+1)(x-t)} dt = -\operatorname{Re} \left[\frac{1}{i(n+1)} (1 - e^{i(n+1)\pi}) \right] = 0$$

und analog

$$\operatorname{Re} \int_x^{x+\pi} e^{i(n+1)(x-t)} dt = 0$$

ist

$$\begin{aligned}
& \frac{1}{2} [f_+(x) - f_-(x)] - \sum_{|k| \leq (n+1)} c_k e^{ikx} \\
&= \frac{1}{2\pi} \left[\int_{x-\pi}^x \{f_-(x) - f(t)\} \underbrace{[A_n(x-t) + 2 \operatorname{Re} e^{i(n+1)(x-t)}]}_{=A_{n+1}(x-t)} dt \right. \\
&\quad \left. + \int_x^{x+\pi} \{f_+(x) - f(t)\} \underbrace{[A_n(x-t) + 2 \operatorname{Re} e^{i(n+1)(x-t)}]}_{=A_{n+1}(x-t)} dt \right] = R_{n+1}(x).
\end{aligned}$$

Dies beweist den Satz. □

Die punktweise Konvergenz kann jetzt gezeigt werden, wenn auch f' stückweise stetig ist.

Satz 3.11 Sei f stückweise stetig differenzierbar, d.h. es gebe $-\pi = x_0 < x_1 < \dots < x_n = \pi$ so, dass alle Restriktionen $f|_{(x_{j-1}, x_j)}$ und ihre Ableitungen $f'|_{(x_{j-1}, x_j)}$ stetig fortsetzbar auf $[x_{j-1}, x_j]$ sind. Dann konvergiert die Fourierreihe **punktweise**, und es ist

$$\frac{1}{2}[f_+(x) + f_-(x)] = \sum_{k \in \mathbb{Z}} c_k e^{ikx}, \quad x \in [-\pi, \pi].$$

Außerdem konvergiert die Reihe im quadratischen Mittel, d.h. mit

$$p_n(x) = \sum_{|k| \leq n} c_k e^{ikx}, \quad x \in [-\pi, \pi],$$

gilt:

$$\int_{-\pi}^{\pi} |f(x) - p_n(x)|^2 dx \rightarrow 0 \quad \text{für } n \rightarrow \infty.$$

Beweis: Wir schätzen das Restglied $R_n(x)$ im letzten Satz für festes x ab. Dazu definieren wir

$$\psi(t) := \begin{cases} \frac{f_-(x) - f(t)}{\sin[(x-t)/2]}, & t \in [x-\pi, x), \\ \frac{f_+(x) - f(t)}{\sin[(x-t)/2]}, & t \in (x, x+\pi], \end{cases}$$

und setzen ψ mit der L'Hospitalschen Regel von rechts oder von links stückweise stetig in den Punkt $t = x$ fort. Nun können wir schreiben:

$$\begin{aligned} R_n(x) &= \frac{1}{2\pi} \int_{x-\pi}^{x+\pi} \psi(t) \sin[(n+1/2)(x-t)] dt = \frac{1}{2\pi} \int_{-\pi}^{\pi} \psi(x-s) \sin[(n+1/2)s] ds \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \psi(x-s) \left[\sin ns \cos \frac{s}{2} + \cos ns \sin \frac{s}{2} \right] ds. \end{aligned}$$

Damit ist $R_n(x)$ die Summe der n -ten Fourierkoeffizienten von $s \mapsto \frac{1}{2}\psi(x-s)\cos\frac{s}{2}$ bzgl. der Sinus- und von $s \mapsto \frac{1}{2}\psi(x-s)\sin\frac{s}{2}$ bzgl. der Kosinusfunktion. Die Besselsche Ungleichung liefert $R_n(x) \rightarrow 0$ für $n \rightarrow \infty$. Damit ist die punktweise Konvergenz gezeigt. Schließlich können wir $R_n(x)$ noch gleichmäßig bzgl. x und n beschränken, denn es gibt $c > 0$ mit $|\psi(t)| \leq c$ für alle $t \in [x-\pi, x+\pi]$ und daher wegen der vorletzten Formel:

$$|R_n(x)| \leq c \quad \text{für alle } x \in [-\pi, \pi].$$

Der Konvergenzsatz von Lebesgue (Satz 5.2), liefert jetzt, dass auch $\int_{-\pi}^{\pi} |R_n(x)|^2 dx \rightarrow 0$ für $n \rightarrow \infty$. Dies liefert die zweite Behauptung, wenn man ausnutzt, dass $\frac{1}{2}[f_+(x) + f_-(x)] = f(x)$ für alle x gilt bis auf endlich viele. $\square$

Beispiele 3.12

- (a) Wir haben schon die Fourierschen Polynome für die Funktion $f(x) = x$, $x \in [-\pi, \pi]$, berechnet:

$$p_n(x) = 2 \sum_{k=1}^n \frac{(-1)^{k+1}}{k} \sin(kx).$$

Die Voraussetzungen des Satzes sind erfüllt, also

$$x = 2 \sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k} \sin(kx), \quad x \in (-\pi, \pi).$$

Für $x = \pi/2$ ist $\sin(k\pi/2) = 0$ für gerades k und $\sin((2j+1)\pi/2) = (-1)^j$ für ungerades $k = 2j+1$. Also erhalten wir die **Gregory-Leibniz-Reihe**

$$\frac{\pi}{4} = 1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + - \dots$$

- (b) Betrachte nun $f(x) = x^2$, $|x| \leq \pi$. Dann ist f gerade, also $b_k = 0$ für alle k und

$$\begin{aligned} a_0 &= \frac{1}{\pi} \int_0^{\pi} x^2 dx = \frac{\pi^2}{3}, \\ a_k &= \frac{2}{\pi} \int_0^{\pi} x^2 \cos kx dx = \frac{2}{\pi k} x^2 \sin kx \Big|_0^{\pi} - \frac{4}{\pi k} \int_0^{\pi} x \sin kx dx \\ &= \frac{4}{\pi k^2} x \cos kx \Big|_0^{\pi} - \frac{4}{\pi k^2} \int_0^{\pi} \cos kx dx = \frac{4}{k^2} (-1)^k, \end{aligned}$$

also

$$x^2 = \frac{\pi^2}{3} + 4 \sum_{k=1}^{\infty} \frac{(-1)^k}{k^2} \cos kx, \quad |x| \leq \pi.$$

Für $x = \pi$ und $x = 0$ erhalten wir speziell:

$$\begin{aligned} \pi^2 &= \frac{\pi^2}{3} + 4 \sum_{k=1}^{\infty} \frac{1}{k^2}, \quad \text{also} \quad \sum_{k=1}^{\infty} \frac{1}{k^2} = \frac{\pi^2}{6}, \\ 0 &= \frac{\pi^2}{3} + 4 \sum_{k=1}^{\infty} \frac{(-1)^k}{k^2}, \quad \text{also} \quad \sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k^2} = \frac{\pi^2}{12}. \end{aligned}$$

Addition der beiden Formeln ergibt:

$$\sum_{k=1}^{\infty} \frac{1}{(2k-1)^2} = \frac{\pi^2}{8}.$$

Abschliessend zeigen wir den Hauptsatz der Fouriertheorie, der die letzte Aussage in Satz 3.11 für alle quadratintegrierbaren Funktionen verallgemeinert.

Satz 3.13 (Fourierscher Entwicklungssatz)

Sei $f \in L^2(-\pi, \pi)$ mit zugehörigen Fourierpolynomen (p_n) . Dann konvergiert (p_n) im **quadratischen Mittel** gegen f , d.h.

$$\|p_n - f\|_{L^2(-\pi, \pi)}^2 = \int_{-\pi}^{\pi} |p_n(x) - f(x)|^2 dx \rightarrow 0, \quad n \rightarrow \infty.$$

Die in diesem Sinne konvergente Reihe

$$\left(\sum_{k \in \mathbb{Z}} c_k e^{ikx} \right) \quad \text{mit} \quad c_k = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-ikx} dx$$

heißt **Fourierreihe**. Es gilt die Parsevalsche Gleichung

$$\sum_{k=-\infty}^{\infty} |c_k|^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx.$$

Bemerkung: Die Parsevalsche Gleichung kann auch für die reellen Fourierkoeffizienten hingeschrieben werden:

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx = |a_0|^2 + \frac{1}{2} \sum_{k=1}^{\infty} (|a_k|^2 + |b_k|^2).$$

Dies nachzurechnen, überlassen wir den Lesern als Übungsaufgabe.

Beweis: Wir benutzen ein Ergebnis der Lebesgueschen Integrationstheorie ohne Beweis. Durch unseren Zugang (Approximation durch Treppenfunktionen) ist es aber sehr plausibel:

Zu jedem $f \in L^2(-\pi, \pi)$ gibt es eine f approximierende Folge $(\varphi_k)_{k \in \mathbb{N}}$ von Treppenfunktionen, d.h. $\|\varphi_k - f\|_{L^2(-\pi, \pi)}^2 \rightarrow 0$ für $k \rightarrow \infty$.

Sei also jetzt $\varepsilon > 0$. Wähle k so groß, dass $\|\varphi_k - f\|_{L^2(-\pi, \pi)} \leq \varepsilon/2$. Wir halten dieses k fest und schreiben $\varphi = \varphi_k$. Die Treppenfunktion φ ist sicher stückweise stetig differenzierbar. Also gilt nach Satz 3.11, dass $\|q_n - \varphi\|_{L^2(-\pi, \pi)} \rightarrow 0$ für $n \rightarrow \infty$, wobei q_n das n -te Fourierpolynom zu φ ist. Insbesondere gibt es ein n_0 mit $\|q_n - \varphi\|_{L^2(-\pi, \pi)} \leq \varepsilon/2$ für alle $n \geq n_0$. Jetzt nutzen wir aus, dass das Fourierpolynom p_n von f die beste Approximation an f unter allen trigonometrischen Polynomen vom Grad höchstens n ist, also

$$\begin{aligned} \|p_n - f\|_{L^2(-\pi, \pi)} &\leq \|q_n - f\|_{L^2(-\pi, \pi)} \leq \|q_n - \varphi\|_{L^2(-\pi, \pi)} + \|\varphi - f\|_{L^2(-\pi, \pi)} \\ &\leq \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon \end{aligned}$$

für alle $n \geq n_0$. Dies ist gerade die Konvergenz von p_n gegen f in $L^2(-\pi, \pi)$.

Mit der oben eingeführten Orthonormalbasis berechnen wir weiterhin

$$\|p_n\|_{L^2(-\pi, \pi)}^2 = \left(\sum_{|k| \leq n} c_k \varphi_k, \sum_{|j| \leq n} c_j \varphi_j \right) = \sum_{|k| \leq n} \sum_{|j| \leq n} c_k \overline{c_j} (\varphi_k, \varphi_j) = \sum_{|k| \leq n} |c_k|^2.$$

Mit der Dreiecksungleichung gilt $\|f\|_{L^2(-\pi,\pi)} - \|p_n - f\|_{L^2(-\pi,\pi)} \leq \|p_n\|_{L^2(-\pi,\pi)} \leq \|f\|_{L^2(-\pi,\pi)} + \|p_n - f\|_{L^2(-\pi,\pi)}$. Also liefert diese Einschachtelung mit der gezeigten Konvergenz im Grenzfall, $n \rightarrow \infty$, die Parsevalsche Gleichung. $\square$

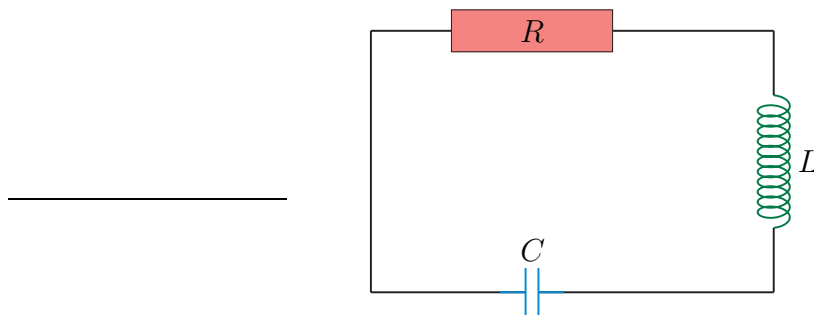
4 Differentialgleichungen

Am Ende des ersten Semesters haben wir einige Differentialgleichungen kennengelernt und ihre Bedeutung in Naturwissenschaft und Technik erahnen können. Mit einem Beispiel aus der Elektrotechnik wollen wir nun daran anknüpfen.

Beispiel 4.1 (*elektrischer Schwingkreis*)

Ein einfacher elektrischer Schwingkreis besteht aus einem Widerstand R , einer Spule der Induktivität L und einem Kondensator der Kapazität C . Nach dem Kirchhoff'schen Gesetz gilt für die Summe der Spannungsabfälle am Widerstand U_R , an der Spule U_L und am Kondensator U_C zu jedem Zeitpunkt t die Gleichung

$$0 = U_R(t) + U_C(t) + U_L(t).$$



Mit $I(t)$ bezeichnen wir den Strom, der zum Zeitpunkt $t > 0$ im Schwingkreis fließt, mit $Q(t)$ die Ladung auf den Kondensatorplatten. Weiter nutzen wir die übliche Schreibweise $\dot{F}$ anstelle von F' für Ableitungen bezüglich der Zeit. Es gelten die Zusammenhänge

$$\dot{Q}(t) = I(t), \quad U_R(t) = RI(t), \quad U_C(t) = \frac{Q(t)}{C}, \quad U_L(t) = L\dot{I}(t).$$

Durch Einsetzen erhalten wir ein **System** von zwei gewöhnlichen Differentialgleichungen erster Ordnung für die unbekannten Funktionen Q und I ,

$$\left. \begin{aligned} \dot{Q}(t) &= I(t) \\ \dot{I}(t) &= -\frac{R}{L}I(t) - \frac{1}{CL}Q(t) \end{aligned} \right\} \quad \text{für } t > 0.$$

Alternativ können wir nach dem Einsetzen das Kirchhoff'sche Gesetz einmal nach der Zeit differenzieren, und erhalten

$$\begin{aligned} 0 &= \ddot{I}(t) + \frac{R}{L}\dot{I}(t) + \frac{1}{CL}\dot{Q}(t) \\ &= \ddot{I}(t) + \frac{R}{L}\dot{I}(t) + \frac{1}{CL}I(t), \quad t > 0. \end{aligned}$$

Dies ist eine Differentialgleichung zweiter Ordnung für I , in der aber Q nicht mehr vorkommt. $\square$

In diesem Kapitel wollen wir uns mit der Theorie solcher Systeme erster Ordnung oder Differentialgleichungen höherer Ordnung beschäftigen und praktische Strategien zu ihrer Lösung kennenlernen. Dabei werden wir feststellen, dass wie im Beispiel beide Problemtypen eng zusammenhängen. Außerdem werden wir im wichtigen linearen Fall sehen, dass die Menge aller Lösungen die algebraische Struktur eines Vektorraums bzw. Unterraums aufweist, sodass Aussagen aus den letzten Kapiteln genutzt werden können. Wir beginnen mit den einfachsten Typen von Differentialgleichungen höherer Ordnung.

4.1 Homogene lineare Differentialgleichungen mit konstanten Koeffizienten

Die Differentialgleichung zweiter Ordnung aus dem Beispiel 4.1 hat eine sehr spezielle Form. Es handelt sich um eine **homogene lineare Differentialgleichung mit konstanten Koeffizienten**, d.h. eine Differentialgleichung n -ter Ordnung der Form

$$a_0 u(x) + a_1 u'(x) + \cdots + a_n u^{(n)}(x) = 0, \quad x \in I \subseteq \mathbb{R},$$

mit Konstanten $a_0, \dots, a_n$, wobei $a_n \neq 0$ vorausgesetzt ist.

Das Wort *linear* deutet darauf hin, dass im homogenen Fall die Summe von zwei Lösungen und ein Vielfaches einer Lösung stets wieder eine Lösung ist. Genauer gilt mit der Definition aus dem letzten Kapitel, dass die Abbildung $\alpha : C^n(I) \rightarrow C(I)$ mit

$$\alpha(u) = a_0 u + a_1 u' + \cdots + a_n u^{(n)}$$

linear ist. Mit $C^n(I)$ wird hier der Vektorraum der n -mal stetig differenzierbaren Funktionen bezeichnet. *Homogen* nennt man eine lineare Differentialgleichung, die die Nullfunktion als Lösung besitzt. Das ist gleichbedeutend damit, dass die *Inhomogenität* null ist. Würde auf der rechten Seite der Gleichung ein von u unabhängiger Term ungleich Null stehen, so heißt die Differentialgleichung *inhomogen*.

Ziel ist es, Lösungen dieser homogenen Differentialgleichung zu finden. Wir werden feststellen, dass es sinnvoll ist, komplexwertige Ausdrücke zu verwenden. Insbesondere dürfen die Zahlen a_j auch komplex sein. Beachten Sie, dass das Definitionsintervall I reell bleibt!

Um Lösungen der Differentialgleichung zu finden, bietet sich ein Exponentialansatz $u(x) = \exp(\lambda x)$ mit $\lambda \in \mathbb{C}$ an. Setzen wir diesen in die Differentialgleichung ein, so erhalten wir

$$(a_0 + a_1 \lambda + \cdots + a_n \lambda^n) e^{\lambda x} = 0.$$

Also löst u die Differentialgleichung, falls λ Nullstelle des **charakteristischen Polynoms**

$$p(\lambda) = \sum_{j=0}^n a_j \lambda^j$$

ist. Wir halten diese Überlegung als Satz fest.

Satz 4.2 Zu jeder Nullstelle $\hat{\lambda} \in \mathbb{C}$ des charakteristischen Polynoms, d.h. $a_0 + a_1 \lambda + \cdots + a_n \lambda^n = 0$, ist $u(x) = \exp(\hat{\lambda} x)$ eine Lösung der homogenen linearen Differentialgleichung.

Bemerkung: Die Bezeichnung *charakteristisches Polynom* knüpft an das vorletzte Kapitel an. Dies sehen wir, wenn wir die Differentialgleichung zu einem System von Differentialgleichungen erster Ordnung umschreiben.

Dazu definieren wir Funktionen $v_j(x) = u^{(j-1)}(x)$ für $j = 1, \dots, n$ und schreiben die Differentialgleichung mit den Funktionen v_j äquivalent um zu

$$\begin{aligned} v_1'(x) &= v_2(x) \\ v_2'(x) &= v_3(x) \\ &\vdots \\ v_{n-1}'(x) &= v_n(x) \\ v_n'(x) &= -\frac{1}{a_n} (a_0 v_1(x) + a_1 v_2(x) + \dots + a_{n-1} v_n(x)) . \end{aligned}$$

Dies ist ein lineares Differentialgleichungssystem. Fassen wir die Funktionen v_j zu einem Vektor zusammen und nutzen wir die Matrixschreibweise, so ist

$$\begin{pmatrix} v_1'(x) \\ v_2'(x) \\ \vdots \\ v_n'(x) \end{pmatrix} = \begin{pmatrix} 0 & 1 & 0 & \dots & 0 \\ & 0 & 1 & 0 & \\ & & \ddots & & \\ & & & 0 & 1 \\ -\frac{a_0}{a_n} & -\frac{a_1}{a_n} & -\frac{a_2}{a_n} & \dots & -\frac{a_{n-1}}{a_n} \end{pmatrix} \begin{pmatrix} v_1(x) \\ v_2(x) \\ \vdots \\ v_n(x) \end{pmatrix} .$$

Der Ansatz $(v_1(x), v_2(x), \dots, v_n(x))^T = e^{\lambda x} w$ mit einem konstanten Vektor $w \in \mathbb{C}^n \setminus \{0\}$, führt auf eine Lösung des Differenzialgleichungssystems, genau dann, wenn λ Eigenwert und w zugehöriger Eigenvektor der Matrix

$$A = \begin{pmatrix} 0 & 1 & 0 & \dots & 0 \\ & 0 & 1 & 0 & \\ & & \ddots & & \\ & & & 0 & 1 \\ -\frac{a_0}{a_n} & -\frac{a_1}{a_n} & -\frac{a_2}{a_n} & \dots & -\frac{a_{n-1}}{a_n} \end{pmatrix}$$

ist. Bis auf den Faktor $1/a_n$ ist das charakteristische Polynom dieser Matrix das ursprüngliche charakteristische Polynom $p(\lambda) = \sum_{j=0}^n a_j \lambda^j$ der homogenen linearen Differentialgleichung n -ter Ordnung (nachrechnen!).

Im Komplexen besitzt das charakteristische Polynom stets n Nullstellen und wir können p in Linearfaktoren zerlegen,

$$p(\lambda) = a_n (\lambda - \lambda_1) (\lambda - \lambda_2) \dots (\lambda - \lambda_n).$$

Jede der Zahlen λ_j , $j = 1, \dots, n$ liefert eine zugehörige Lösung $\exp(\lambda_j x)$. Auch Summen von skalaren Vielfachen dieser Funktionen, also **Linearkombinationen**, sind aufgrund der Linearität wieder Lösungen:

$$u(x) = c_1 e^{\lambda_1 x} + c_2 e^{\lambda_2 x} + \dots + c_n e^{\lambda_n x}.$$

Sind alle λ_j paarweise verschieden, so werden wir uns später überlegen, dass wir damit tatsächlich alle Lösungen der homogenen linearen Differentialgleichung mit konstanten Koeffizienten gefunden haben. Daher nennt man die Linearkombinationen aller dieser Lösungen auch die **allgemeine Lösung** der Differentialgleichung.

Zunächst aber greifen wir das Beispiel 4.1 wieder auf. In der Physik und Technik spielt die dort auftretende Differentialgleichung eine ausgezeichnete Rolle. Man nennt sie auch **gedämpfte lineare Schwingungsdifferentialgleichung**. Wir diskutieren ihre Lösungen:

Beispiel 4.3 Die Differentialgleichung des elektrischen Schwingkreises lautet

$$\ddot{I}(t) + \frac{R}{L} \dot{I}(t) + \frac{1}{CL} I(t) = 0,$$

wobei $L, C > 0$ vorausgesetzt wird. Hier ist

$$p(\lambda) = \lambda^2 + \frac{R}{L} \lambda + \frac{1}{CL} = \left(\lambda + \frac{R}{2L} \right)^2 - \frac{1}{4L^2} \left(R^2 - \frac{4L}{C} \right).$$

Wir unterscheiden drei Fälle:

1. Fall: $R^2 > 4L/C$. Dann sind die beiden Nullstellen von p gegeben durch $\lambda_{1,2} = \frac{1}{2L} [-R \pm \sqrt{R^2 - 4L/C}]$. Beide Nullstellen sind negativ und die allgemeine Lösung der Differentialgleichung lautet

$$I(t) = \alpha_1 e^{-\left(\frac{R}{2L} + \frac{1}{2L} \sqrt{R^2 - 4L/C}\right)t} + \alpha_2 e^{-\left(\frac{R}{2L} - \frac{1}{2L} \sqrt{R^2 - 4L/C}\right)t}, \quad \alpha_1, \alpha_2 \in \mathbb{R}.$$

Wir erkennen, dass die Lösung maximal eine Nullstelle besitzt und für $t \rightarrow \infty$ exponentiell gegen 0 abfällt. Daher heißt dieser Fall **Kriechfall** oder **stark gedämpfter Fall**.

2. Fall: $R^2 < 4L/C$. Dann sind beide Nullstellen von p komplex und gegeben durch $\lambda_{1,2} = \frac{1}{2L} [-R \pm i\sqrt{4L/C - R^2}]$. Die allgemeine komplexe bzw. reelle Lösung der Differentialgleichung lautet

$$I(t) = e^{-\frac{R}{2L}t} \left[\gamma_1 e^{\frac{i}{2L}\sqrt{4L/C - R^2}t} + \gamma_2 e^{-\frac{i}{2L}\sqrt{4L/C - R^2}t} \right], \quad \gamma_1, \gamma_2 \in \mathbb{C},$$

bzw.

$$I(t) = e^{-\frac{R}{2L}t} \left[\alpha_1 \cos\left(\frac{1}{2L}\sqrt{4L/C - R^2}t\right) + \alpha_2 \sin\left(\frac{1}{2L}\sqrt{4L/C - R^2}t\right) \right],$$

mit $\alpha_1, \alpha_2 \in \mathbb{R}$. Auch in diesem Fall fällt $I(t)$ exponentiell ab, schwingt aber mit der Frequenz $\sqrt{4L/C - R^2}/(2L)$. Dieser Fall heißt **gedämpfte Schwingung**. Ist $R = 0$, so ist die Schwingung **ungedämpft**.

3. Fall: $R^2 = 4L/C$. Dieser Fall heißt **aperiodischer Grenzfall**. Hier fallen beide Nullstellen zusammen: $\lambda_{1,2} = -R/(2L)$. Diesen Fall haben wir in der allgemeinen Theorie noch nicht behandelt (kommt gleich). Wir sehen aber direkt, dass die beiden Funktionen

$$I_1(t) = e^{-\frac{R}{2L}t} \quad \text{und} \quad I_2(t) = t e^{-\frac{R}{2L}t}$$

Lösungen sind. Die allgemeine Lösung hat die Form

$$I(t) = e^{-\frac{R}{2L}t} [\alpha_1 + \alpha_2 t], \quad \alpha_1, \alpha_2 \in \mathbb{R}.$$

Auch hier fällt $I(t)$ exponentiell ab, hat aber genau eine Nullstelle $t_0 > 0$, falls α_1 und α_2 verschiedenes Vorzeichen haben. $\square$

Auch wenn alle Koeffizienten der Differentialgleichung reell sind, können wie im Beispiel im Fall $R^2 < 4L/C$ komplexe Nullstellen des charakteristischen Polynoms und somit auch komplexwertige Lösungen auftreten.

Sind alle Koeffizienten reell, so treten komplexe Nullstellen immer paarweise als komplex konjugierte Zahlen auf, etwa $\lambda_{1/2} = \mu \pm i\eta$. Dazu gehören die komplex konjugierten Lösungen

$$u(x) = c_1 e^{(\mu+i\eta)x} + c_2 e^{(\mu-i\eta)x}.$$

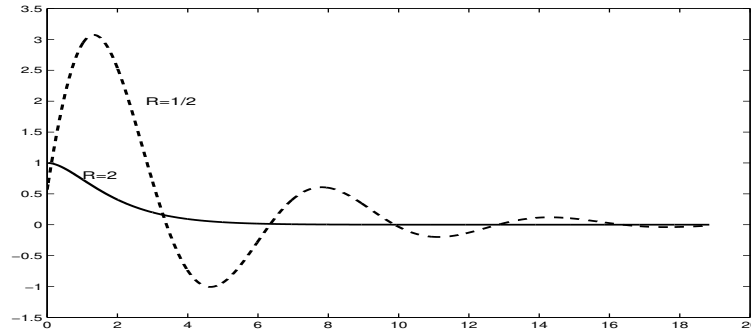


Abbildung 1: Lösungen zu $I''(t) + RI'(t) + I(t) = 0$, $I(0) = 1$

Sind wir an reellwertigen Lösungen interessiert, können wir ausnutzen, dass sich Real- und Imaginärteil komplexer Zahlen mit Hilfe der komplex konjugierten Zahl ausdrücken lassen:

$$\begin{aligned}\operatorname{Re}(e^{(\mu+i\eta)x}) &= \frac{1}{2} (e^{(\mu+i\eta)x} + e^{(\mu-i\eta)x}) = e^{\mu x} \cos(\eta x), \\ \operatorname{Im}(e^{(\mu+i\eta)x}) &= \frac{1}{2i} (e^{(\mu+i\eta)x} - e^{(\mu-i\eta)x}) = e^{\mu x} \sin(\eta x).\end{aligned}$$

Da es sich bei Real- und Imaginärteil von $e^{(\mu+i\eta)x}$ um Linearkombinationen von Lösungen handelt, sind diese auch selbst wieder Lösungen der Differentialgleichung. Wir haben die alternative Darstellung

$$u(x) = \tilde{c}_1 e^{\mu x} \cos(\eta x) + \tilde{c}_2 e^{\mu x} \sin(\eta x)$$

gefunden. Lassen wir für $\tilde{c}_1, \tilde{c}_2$ nur reelle Zahlen zu, so haben wir die Darstellung aller reellwertigen Lösungen.

Wir müssen noch den Fall behandeln, dass das charakteristische Polynom p doppelte oder mehrfache Nullstellen besitzt. Sei z.B. λ_1 k -fache Nullstelle von $p(\lambda) = a_0 + a_1\lambda + \dots + a_n\lambda^n$, d.h. $p(\lambda_1) = p'(\lambda_1) = \dots = p^{(k-1)}(\lambda_1) = 0$. Dann wissen wir, dass $u_1(x) = c e^{\lambda_1 x}$ Lösung der homogenen Differentialgleichung $a_0 u(x) + a_1 u'(x) + \dots + a_n u^{(n)}(x) = 0$ ist. Nun modifizieren wir u_1 und betrachten den Ansatz $u(x) = c(x) e^{\lambda_1 x}$. Induktiv folgt

$$u^{(j)}(x) = \sum_{k=0}^j \binom{j}{k} c^{(k)}(x) \lambda_1^{j-k} e^{\lambda_1 x}$$

Damit u Lösung der homogenen Differentialgleichung ist, muss somit

$$\begin{aligned}0 &= \sum_{j=0}^n a_j \left(\sum_{k=0}^j \binom{j}{k} c^{(k)}(x) \lambda_1^{j-k} e^{\lambda_1 x} \right) \\ &= \left(\sum_{j=1}^n a_j \sum_{k=1}^j \binom{j}{k} c^{(k)}(x) \lambda_1^{j-k} \right) e^{\lambda_1 x} + \underbrace{c(x) \sum_{j=0}^n a_j \lambda_1^j e^{\lambda_1 x}}_{=0}\end{aligned}$$

gelten, da λ_1 Nullstelle des charakteristischen Polynoms ist. Die Methode, die Konstante c in u_1 zu variieren, heißt **Reduktion der Ordnung**, da wir für c' eine Differentialgleichung

$$\sum_{j=1}^n a_j \sum_{k=1}^j \binom{j}{k} \lambda_1^{j-k} c^{(k)}(x) = 0$$

der Ordnung $n - 1$ erhalten. Vertauschen wir die Summationsreihenfolgen, so ergibt sich für $c'(x)$ die Forderung

$$\begin{aligned} 0 &= \sum_{j=0}^n a_j \sum_{k=1}^j \binom{j}{k} c^{(k)}(x) \lambda_1^{j-k} = \sum_{k=1}^n \left(\sum_{j=k}^n a_j \binom{j}{k} \lambda_1^{j-k} \right) c^{(k)}(x) \\ &= \sum_{k=1}^n \frac{p^{(k)}(\lambda_1)}{k!} c^{(k)}(x). \end{aligned}$$

Wenn λ_1 doppelte Nullstelle ist, so ist $p'(\lambda_1) = 0$ und die Forderung ist für $c(x) = x$ erfüllt. Also erhalten wir als eine weitere Lösung $u_2(x) = x e^{\lambda_1 x}$.

Wenn die Ordnung der Nullstelle λ_1 größer ist als 2, können wir offensichtlich Polynome höheren Grades für c wählen. Insgesamt haben wir den Beweis des folgenden Satzes skizziert.

Satz 4.4 Ist $\hat{\lambda} \in \mathbb{C}$ k -fache Nullstelle des charakteristischen Polynoms $p(\lambda) = a_0 + a_1 \lambda + \dots + a_n \lambda^n$, so sind die Funktionen der Form

$$u(x) = c_1 e^{\hat{\lambda} x} + c_2 x e^{\hat{\lambda} x} + \dots + c_k x^{k-1} e^{\hat{\lambda} x}$$

mit $c_1, \dots, c_n \in \mathbb{C}$ Lösungen der homogenen linearen Differentialgleichung.

Wir betrachten nun noch einen Differentialgleichungstyp, der sich durch eine Substitution auf eine lineare Differentialgleichung mit konstanten Koeffizienten transformieren lässt. Seien $a_0, \dots, a_n \in \mathbb{R}$ mit $a_n \neq 0$ gegeben. Dann ist durch

$$\sum_{j=0}^n a_j x^j u^{(j)}(x) = 0, \quad x > 0,$$

eine homogene **Eulersche Differentialgleichung** gegeben. Diese Differentialgleichung ist von n -ter Ordnung und hat die nichtkonstanten speziellen Koeffizientenfunktionen $a_j x^j$. Anfangsbedingungen darf man nicht bei $x = 0$ stellen, da dort der höchste Koeffizient $a_n x^n$ verschwindet.

Um Eulersche Differentialgleichungen zu lösen, bietet sich die Variablentransformation $x = \exp t$, also $v(t) = u(\exp t)$ für $t \in \mathbb{R}$ an. Dann ist $u(x) = v(\ln x)$, $x > 0$, $u'(x) = v'(t)/x$, $u''(x) = (v''(t) - v'(t))/x^2$ mit $t = \ln x$. Induktiv lässt sich nachweisen, dass jede Ableitung $u^{(k)}$ von der Form

$$u^{(k)}(x) = \frac{1}{x^k} \sum_{j=1}^k d_{kj} v^{(j)}(\ln x)$$

mit gewissen Konstanten $d_{kj} \in \mathbb{R}$ ist. Setzen wir dies in die Differentialgleichung ein, so ergibt sich eine lineare Differentialgleichung mit konstanten Koeffizienten von der Form

$$\sum_{k=0}^n \beta_k v^{(k)}(\ln x) = 0, \quad x > 0,$$

oder mit $x = \exp t$,

$$\sum_{k=0}^n \beta_k v^{(k)}(t) = 0, \quad t \in \mathbb{R}. \quad (4.5)$$

Beispiel 4.5

Sei $u''(x) - 2u(x)/x^2 = 0$, d.h. $x^2 u''(x) - 2u(x) = 0$, gegeben mit Anfangsbedingungen $u(1) = 0$, $u'(1) = 2$. Setze $u(x) = v(\ln x)$. Dann gilt $u'(x) = v'(t)/x$ und $u''(x) = (v''(t) - v'(t))/x^2$ mit $t = \ln x$. Also erhalten wir $v''(t) - v'(t) - 2v(t) = 0$, $t \in \mathbb{R}$. Das charakteristische Polynom ist $p(\lambda) = \lambda^2 - \lambda - 2$ mit Nullstellen -1 und 2 . Daher ist die allgemeine Lösung der homogenen Differentialgleichung

$$v(t) = a e^{-t} + b e^{2t}, \quad t \in \mathbb{R}, \quad a, b \in \mathbb{R}.$$

Die allgemeine Lösung für die Eulersche Differentialgleichung ist somit nach Rücksubstitution

$$u(x) = \frac{a}{x} + b x^2, \quad x > 0, \quad a, b \in \mathbb{R}.$$

Anpassen der Anfangsbedingung liefert $a = -\frac{2}{3}$ und $b = \frac{2}{3}$. □

Bemerkung: Eine andere, leichter zu merkende Möglichkeit, die Eulerschen Differentialgleichungen zu lösen, ist der direkte **Potenzansatz** $u(x) = x^\lambda$ für $\lambda \in \mathbb{C}$. Es gilt

$$u^{(k)}(x) = \lambda(\lambda - 1) \cdots (\lambda - k + 1) x^{\lambda - k}, \quad x > 0, \quad k = 1, 2, \dots, n.$$

Einsetzen in die Differentialgleichung und Koeffizientenvergleich führt auf die charakteristische Gleichung

$$a_0 + \sum_{k=1}^n a_k \underbrace{\lambda(\lambda - 1) \cdots (\lambda - k + 1)}_{k \text{ Faktoren}} = 0.$$

Die linke Seite ist wieder ein Polynom in λ vom Grad n . Es ist das charakteristische Polynom der Differentialgleichung (4.5). Das Polynom besitzt im Komplexen wieder n Nullstellen. Sind alle diese Nullstellen $\lambda_1, \dots, \lambda_n \in \mathbb{C}$ verschieden, so ist die allgemeine Lösung der Eulerdifferentialgleichung durch

$$u(x) = c_1 x^{\lambda_1} + c_2 x^{\lambda_2} + \cdots + c_n x^{\lambda_n}$$

gegeben.

Im Fall mehrfacher Nullstellen ergeben sich (s. obige Transformation) logarithmische Faktoren für die weiteren linear unabhängigen Lösungen. Wenn zum Beispiel λ dreifache Nullstelle ist, so erhalten wir bezüglich der Variablen $t = \ln x$ die Lösungen $v_1(t) = e^{\lambda t}$, $v_2(t) = t e^{\lambda t}$ und $v_3(t) = t^2 e^{\lambda t}$ und somit lauten die entsprechenden Lösungen der Euler Differentialgleichung $u_1(x) = x^\lambda$, $u_2(x) = \ln(x) x^\lambda$ und $u_3(x) = \ln^2(x) x^\lambda$.

4.2 Der Satz von Picard-Lindelöf

Wir wenden uns nun der Frage zu, ob ein Anfangswertproblem für eine Differentialgleichung höherer Ordnung oder für ein Differentialgleichungssystem stets eine Lösung besitzt. Dabei treten in Anwendungen durchaus sehr viel kompliziertere Gleichungen auf, als wir im letzten Abschnitt behandelt haben, wie das folgende Beispiel zeigt.

Beispiel 4.6 (Räuber-Beute Modell)

In einem abgeschlossenen System gebe es zwei Populationen, eine Räuber- und eine Beutepopulation (und sonst nichts) der Stärken $R(t)$ und $B(t)$ zur Zeit t . Die Räuberpopulation ernähre sich ausschließlich von den Beutetieren und habe keinen Feind, die Beutetiere haben eine andere unbegrenzte Nahrungsquelle und als einzigen Feind die Räuber. Ohne Beutetiere würden die

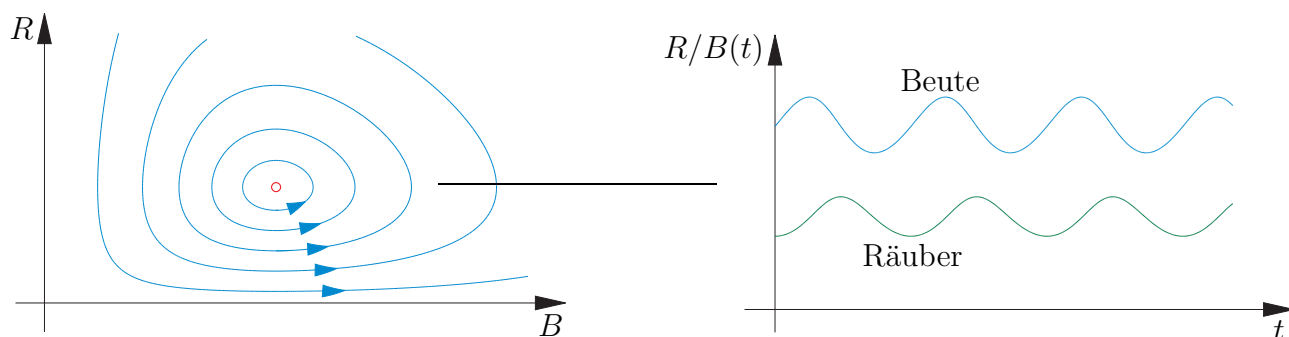
Räuber nach dem Gesetz $R'(t) = -\alpha_1 R(t)$ aussterben (Anzahl der sterbenden Räuber pro Zeiteinheit ist proportional zur Anzahl der Räuber). Die Häufigkeit der „Begegnungen“ zwischen Räuber und Beute sei proportional zu $R(t)B(t)$. Daher ergibt sich für die Änderungsrate der Räuber die Differentialgleichung

$$R'(t) = -\alpha_1 R(t) + \beta_1 R(t)B(t), \quad t \geq 0.$$

Analog erhält man für die Änderungsrate der Beutetiere die Differentialgleichung

$$B'(t) = \alpha_2 B(t) - \beta_2 R(t)B(t), \quad t \geq 0.$$

Diese beiden Differentialgleichungen bilden ein (nichtlineares) **System** erster Ordnung von zwei Differentialgleichungen mit zwei unbekannten Funktionen R und B . Dieses erste Modell zur Beschreibung von konkurrierenden Populationen heißt **Lotka - Volterra Modell**. Obwohl das System recht einfach ist, lassen sich geschlossene Ausdrücke für die Lösungen nicht angeben. In der Abbildung unten links sind für verschiedene Anfangswerte die Kurven $(B(t), R(t))$ (sogenannte Trajektorien) eingezeichnet, rechts sind die Funktionen B und R für einen Anfangswert zu sehen. $\square$



Um allgemein ein Differentialgleichungssystem zu beschreiben, nutzen wir die Notationen des Vektorraums $\mathbb{R}^n$. Seien $I \subset \mathbb{R}$ ein Intervall, $x_0 \in I$, $u^0 \in \mathbb{R}^n$ und eine Funktion $f : I \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ gegeben. Dann ist durch

$$\begin{aligned} u_1'(x) &= f_1(x, u_1(x), \dots, u_n(x)), & x \in I, \\ u_2'(x) &= f_2(x, u_1(x), \dots, u_n(x)), & x \in I, \\ &\vdots & \\ u_n'(x) &= f_n(x, u_1(x), \dots, u_n(x)), & x \in I, \end{aligned}$$

ein **Differentialgleichungssystem** erster Ordnung gegeben. Anfangsbedingungen sind durch $u_j(x_0) = u_j^0$, $j = 1, \dots, n$ beschrieben. Wir schreiben ein Anfangswertproblem (AWA) kürzer in Vektorform als

$$u'(x) = f(x, u(x)), \quad x \in I, \quad u(x_0) = u^0. \quad (4.6)$$

Somit sieht die allgemeine AWA formal aus wie die früheren Beispiele gewöhnlicher Differentialgleichungen erster Ordnung, die natürlich mit dieser Notation auch erfasst sind. Eine differenzierbare Funktion $u : \mathbb{R} \rightarrow \mathbb{R}^n$, die (4.6) erfüllt, heißt Lösung des Anfangswertproblems. Wir schreiben hier u^0 für den Vektor der Anfangswerte statt u_0 , um nicht mit den Komponenten von u durcheinander zu kommen.

Nach dem gezeigten Beispiel lässt sich vermuten, dass nur in speziellen Fällen geschlossene Lösungen zu einer AWA angegeben werden können. Umso bedeutungsvoller sind theoretische Aussagen,

die die Existenz (genau) einer Lösung zu einem Differentialgleichungssystem begründen, ohne diese anzugeben. Ein solcher Satz sei hier ohne Beweis beschrieben.

Satz 4.7 (Picard-Lindelöf)

Gegeben seien $x_0 \in \mathbb{R}$, $I = [x_0 - a, x_0 + a] \subset \mathbb{R}$ mit $a > 0$, $u^0 \in \mathbb{R}^n$ und $D := \{u \in \mathbb{R}^n : \|u - u^0\| \leq b\}$ mit $b > 0$. Weiter sei $f : I \times D \rightarrow \mathbb{R}^n$ stetig und bzgl. u Lipschitzstetig, d.h. es existiere eine Konstante $L > 0$ mit

$$\|f(x, u) - f(x, v)\| \leq L \|u - v\| \quad \text{für alle } u, v \in D, \quad x \in I.$$

Dann besitzt die Anfangswertaufgabe

$$u'(x) = f(x, u(x)), \quad x \in J, \quad u(x_0) = u^0,$$

genau eine stetig differenzierbare Lösung $u : J \rightarrow \mathbb{R}^n$ auf $J := [x_0 - h, x_0 + h] \subset I$ mit $h = \min\{a, \frac{b}{M}\}$ und $M = \max\{\|f(x, u)\| : x \in I, u \in D\}$.

Bemerkungen: (a) Der Satz lässt sich auch bei Differentialgleichungen höherer Ordnung anwenden. Betrachten wir etwa eine gewöhnliche Differentialgleichung (also für $v : I \subseteq \mathbb{R} \rightarrow \mathbb{R}$) n -ter Ordnung,

$$v^{(n)}(x) = g(x, v(x), v'(x), \dots, v^{(n-1)}(x)).$$

Setzen wir $u_1(x) = v(x)$, $u_2(x) = v'(x)$, $\dots$, $u_n(x) = v^{(n-1)}(x)$ so ist die gewöhnliche Differentialgleichung äquivalent zu dem n -dimensionalen System erster Ordnung,

$$\begin{pmatrix} u_1'(x) \\ u_2'(x) \\ \vdots \\ u_{n-1}'(x) \\ u_n'(x) \end{pmatrix} = \begin{pmatrix} u_2(x) \\ u_3(x) \\ \vdots \\ u_n(x) \\ g(x, u_1(x), \dots, u_n(x)) \end{pmatrix}.$$

Hier ist $f : I \times \mathbb{R}^n \rightarrow \mathbb{R}^n$ also die Funktion mit $f(x, u_1, \dots, u_n) = (u_2, u_3, \dots, u_n, g(x, u_1, \dots, u_n))^T$.

(b) Die Aussage kann völlig analog auch für andere Normen des Vektorraums $\mathbb{R}^n$ formuliert werden, etwa

$$\|z\|_\infty = \max_{j=1, \dots, n} |z_j| \quad (\text{Maximumnorm})$$

$$\|z\|_1 = \sum_{j=1}^n |z_j| \quad (\text{„Taxifahrernorm in New York“}).$$

4.3 Lineare Differentialgleichungssysteme

Im Abschnitt 4.1 haben wir lineare Differentialgleichungen n -ter Ordnung kennengelernt. Wir stellen zunächst den Zusammenhang zu Systemen erster Ordnung her.

Beispiel 4.8

Wir betrachten eine **lineare Differentialgleichung n -ter Ordnung** mit variablen Koeffizienten, d.h.

$$a_n(x) u^{(n)}(x) + a_{n-1}(x) u^{(n-1)}(x) + \dots + a_1(x) u'(x) + a_0(x) u(x) = b(x), \quad x \in I, \quad (4.7)$$

mit den Anfangsbedingungen $u^{(j)}(x_0) = u_j$, $j = 0, 1, \dots, n-1$. Die Funktionen a_j und b seien stetig und können reell oder komplex sein. Wir nehmen an, dass $a_n(x) \neq 0$ für alle $x \in I$ ist. Wie in der Bemerkung nach Satz 4.7 beschrieben, können wir eine solche Anfangswertaufgabe als ein System erster Ordnung schreiben, indem wir folgende Funktionen definieren,

$$v_1 := u, \quad v_2 := u' = v_1', \quad v_3 := u'' = v_2', \quad \dots, \quad v_n := u^{(n-1)} = v_{n-1}'.$$

Damit lautet die Differentialgleichung

$$a_n(x) v_n'(x) + a_{n-1}(x) v_n(x) + a_{n-2}(x) v_{n-1}(x) + \dots + a_1(x) v_2(x) + a_0(x) v_1(x) = b(x),$$

$x \in I$ und mit $v = (v_1, \dots, v_n)^\top$ folgt

$$v'(x) = \begin{pmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & 1 \\ -\frac{a_0(x)}{a_n(x)} & -\frac{a_1(x)}{a_n(x)} & \dots & \dots & -\frac{a_{n-1}(x)}{a_n(x)} \end{pmatrix} v(x) + \begin{pmatrix} 0 \\ \vdots \\ 0 \\ \frac{b(x)}{a_n(x)} \end{pmatrix}.$$

Auch die Anfangswerte schreiben wir als Vektor $v(x_0) = v^0 := (u_0, \dots, u_{n-1})^\top$. □

Das System aus dem Beispiel ist ein Spezialfall des wichtigen Falls linearer Differentialgleichungssysteme. Ein solches System hat die Form

$$\begin{aligned} u_1'(x) &= a_{11}(x)u_1(x) + a_{12}(x)u_2(x) + \dots + a_{1n}(x)u_n(x) + b_1(x), & x \in I, \\ u_2'(x) &= a_{21}(x)u_1(x) + a_{22}(x)u_2(x) + \dots + a_{2n}(x)u_n(x) + b_2(x), & x \in I, \\ &\vdots & \\ u_n'(x) &= a_{n1}(x)u_1(x) + a_{n2}(x)u_2(x) + \dots + a_{nn}(x)u_n(x) + b_n(x), & x \in I, \end{aligned}$$

mit Anfangsbedingungen $u_j(x_0) = u_j^0$, $j = 1, \dots, n$. Diese schreiben wir bequem mit Hilfe von Matrizen als

$$u'(x) = A(x)u(x) + b(x), \quad x \in I, \quad u(x_0) = u^0, \tag{4.8}$$

wobei $I \subset \mathbb{R}$ wieder ein Intervall, $x_0 \in I$, $u^0 \in \mathbb{R}^n$ seien, und Funktionen $A : I \rightarrow \mathbb{R}^{n \times n}$ und $b : I \rightarrow \mathbb{R}^n$ mit stetigen Funktionen in den Komponenten. Für jedes $x \in I$ ist $A(x)$ also eine (n, n) -Matrix und $b(x)$ ein Spaltenvektor.

Achtung: Wir werden im Folgenden Begriffe für lineare Systeme erster Ordnung (4.8) einführen und dann, immer unter der Voraussetzung $a_n(x) \neq 0$ für $x \in I$, auf lineare Differentialgleichungen n -ter Ordnung (4.7) übertragen.

Definition 4.9 Wenn $b = 0$ ist, so heißt das System $u' = Au$ **homogen**. Im Fall $b \neq 0$ wird es als **inhomogen** bezeichnet. Ein Differentialgleichungssystem der Form $u' = Au + b$ heißt **linear**; denn für zwei Lösungen u_1, u_2 des homogenen Systems ($b = 0$) lösen auch beliebige Summen $u = \lambda u_1 + \mu u_2$ mit $\lambda, \mu \in \mathbb{R}$ das homogene Differentialgleichungssystem.

Es lässt sich der Satz von Picard-Lindelöf anwenden, wobei eine Intervallbreite $h > 0$ unabhängig von x_0 gewählt werden kann. Dies liefert die folgende globale Existenz- und Eindeutigkeitsaussage.

Satz 4.10 Die AWA (4.8) besitzt genau eine stetig differenzierbare Lösung $u : I \rightarrow \mathbb{R}^n$.

Wie oben angedeutet, können wir diesen Satz insbesondere auf das Beispiel 4.8, also auf eine lineare Differentialgleichung n -ter Ordnung anwenden. Als Ergebnis halten wir fest:

Fazit: Anfangswertprobleme zu linearen Differentialgleichungen n -ter Ordnung besitzen für beliebige Anfangswerte $(u(x_0), u'(x_0), \dots, u^{(n-1)}(x_0)) = u^0 \in \mathbb{R}^n$ stets genau eine Lösung $u \in C^n(I)$, wenn $a_j, b \in C(I)$ gilt und $a_n(x) \neq 0$ für $x \in I$ ist.

Wie bei linearen Gleichungssystemen (s. Satz 2.9), sind die Lösungen des homogenen Systems, also für $b(x) = 0$, wesentlicher Bestandteil der allgemeinen Lösungstheorie zu $u'(x) = A(x)u(x) + b(x)$. Im Fall $n = 1$ können wir die allgemeine Lösung der homogenen Gleichung $u'(x) = a(x)u(x)$ angeben als $u(x) = c \exp \int_{x_0}^x a(t)dt$ mit einer Konstanten $c \in \mathbb{R}$ (s. HM I, Satz 6.25). Für $n > 1$ ist die Lösung im allgemeinen nicht mehr geschlossen darstellbar. Trotzdem können wir gewisse Aussagen zur Struktur von Lösungen machen. Zunächst bemerken wir, dass wegen der Linearität des Systems die Menge aller Lösungen der homogenen Differentialgleichung $u'(x) = A(x)u(x)$, $x \in I$, wieder ein Vektorraum ist, d.h. Summen und Vielfache von Lösungen sind wieder Lösungen. Im Fall $n = 1$ ist dieser Raum eindimensional mit Basis $u^1(x) = \exp \int_{x_0}^x a(t)dt$, $x \in I$. Im allgemeinen Fall ist der Raum n -dimensional. Eine explizite Basis kann allerdings im Allgemeinen nicht angegeben werden.

Satz 4.11 und Definition

- (a) Der Vektorraum V aller Lösungen des homogenen Differentialgleichungssystems $u'(x) = A(x)u(x)$ auf I ist n -dimensional. Eine Basis $\{u^1, \dots, u^n\}$ von V heißt **Fundamentalsystem** des Differentialgleichungssystems. Die Matrix $\Phi(x) = [u^1(x) \cdots u^n(x)] \in \mathbb{R}^{n \times n}$, $x \in I$ bestehend aus den Basisfunktionen wird **Fundamentalmatrix** genannt.
- (b) Sei $u^k : I \rightarrow \mathbb{R}^n$ die eindeutige Lösung von $(u^k)'(x) = A(x)u^k(x)$, $x \in I$, und $u^k(x_0) = e^k = (0, \dots, 0, 1, 0, \dots, 0)^\top$ der k -te Einheitsvektor im $\mathbb{R}^n$. Dann ist $\{u^1, \dots, u^n\}$ ein Fundamentalsystem der Differentialgleichung und die Fundamentalmatrix Φ dieses speziellen Fundamentalsystems hat die Eigenschaft $\Phi(x_0) = I_n$ ($n \times n$ -Einheitsmatrix). Es heißt **kanonisches Fundamentalsystem**.

Beweis: Seien die u^k , $k = 1, \dots, n$, wie in Teil (b) erklärt, d.h. die u^k lösen die Differentialgleichung $(u^k)'(x) = A(x)u^k(x)$, $x \in I$, mit Anfangsbedingungen $u^k(x_0) = e^k$. Wir zeigen, dass $\{u^1, \dots, u^n\}$ ein Fundamentalsystem bildet:

- (i) Die $\{u^1, \dots, u^n\}$ sind linear unabhängig: Ist nämlich $\sum_{k=1}^n \alpha_k u^k(x) = 0$ für alle $x \in I$, so gilt dies auch für $x = x_0$, also $0 = \sum_{k=1}^n \alpha_k e^k = (\alpha_1, \dots, \alpha_n)^\top$. Daher verschwinden alle α_k .
- (ii) Die Funktionen $\{u^1, \dots, u^n\}$ spannen den ganzen Lösungsraum auf: Sei u irgendeine Lösung der Differentialgleichung $u'(x) = A(x)u(x)$, $x \in I$. Es gilt $u(x_0) = \sum_{k=1}^n \alpha_k e^k$ mit $\alpha_k = u_k(x_0)$,

d.h. der k -ten Komponente von $u(x_0)$. Die Funktion $v(x) := \sum_{k=1}^n \alpha_k u^k(x)$, $x \in I$, löst ebenfalls die Differentialgleichung mit den Anfangsbedingungen $v(x_0) = u(x_0)$. Nach dem Satz von Picard-Lindelöf müssen u und v übereinstimmen. Also ist $u(x) = \sum_{k=1}^n \alpha_k u^k(x)$, $x \in I$, d.h. $\{u^1, \dots, u^n\}$ spannen den gesamten Lösungsraum auf.

Damit haben wir die Teile (a) und (b) bewiesen. $\square$

Beispiel 4.12 Wir erhalten ein Fundamentalsystem zur linearen Differentialgleichung n -ter Ordnung,

$$a_n(x) u^{(n)}(x) + a_{n-1}(x) u^{(n-1)}(x) + \dots + a_1(x) u'(x) + a_0(x) u(x) = 0, \quad x \in I,$$

indem wir der Reihe nach die (AWA) mit Anfangswerten

$$(u(x_0), u'(x_0), \dots, u^{(n-1)}(x_0)) = (1, 0, \dots, 0), (0, 1, 0, \dots, 0), \dots, (0, 0, \dots, 1)$$

lösen. Bezeichnen wir diese Lösungen der Differentialgleichung mit u_j , $j = 1, \dots, n$, so ergibt sich mit der Notation der Differentialgleichung als ein System (s. Beispiel 4.8) die Fundamentalmatrix

$$\Phi(x) = \begin{pmatrix} u_1(x) & u_2(x) & \dots & u_n(x) \\ u'_1(x) & u'_2(x) & \dots & u'_n(x) \\ \vdots & \vdots & & \vdots \\ u_1^{(n-1)}(x) & u_2^{(n-1)}(x) & \dots & u_n^{(n-1)}(x) \end{pmatrix}, \quad x \in I.$$

Sei nun $\sum_{j=1}^n \rho_j u_j(x) = 0$ für alle $x \in I$. Dann gilt, wenn wir diese Gleichung k -mal differenzieren

für $k = 0, 1, \dots, n-1$, und $x = x_0$ einsetzen, $\sum_{j=1}^n \rho_j u_j^{(k)}(x_0) = 0$ für alle $k = 0, \dots, n-1$. Der

Vektor $(\rho_1, \dots, \rho_n)^\top$ löst also das homogene lineare Gleichungssystem $I_n \rho = 0$. Daher müssen alle $\rho_j = 0$ sein, d.h. wir haben ein Fundamentalsystem gefunden. Üblicherweise wird in diesem Fall die Menge der Funktionen $\{u_j : j = 1, \dots, n\}$ als Fundamentalsystem bezeichnet, ohne dass die Vektoren bestehend aus den Funktionen und ihren Ableitungen aufgelistet werden. $\square$

Um zu prüfen, ob einem homogenen System ein Fundamentalsystem vorliegt, bietet sich die Berechnung der *Wronskideterminante* an, wie der folgende Satz belegt.

Satz 4.13 Seien $\{u^1, \dots, u^n\}$ Lösungen eines homogenen, linearen Differentialgleichungssystems 1-ter Ordnung,

$$u'(x) = A(x) u(x), \quad x \in I \subseteq \mathbb{R}$$

und $W(x) := \det(u^1(x) \dots u^n(x))$. Dann ist entweder $W(x) = 0$ für alle $x \in I$, oder $W(x) \neq 0$ für alle $x \in \mathbb{R}$. Die Funktionen bilden genau dann ein Fundamentalsystem, wenn $W(x) \neq 0$ für ein (und damit alle) $x \in I$ ist. Die Determinante $W(x)$ wird **Wronskideterminante** genannt.

Beweis: Sei $W(x_1) = 0$ für ein $x_1 \in I$. Wir zeigen, dass dann auch $W(x) = 0$ für alle $x \in I$ gelten muss: Ist $W(x_1) = 0$, so heißt dies, dass die Vektoren $\{u^1(x_1), \dots, u^n(x_1)\}$ linear abhängig sind.

Es gibt also Zahlen $\alpha_1, \dots, \alpha_n$, nicht alle Null, mit $\sum_{k=1}^n \alpha_k u^k(x_1) = 0$.

Die Funktion $u(x) := \sum_{k=1}^n \alpha_k u^k(x)$ löst die Differentialgleichung und verschwindet an der Stelle $x = x_1$. Nach dem Satz 4.10 muss u überall verschwinden, d.h. $u(x) = 0$ für alle $x \in I$, da $u = 0$ die einzige Lösung ist. Dies wiederum bedeutet, dass die Vektoren $\{u^1(x), \dots, u^n(x)\}$ für alle $x \in I$ linear abhängig sind und somit $W(x) = 0$ für alle $x \in I$ ist. $\square$

Wir können den Satz 4.11 speziell auf die linearen Differentialgleichungen mit konstanten Koeffizienten, die wir in Abschnitt 4.1 behandelt haben, anwenden. Mit der folgenden Aussage wird so deutlich, dass wir mit dem Vorgehen aus dem vorletzten Abschnitt wirklich den gesamten Lösungsraum zur homogenen Differentialgleichung bei konstanten Koeffizienten ermitteln können.

Satz 4.14 *Es seien $\lambda_1, \dots, \lambda_m \in \mathbb{C}$ ($m \leq n$) die paarweise verschiedenen Nullstellen des charakteristischen Polynoms $p(\lambda) = a_0 + a_1\lambda + \dots + a_n\lambda^n$. Weiter sei $k_j \in \mathbb{N}$ die Vielfachheit der Nullstelle λ_j , $j = 1, \dots, m$, d.h. $p^{(l)}(\lambda_j) = 0$ für $l = 0, \dots, k_j - 1$ und $p^{(k_j)}(\lambda_j) \neq 0$ (beachte: $\sum_{j=1}^m k_j = n$). Die Funktionen*

$$\begin{aligned} & \{ e^{\lambda_1 x}, x e^{\lambda_1 x}, \dots, x^{k_1-1} e^{\lambda_1 x} \} \cup \{ e^{\lambda_2 x}, x e^{\lambda_2 x}, \dots, x^{k_2-1} e^{\lambda_2 x} \} \cup \dots \cup \\ & \cup \{ e^{\lambda_m x}, x e^{\lambda_m x}, \dots, x^{k_m-1} e^{\lambda_m x} \} \end{aligned}$$

bilden ein Fundamentalsystem der Differentialgleichung

$$a_0 u(x) + \dots + a_n u^{(n)}(x) = 0, \quad x \in \mathbb{R}.$$

Beweis: Es ist nur noch zu zeigen, dass alle genannten Funktionen über einem beliebigen Intervall $I \subseteq \mathbb{R}$ linear unabhängig sind. Eine Linearkombination lässt sich schreiben als

$$u(x) = \sum_{j=1}^m \sum_{l=0}^{k_j-1} \alpha_{jl} x^l e^{\lambda_j x} = \sum_{j=1}^m p_j(x) e^{\lambda_j x}$$

mit Polynomen p_j vom Grad höchstens $k_j - 1$, $j = 1, \dots, m$. Zu zeigen ist, dass aus der Annahme $u(x) = 0$ für alle $x \in I$ folgt, dass alle p_j die Nullpolynome sind. Dies erledigen wir durch eine doppelte vollständige Induktion.

Induktionsanfang 1: Im Fall $m = 1$ lautet die Annahme

$$p_1(x) e^{\lambda_1 x} = 0 \quad \text{für alle } x \in I.$$

Da die Exponentialfunktion keine Nullstellen besitzt, folgt $p_1(x) = 0$ für alle $x \in I$, d.h. p_1 ist das Nullpolynom.

Induktionsschritt 1 ($m \rightarrow m + 1$): Wir führen für diesen Schritt eine zweite Induktion über den Grad l des Polynoms p_{m+1} durch.

Induktionsanfang 2: Für $l = 0$ ist $p_{m+1}(x) = \alpha_{m+1,0}$ eine Konstante. Es folgt

$$\alpha_{m+1,0} = - \sum_{j=1}^m p_j(x) e^{(\lambda_j - \lambda_{m+1})x}.$$

Wir differenzieren diese Gleichung,

$$0 = - \sum_{j=1}^m [(\lambda_j - \lambda_{m+1}) p_j(x) + p'_j(x)] e^{(\lambda_j - \lambda_{m+1})x} \quad \text{für alle } x \in I.$$

Nun können wir die Induktionsvoraussetzung für m anwenden und erhalten

$$(\lambda_j - \lambda_{m+1}) p_j(x) - p'_j(x) = 0 \quad \text{für alle } x \in I.$$

Ein Koeffizientenvergleich ergibt, dass p_j konstant ist. Da $\lambda_j \neq \lambda_{m+1}$, $j = 1, \dots, m$, folgt die Behauptung.

Induktionsschritt 2 ($l \rightarrow l+1$): Das Polynom p_{m+1} habe den Grad $l+1$. Wir differenzieren die Gleichung

$$\sum_{j=1}^{m+1} p_j(x) e^{\lambda_j x} = 0, \quad x \in I,$$

einmal und erhalten

$$\sum_{j=1}^{m+1} (p'_j(x) + \lambda_j p_j(x)) e^{\lambda_j x} = 0, \quad x \in I.$$

Nun subtrahieren wir das λ_{m+1} -fache der ersten Gleichung von der zweiten und erhalten

$$\sum_{j=1}^m (p'_j(x) + (\lambda_j - \lambda_{m+1}) p_j(x)) e^{\lambda_j x} + p'_{m+1}(x) e^{\lambda_{m+1} x} = 0, \quad x \in I.$$

Der Grad von p'_{m+1} ist l , daher können wir die Induktionsvoraussetzung für l anwenden. Insbesondere ist p'_{m+1} also das Nullpolynom und daher p_{m+1} eine Konstante. Dies ist aber genau der Fall aus dem Induktionsanfang 2, d.h. $p_{m+1} = 0$.

Die Induktion 2 liefert den Induktionsschritt der Induktion 1. Somit ist die Aussage des Satzes bewiesen. $\square$

Um auch bei anderen Differentialgleichungssystemen mit konstanten Koeffizienten ein Fundamentalsystem zu ermitteln, lassen sich Eigenwerte und ihre Eigenvektoren direkt nutzen; denn der Exponentialansatz bei Differentialgleichungen n -ter Ordnung lässt sich auch auf allgemeinere lineare Systeme mit konstanten Koeffizienten übertragen.

Sei $A \in \mathbb{C}^{n \times n}$ eine konstante komplexe Matrix. Wir untersuchen das homogene Differentialgleichungssystem

$$u'(x) = A u(x), \quad x \in \mathbb{R}$$

und wollen ein Fundamentalsystem konstruieren. Als Ansatz suchen wir Lösungen in der Form $u(x) = e^{\lambda x} v$ mit einer Zahl $\lambda \in \mathbb{C}$ und einem Vektor $v \in \mathbb{C}^n$. Dies führt auf die Gleichung $Av = \lambda v$. Also können wir eine Lösung u angeben, wenn wir einen Eigenwert und einen zugehörigen Eigenvektor der Matrix A bestimmen. Übrigens: Schreiben wir eine lineare Differentialgleichung n -ter Ordnung mit konstanten Koeffizienten als ein System (s. Beispiel 4.8), so fallen die beiden Definitionen des charakteristischen Polynoms zusammen (nachprüfen!).

Fazit: Ist $A \in \mathbb{C}^{n \times n}$ gegeben, so müssen wir das charakteristische Polynom $p(\lambda) = \det(A - \lambda E)$ aufstellen und dessen Nullstellen berechnen. Das sind dann die Eigenwerte von A . Wir wissen: Ist $\hat{\lambda}$ ein Eigenwert und $\hat{v} \in \mathbb{C}^n$ ein zugehöriger Eigenvektor, so ist $u(x) = \exp(\hat{\lambda}x) \hat{v}$ eine Lösung der Differentialgleichung $u'(x) = A u(x)$. Jetzt erkennen wir auch, weshalb wir komplexe Zahlen zulassen. Ein Polynom braucht im Reellen ja keine Nullstelle zu besitzen. Im Komplexen besagt der Fundamentalsatz der Algebra, dass jedes Polynom p vom Grad n genau n komplexe Nullstellen besitzt, wenn sie gemäß ihrer Vielfachheit gezählt werden.

Beispiele 4.15

(a) Sei zunächst $A = \frac{1}{3} \begin{pmatrix} 4 & -2 \\ 1 & 1 \end{pmatrix}$. Dann ist $p(\lambda) = \det \begin{pmatrix} 4/3 - \lambda & -2/3 \\ 1/3 & 1/3 - \lambda \end{pmatrix} = 2/3 - 5/3\lambda + \lambda^2$ mit Nullstellen $\lambda_1 = 1$ und $\lambda_2 = 2/3$. Einen zu $\lambda_1 = 1$ gehörigen Eigenvektor erhalten wir aus $(A - 1I)v = 0$, also etwa $v^1 = (2, 1)^\top$. Analog ist $v^2 = (1, 1)^\top$ Eigenvektor zu $\lambda_2 = 2/3$. Also sind

$$u^1(x) = e^x \begin{pmatrix} 2 \\ 1 \end{pmatrix} \quad \text{und} \quad u^2(x) = e^{2x/3} \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

Lösungen des Systems $u'(x) = Au(x)$. Da für $x = 0$ die beiden Vektoren linear unabhängig sind, bilden diese beiden Funktionen nach obigem Satz ein Fundamentalsystem mit Fundamentalmatrix

$$\Phi(x) = \begin{pmatrix} 2\exp(x) & \exp(2x/3) \\ \exp(x) & \exp(2x/3) \end{pmatrix} \in \mathbb{R}^{2 \times 2}, \quad x \in \mathbb{R}.$$

Hier haben wir also ein reelles Fundamentalsystem gefunden.

(b) Sei jetzt $A = \begin{pmatrix} 3 & 2 \\ -2 & 1 \end{pmatrix}$. Dann ist $p(\lambda) = (3 - \lambda)(1 - \lambda) + 4 = \lambda^2 - 4\lambda + 7$ mit Nullstellen $\lambda_{1,2} = 2 \pm \sqrt{3}i$. Aus $(A - \lambda_1 I)v = 0$ berechnen wir den Eigenvektor $v^1 = (-2, 1 - \sqrt{3}i)^\top$ zu $\lambda_1 = 2 + \sqrt{3}i$ und analog $v^2 = (-2, 1 + \sqrt{3}i)^\top$ zu $\lambda_2 = 2 - \sqrt{3}i$. Daher erhalten wir das komplexe Fundamentalsystem

$$u^1(x) = e^{(2+\sqrt{3}i)x} \begin{pmatrix} -2 \\ 1 - \sqrt{3}i \end{pmatrix}, \quad u^2(x) = e^{(2-\sqrt{3}i)x} \begin{pmatrix} -2 \\ 1 + \sqrt{3}i \end{pmatrix}, \quad x \in \mathbb{R}.$$

Da die Matrix A reell ist, sind Real- und Imaginärteil dieser komplexen Lösungen ebenfalls Lösungen von $u'(x) = Au(x)$, also erhalten wir ein reelles Fundamentalsystem durch:

$$\begin{aligned} \tilde{u}^1(x) &= \operatorname{Re} u^1(x) = e^{2x} \left[\cos(\sqrt{3}x) \begin{pmatrix} -2 \\ 1 \end{pmatrix} - \sin(\sqrt{3}x) \begin{pmatrix} 0 \\ -\sqrt{3} \end{pmatrix} \right] \\ &= e^{2x} \begin{pmatrix} -2\cos(\sqrt{3}x) \\ \cos(\sqrt{3}x) + \sqrt{3}\sin(\sqrt{3}x) \end{pmatrix}, \\ \tilde{u}^2(x) &= \operatorname{Im} u^1(x) = e^{2x} \left[\cos(\sqrt{3}x) \begin{pmatrix} 0 \\ -\sqrt{3} \end{pmatrix} + \sin(\sqrt{3}x) \begin{pmatrix} -2 \\ 1 \end{pmatrix} \right] \\ &= e^{2x} \begin{pmatrix} -2\sin(\sqrt{3}x) \\ \sin(\sqrt{3}x) - \sqrt{3}\cos(\sqrt{3}x) \end{pmatrix}. \end{aligned}$$

□

Allgemein können wir uns folgenden Satz merken:

Satz 4.16 Sei $A \in \mathbb{C}^{n \times n}$. Seien die Eigenwerte $\lambda_1, \dots, \lambda_n \in \mathbb{C}$ von A , d.h. die Nullstellen des charakteristischen Polynoms $p(\lambda) = \det(A - \lambda I)$, alle paarweise verschieden. Dann bilden die Funktionen $\{u^1, \dots, u^n\}$ ein Fundamentalsystem der Differentialgleichung $u'(x) = Au(x)$, $x \in \mathbb{R}$, wobei $u^j(x) = \exp(\lambda_j x)v^j$ und $v^j \neq 0$ ein Eigenvektor zu λ_j ist.

Beweis: Wir wissen, dass der Lösungsraum n -dimensional ist und dass die angegebenen Funktionen wirklich Lösungen der Differentialgleichung sind. Wir müssen also nur noch zeigen, dass

sie auch linear unabhängig sind, d.h. dass die Wronskideterminante $W(x) = \det \Phi(x)$ wenigstens für ein x ungleich Null ist. Für $x = 0$ ist $u^j(0) = v^j$, und nach Satz 2.29 sind die Eigenvektoren $\{v^1, \dots, v^n\}$ linear unabhängig. Daher ist $\Phi(0)$ invertierbar und somit gilt $W(0) = \det \Phi(0) \neq 0$. $\square$

Bemerkung: Wenn doppelte (mehrfache) Nullstellen auftreten, wir aber linear unabhängige Eigenvektoren zu diesem Eigenwert finden können, erhalten wir mit dem Ansatz entsprechend zwei (oder mehr) linear unabhängige Lösungen der Differentialgleichungen. Wenn dies nicht möglich ist, sind weitere Schritte notwendig, um linear unabhängige Lösungen zu finden. Darauf wollen wir hier nicht weiter eingehen, da uns bald mit der Laplacetransformation eine weitere Technik zum Lösen solcher Differentialgleichungssysteme zur Verfügung stehen wird.

Beispiel 4.17 (gekoppelter Schwingkreis)

Wir betrachten jetzt das folgende System zweiter Ordnung

$$\begin{aligned} u_1''(t) &= -a_{11}u_1(t) + a_{12}u_2(t), & t \in \mathbb{R}, \\ u_2''(t) &= a_{21}u_1(t) - a_{22}u_2(t), & t \in \mathbb{R}, \end{aligned}$$

wobei alle Koeffizienten a_{ij} positiv seien. Setzen wir $u_3 = u_1'$ und $u_4 = u_2'$, so erhalten wir ein System erster Ordnung von vier Gleichungen und vier Unbekannten. Dieses schreiben wir gleich in der Form

$$\begin{pmatrix} u_1'(t) \\ u_2'(t) \\ u_3'(t) \\ u_4'(t) \end{pmatrix} = \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ -a_{11} & a_{12} & 0 & 0 \\ a_{21} & -a_{22} & 0 & 0 \end{pmatrix} \begin{pmatrix} u_1(t) \\ u_2(t) \\ u_3(t) \\ u_4(t) \end{pmatrix}.$$

Die Determinante $p(\lambda) = \det(A - \lambda I)$ entwickeln wir nach der ersten Zeile

$$\begin{aligned} \det(A - \lambda I) &= \begin{vmatrix} -\lambda & 0 & 1 & 0 \\ 0 & -\lambda & 0 & 1 \\ -a_{11} & a_{12} & -\lambda & 0 \\ a_{21} & -a_{22} & 0 & -\lambda \end{vmatrix} \\ &= -\lambda \begin{vmatrix} -\lambda & 0 & 1 \\ a_{12} & -\lambda & 0 \\ -a_{22} & 0 & -\lambda \end{vmatrix} + \begin{vmatrix} 0 & -\lambda & 1 \\ -a_{11} & a_{12} & 0 \\ a_{21} & -a_{22} & -\lambda \end{vmatrix} \\ &= -\lambda[-\lambda^3 - \lambda a_{22}] + [\lambda^2 a_{11} + d] \quad \text{mit } d := a_{11}a_{22} - a_{12}a_{21} \\ &= \lambda^4 + (a_{11} + a_{22})\lambda^2 + d \\ &= \left[\lambda^2 + \frac{1}{2}(a_{11} + a_{22}) \right]^2 - \frac{1}{4}(a_{22} - a_{11})^2 - a_{12}a_{21}. \end{aligned}$$

Nullsetzen ergibt eine quadratische Gleichung in λ^2 , und wir erhalten

$$\lambda^2 = -\frac{1}{2}(a_{22} + a_{11}) \pm \sqrt{\frac{1}{4}(a_{22} - a_{11})^2 + a_{12}a_{21}}.$$

Die rechte Seite ist auf jeden Fall reell, kann aber je nach den Werten der a_{ij} positiv oder negativ sein.

Betrachten wir den Fall $a_{11} = a_{12} = a_{21} = a_{22} = 1$. Dann erhalten wir $\lambda_1 = \sqrt{2}i$, $\lambda_2 = -\sqrt{2}i$ und $\lambda_{3,4} = 0$. Weiter folgt aus dem Gleichungssystem $(A - \lambda_1 I)v^1 = 0$ ein Eigenvektor $v^1 =$

$(1, -1, \sqrt{2}i, -\sqrt{2}i)^\top$. Da $\lambda_2 = \overline{\lambda_1}$ gilt, muss $v^2 = \overline{v^1}$ Eigenvektor zu diesem Eigenwert sein. Beide lassen sich zu reellen Lösungen des Systems kombinieren und wir erhalten etwa $u = \operatorname{Re}(v^1 e^{\lambda_1 x}) = \begin{pmatrix} \cos(\sqrt{2}x), -\cos(\sqrt{2}x), -\sqrt{2}\sin(\sqrt{2}x), \sqrt{2}\sin(\sqrt{2}x) \end{pmatrix}^\top$. Analog für den Imaginärteil. Also ergeben sich die reellen Lösungen $u_1(x) = c_1 \cos(\sqrt{2}x) + c_2 \sin(\sqrt{2}x)$ und $u_2(x) = -u_1(x)$. Weitere Lösungen des Systems erhalten wir mit dem Eigenwert $\lambda = 0$. Es ergibt sich $u_1(x) = u_2(x) = 1$ als Lösung. Die vierte linear unabhängige Lösung des Systems ist ein wenig aufwendiger zu sehen, da die Nullstelle $\lambda = 0$ doppelt auftritt. Diese Situation haben wir nicht genauer untersucht. Es sei aber angemerkt, dass $u_1(x) = x = u_2(x)$ eine weitere unabhängige Lösung ist. Die allgemeine Lösung des Systems ergibt sich nun aus der Linearkombination all dieser Lösungen zu

$$\begin{aligned} u_1(x) &= a + bx + c \cos(\sqrt{2}x) + d \sin(\sqrt{2}x) \\ u_2(x) &= a + bx - c \cos(\sqrt{2}x) - d \sin(\sqrt{2}x) \end{aligned}$$

mit $a, b, c, d \in \mathbb{R}$. □

Analog zu den linearen Gleichungssystemen und dem Kern einer Matrix besteht bei linearen **inhomogenen** Differentialgleichungssystemen ein struktureller Zusammenhang zwischen Lösungen und den Lösungen der zugehörigen homogenen Gleichung.

Satz 4.18 Sei $\hat{u}$ irgendeine („partikuläre“) Lösung des inhomogenen Differentialgleichungssystems $u'(x) = A(x)u(x) + b(x)$. Dann sind alle Lösungen des Systems von der Form $u = \hat{u} + v$, wobei v eine Lösung des homogenen Systems $v'(x) = A(x)v(x)$ ist.

Beweis: Dies ist leicht einzusehen. Denn sei $\hat{w}$ eine weitere Lösung des inhomogenen Systems, so folgt

$$(\hat{u} - \hat{w})'(x) = A(x)\hat{u}(x) + b(x) - A(x)\hat{w}(x) - b(x) = A(x)(\hat{u} - \hat{w})(x)$$

Also ist $v = \hat{u} - \hat{w}$ Lösung des homogenen Systems. □

4.4 Bestimmung partikulärer Lösungen

Mit dem Satz 4.18 können wir das Lösen von linearen Differentialgleichungssystemen in zwei Schritte aufteilen. Zum einen ist ein Fundamentalsystem zur homogenen Gleichung gesucht und andererseits wird irgendeine partikuläre Lösung benötigt. Methoden, um unabhängige Lösungen zu homogenen Gleichungen zu berechnen, haben wir im letzten Abschnitt gesehen. Daher beschäftigen wir uns in diesem Teil mit Wegen zur Bestimmung einer partikulären Lösung bei inhomogenen Differentialgleichungen.

Wir betrachten lineare inhomogene Differentialgleichungen, d.h.

$$a_n(x)u^{(n)}(x) + a_{n-1}(x)u^{(n-1)}(x) + \cdots + a_1(x)u'(x) + a_0(x)u(x) = b(x), \quad x \in I,$$

mit stetigen Funktionen a_j und b und $a_n(x) \neq 0$ für alle $x \in I$. Aus der allgemeinen Theorie (Satz 4.18) wissen wir, dass Lösungen der inhomogenen Differentialgleichung aus der Summe einer partikulären Lösung der inhomogenen Gleichung und der allgemeinen Lösung der homogenen Gleichung bestehen. Ziel ist es also, neben dem Fundamentalsystem der zugehörigen homogenen Differentialgleichung eine partikuläre Lösung zu finden.

Im Fall von Differentialgleichungen mit konstanten Koeffizienten ist häufig ein **Ansatz vom Typ der rechten Seite** erfolgreich.

Zunächst machen wir die folgende Beobachtung: Ist $b(x) = b_1(x) + b_2(x)$, die Summe von zwei Funktionen, und können wir für die rechten Seiten b_1, b_2 partikuläre Lösungen $\hat{u}_1$ und $\hat{u}_2$ finden, so ist $\hat{u} = \hat{u}_1 + \hat{u}_2$ eine partikuläre Lösung zur Differentialgleichung mit Inhomogenität b .

Sei $\boxed{b(x) = q(x) \exp(\mu x)}$ für ein $\mu \in \mathbb{C}$ und ein Polynom q vom Grad $m \in \mathbb{N}_0$. Dann wählen wir den Ansatz

$$u(x) = r(x)x^k e^{\mu x}, \quad x \in \mathbb{R},$$

wobei r als Polynom vom Grad wie q zuzulassen ist. Entscheidend ist nun, zwei Fälle zu unterscheiden:

- (a) Wenn $\mu \in \mathbb{C}$ keine Nullstelle des charakteristischen Polynoms $p(\lambda) = \sum_{j=0}^n a_j \lambda^j$ ist, so setzen wir $k = 0$.
- (b) Wenn $\mu \in \mathbb{C}$ Nullstelle des charakteristischen Polynoms ist, so wähle für k die Vielfachheit dieser Nullstelle. In diesem Fall liegt Resonanz vor. Allgemeiner wird die Situation, dass die anregende Frequenz μ und die Eigenfrequenzen des Systems zusammenfallen, d.h. $\operatorname{Im} \mu = \operatorname{Im} \lambda_1 = -\operatorname{Im} \lambda_2$, **Resonanzfall** genannt.

Bemerkung: Wir haben nicht ausgeschlossen, dass $r(x)$ und μ komplexwertig sind. Zerlegen wir alle Terme in Real- und Imaginärteil, so erhalten wir mit $\mu = \mu_1 + i\mu_2$ und $q(x) = q_1(x) + i q_2(x)$ ($\mu_1, \mu_2, q_1(x), q_2(x) \in \mathbb{R}$)

$$b(x) = e^{\mu_1 x} \left[(q_1(x) \cos(\mu_2 x) - q_2(x) \sin(\mu_2 x)) + i(q_2(x) \cos(\mu_2 x) + q_1(x) \sin(\mu_2 x)) \right]$$

und der Ansatz lautet

$$u(x) = x^k e^{\mu_1 x} \left[(r_1(x) \cos(\mu_2 x) - r_2(x) \sin(\mu_2 x)) + i(r_2(x) \cos(\mu_2 x) + r_1(x) \sin(\mu_2 x)) \right].$$

Zusammenfassend bedeutet dies: Für reelle lineare Differentialgleichungen mit konstanten Koeffizienten und Inhomogenität von der Form

$$b(x) = e^{\mu_1 x} \left[q_1(x) \cos(\mu_2 x) + q_2(x) \sin(\mu_2 x) \right]$$

mit Polynomen q_1, q_2 empfiehlt sich der allgemeine Ansatz

$$u(x) = x^k e^{\mu_1 x} \left[r_1(x) \cos(\mu_2 x) + r_2(x) \sin(\mu_2 x) \right].$$

Dabei ist k durch die Vielfachheit von $\mu_1 + i\mu_2$ als Nullstelle des charakteristischen Polynoms gegeben und die Polynome r_1, r_2 sind zu bestimmen.

Beispiel 4.19 Betrachte die Differentialgleichung

$$u''(x) + u(x) = \cos \omega x.$$

Also ist $\mu_1 = 0$, $\mu_2 = \omega$, $q_1(x) = 1$ und $q_2(x) = 0$ im allgemeinen Ansatz zu setzen.

1.Fall: Wenn $\omega \neq \pm 1$, d.h. $\mu = \pm i\omega$ nicht Nullstelle des charakteristischen Polynoms $p(\lambda) = \lambda^2 + 1$ ist, so erhalten wir aus dem Ansatz $u(x) = a \cos \omega x + b \sin \omega x$ die Identität

$$u''(x) + u(x) = a(1 - \omega^2) \cos \omega x + b(1 - \omega^2) \sin \omega x.$$

Durch Koeffizientenvergleich sehen wir, dass u eine Lösung der inhomogenen Differentialgleichung ist, wenn wir $a = \frac{1}{1-\omega^2}$ und $b = 0$ wählen.

2.Fall: Wenn etwa $\omega = 1$ ist, so müssen wir den Ansatz modifizieren zu $u(x) = ax \cos \omega x + bx \sin \omega x$, da $\lambda = i$ die Vielfachheit 1 als Nullstelle von $p(\lambda) = (\lambda + i)(\lambda - i)$ besitzt. Nun erhalten wir

$$u''(x) + u(x) = -2a \sin x + 2b \cos x.$$

Mit $a = 0$ und $b = \frac{1}{2}$ ist $u(x) = \frac{1}{2}x \sin x$ eine partikuläre Lösung. Beachte den wichtigen qualitativen Unterschied des linearen Anstiegs der Amplitude bei wachsendem x im Resonanzfall. $\square$

Eine weitere und allgemeinere Methode zur Lösung inhomogener linearer Differentialgleichungen ist die **Methode der Variation der Konstanten**. Bei dieser Methode benötigen wir ein Fundamentalsystem $\{u_1, \dots, u_n\}$ der homogenen Differentialgleichung. Damit machen wir den Ansatz

$$u(x) = \sum_{j=1}^n v_j(x) u_j(x), \quad x \in I$$

mit unbekannten Funktionen v_j (keine Konstanten, deshalb der Name). Ableiten ergibt

$$u'(x) = \sum_{j=1}^n v_j'(x) u_j(x) + \sum_{j=1}^n v_j(x) u_j'(x).$$

Nun überschlagen wir, was gegeben ist und was wir suchen. Wir haben eine Differentialgleichung aber n unbekannte Funktionen v_j . Da eine Gleichung nur eine Funktion eindeutig festlegen kann, sind vermutlich noch weitere $n - 1$ Bedingungen bei dem Ansatz erforderlich. Dies gibt uns Freiheiten bei der Wahl der Funktionen. Diese nutzen wir, um es uns einfacher zu machen, indem wir setzen

$$\sum_{j=1}^n v_j'(x) u_j(x) = 0, \quad \text{dann ist} \quad u'(x) = \sum_{j=1}^n v_j(x) u_j'(x)$$

und für die zweite Ableitung ergibt sich

$$u''(x) = \sum_{j=1}^n v_j'(x) u_j'(x) + \sum_{j=1}^n v_j(x) u_j''(x).$$

Die nächste Forderung $\sum_{j=1}^n v_j'(x) u_j'(x) = 0$ vereinfacht wiederum den Ausdruck. Setzen wir diese Prozedur fort, so ergeben sich $(n - 1)$ Forderungen

$$\sum_{j=1}^n v_j'(x) u_j^{(k)}(x) = 0 \quad \text{für } k = 0, \dots, n - 2$$

und es ist

$$u^{(n-1)}(x) = \sum_{j=1}^n v_j(x) u_j^{(n-1)}(x).$$

Mit $u^{(n)}(x) = \sum_{j=1}^n v'_j(x) u_j^{(n-1)}(x) + \sum_{j=1}^n v_j(x) u_j^{(n)}(x)$ ergibt sich

$$\begin{aligned} \sum_{k=0}^n a_k(x) u^{(k)}(x) &= a_n(x) u^{(n)}(x) + \sum_{j=1}^n v_j(x) \sum_{k=0}^{n-1} a_k(x) u_j^{(k)}(x) \\ &= \sum_{j=1}^n v_j(x) \underbrace{\sum_{k=0}^n a_k(x) u_j^{(k)}(x)}_{=0} + a_n(x) \sum_{j=1}^n v'_j(x) u_j^{(n-1)}(x). \end{aligned}$$

Dieser Ausdruck ist $b(x)$, wenn wir als letzte Forderung

$$\sum_{j=1}^n v'_j(x) u_j^{(n-1)}(x) = \frac{b(x)}{a_n(x)}$$

stellen. Zusammen liefert dies n Bedingungen für die n Funktionen $v'_1, \dots, v'_n$. Das lineare Gleichungssystem ist zu lösen und die Funktionen $v_1, \dots, v_n$ müssen noch durch Integration bestimmt werden. An einem Beispiel illustrieren wir dieses Vorgehen.

Beispiel 4.20 Wir betrachten die Differentialgleichung

$$x^2 u''(x) - 2u(x) = x, \quad x > 0.$$

Die homogene Gleichung $x^2 u''(x) - 2u(x) = 0$ (Euler-DGL) hat die allgemeine Lösung

$$u_{\text{hom}}(x) = c_1 x^2 + c_2 \frac{1}{x}.$$

Variation der Konstanten, d.h. der Ansatz $u(x) = v_1(x)x^2 + v_2(x)x^{-1}$ führt auf

$$u'(x) = \underbrace{v'_1(x)x^2 + v'_2(x)x^{-1}}_{=0 \text{ (1. Forderung)}} + 2v_1(x)x - v_2(x)x^{-2}.$$

Damit u die Differentialgleichung erfüllt, erhalten wir mit

$$u''(x) = 2v'_1(x)x - v'_2(x)x^{-2} + 2v_1(x) + 2v_2(x)x^{-3}$$

zusammen die beiden Bedingungen

$$\begin{aligned} v'_1(x)x^2 + v'_2(x)x^{-1} &= 0 \\ 2v'_1(x)x^3 - v'_2(x) &= x. \end{aligned}$$

Setzen wir die zweite Gleichung in die erste ein, so ergibt sich $v'_1(x) = \frac{1}{3x^2}$. Also ist

$$v_1(x) = -\frac{1}{3x},$$

wobei wir die Integrationskonstante gleich Null wählen können, da wir nur an einer partikulären Lösung interessiert sind. Weiter erhalten wir aus $v'_2(x) = -v'_1(x)x^3 = -\frac{1}{3}x$ eine Stammfunktion

$$v_2(x) = -\frac{1}{6}x^2.$$

Einsetzen von v_1 und v_2 in den Ansatz liefert eine partikuläre Lösung der inhomogenen Differentialgleichung,

$$u(x) = v_1(x)x^2 + v_2(x)\frac{1}{x} = -\frac{1}{2}x.$$

Also ist $u(x) = -\frac{1}{2}x + c_1 x^2 + c_2 x^{-1}$ die allgemeine Lösung dieser Differentialgleichung. $\square$

4.5 Die Potenzreihenmethode

Nachdem wir nun Techniken zum Lösen von linearen Differentialgleichungen mit konstanten Koeffizienten kennengelernt haben, wenden wir uns der Situation mit nicht konstanten Koeffizienten zu. Einen speziellen Typ haben wir schon gesehen, die Eulerschen Differentialgleichungen, bei denen der Potenzansatz $u(x) = x^\lambda$ hilfreich ist.

Eine weitere Methode, die wir in einem Spezialfall schon angewendet haben, kann auch hier nützlich sein. Mit der Idee der **Reduktion der Ordnung** lässt sich manchmal bei homogenen Differentialgleichungen mit variablen Koeffizienten eine weitere Lösung finden, wenn eine Lösung u_1 schon bekannt ist. Wir machen das am folgenden Beispiel klar:

Beispiel 4.21

Betrachte $u''(x) + x u'(x) + u(x) = 0$, $x \in \mathbb{R}$. Eine Lösung ist $u_1(x) = \exp(-x^2/2)$ (nachprüfen!). Wir machen nun den Produktansatz $u_2(x) = c(x)u_1(x)$. Dann ist

$$u_2'(x) = [c'(x) - x c(x)] e^{-x^2/2}, \quad u_2''(x) = [c''(x) - c(x) - 2x c'(x) + x^2 c(x)] e^{-x^2/2},$$

und eingesetzt folgt

$$[c''(x) - 2x c'(x) + (x^2 - 1)c(x) + x c'(x) - x^2 c(x) + c(x)] e^{-x^2/2} = 0,$$

also ist $c''(x) - x c'(x) = 0$, da u_1 keine Nullstellen besitzt. Dies ist eine lineare Differentialgleichung erster Ordnung für $v = c'$. Die Differentialgleichung wird gelöst durch (herleiten!) $v(x) = c \exp(x^2/2)$. Daher ist

$$u_2(x) = e^{-x^2/2} \int_0^x e^{t^2/2} dt, \quad x \in \mathbb{R},$$

eine zweite Lösung der Differentialgleichung. Wir müssen noch die lineare Unabhängigkeit überprüfen. Sei

$$a_1 e^{-x^2/2} + a_2 e^{-x^2/2} \int_0^x e^{t^2/2} dt = 0 \quad \text{für alle } x \in \mathbb{R}$$

und nehme $a_2 \neq 0$ an. Dann folgt, dass das Integral gleich der Konstanten $-a_1/a_2$ ist, ein Widerspruch dazu, dass seine Ableitung v nicht die Nullfunktion ist. Also ist $a_2 = 0$ und dann sofort auch $a_1 = 0$. Somit bildet $\{u_1, u_2\}$ ein Fundamentalsystem der Differentialgleichung, auch wenn wir u_2 nicht geschlossen angeben können. $\square$

Das allgemeinste Vorgehen bei linearen Differentialgleichungen n -ter Ordnung, d.h.

$$a_0(x) u(x) + a_1(x) u'(x) + \cdots + a_n(x) u^{(n)}(x) = b(x), \quad x \in I \subseteq \mathbb{R},$$

wobei $a_n(x) \neq 0$ für alle $x \in I$ ist, ist die **Potenzreihenmethode**. Wir setzen voraus, dass alle Koeffizientenfunktionen a_j und die rechte Seite b durch Potenzreihen um $x_0 \in I$ entwickelbar sind, die in I konvergieren. Es sei also etwa $b(x) = \sum_{k=0}^{\infty} \beta_k (x - x_0)^k$, $x \in I$. Außerdem seien

Anfangsbedingungen $u^{(j)}(x_0) = u_j$, $j = 0, \dots, n-1$, am Entwicklungspunkt $x_0 \in I$ gegeben. Wir machen nun einen **Ansatz in der Form einer Potenzreihe**

$$u(x) = \sum_{k=0}^{\infty} c_k (x - x_0)^k, \quad x \in I$$

um den Entwicklungspunkt x_0 , an dem die Anfangswerte u festlegen. Wir wissen, dass die Potenzreihe im Innern des Konvergenzintervalls gliedweise differenziert werden darf. Es lässt sich zeigen, dass Einsetzen der Ableitungen von u und ein Koeffizientenvergleich zum Ziel führen. Wir wollen dies an zwei Beispielen vorführen:

Beispiele 4.22

(a) Betrachte die Differentialgleichung

$$u''(x) + x u'(x) + u(x) = 0.$$

Mit dem Ansatz $u(x) = \sum_{k=0}^{\infty} c_k x^k$ ist

$$x u'(x) = \sum_{k=1}^{\infty} c_k k x^k = \sum_{k=0}^{\infty} c_k k x^k$$

und

$$u''(x) = \sum_{k=2}^{\infty} c_k k(k-1) x^{k-2} = \sum_{k=0}^{\infty} c_{k+2} (k+2)(k+1) x^k.$$

Setzen wir dies in die Differentialgleichung ein, so folgt

$$\sum_{k=0}^{\infty} (c_{k+2} (k+2)(k+1) + c_k k + c_k) x^k = 0$$

Ein Koeffizientenvergleich ergibt

$$c_{k+2} (k+2)(k+1) + c_k (k+1) = 0, \quad \text{also} \quad c_{k+2} = -\frac{1}{k+2} c_k, \quad k = 0, 1, 2, \dots$$

Wir erhalten zwei linear unabhängige Lösungen u_1 und u_2 , die eine (nennen wir sie u_1) für $c_0 = 1$, $c_1 = 0$ und die andere u_2 für $c_0 = 0$, $c_1 = 1$. Dann besteht u_1 nur aus Monomen gerader Potenz und u_2 aus Monomen ungerader Potenz. Ersetzt man für u_1 den Index k durch $2j$, so folgt die Rekursionsformel

$$c_0 = 1, \quad c_{2j+2} = -\frac{1}{2(j+1)} c_{2j}, \quad j = 0, 1, 2, \dots$$

also (Induktion!) $c_{2j} = (-1/2)^j / j!$, $j = 0, 1, \dots$. Daher ist

$$u_1(x) = \sum_{j=0}^{\infty} \frac{1}{j!} \left(-\frac{x^2}{2} \right)^j = e^{-x^2/2}.$$

(b) Betrachte jetzt die Differentialgleichung zweiter Ordnung

$$u''(x) + \frac{1}{1+x^2} u(x) = 0, \quad |x| < 1,$$

mit den Anfangsbedingungen $u(0) = 1$ und $u'(0) = 0$. Statt $1/(1+x^2)$ in eine Potenzreihe zu entwickeln, ist es einfacher, die Differentialgleichung in der Form

$$(1+x^2) u''(x) + u(x) = 0$$

zu behandeln. Mit dem Ansatz $u(x) = \sum_{k=0}^{\infty} c_k x^k$ rechnen wir aus:

$$\begin{aligned} u''(x) &= \sum_{k=2}^{\infty} c_k k(k-1) x^{k-2} = \sum_{k=0}^{\infty} c_{k+2} (k+2)(k+1) x^k, \\ x^2 u''(x) &= \sum_{k=2}^{\infty} c_k k(k-1) x^k = \sum_{k=0}^{\infty} c_k k(k-1) x^k. \end{aligned}$$

Daher haben wir

$$\sum_{k=0}^{\infty} [c_{k+2}(k+2)(k+1) + c_k k(k-1) + c_k] x^k = 0.$$

Da dies für alle x gilt, müssen alle Koeffizienten Null sein, also ist

$$c_{k+2} = -\frac{k(k-1)+1}{(k+2)(k+1)} c_k, \quad k = 0, 1, 2, \dots$$

Kennen wir c_0 und c_1 , so sind durch diese Formel alle c_k bestimmt. Aus den beiden Anfangsbedingungen $1 = u(0) = c_0$ und $0 = u'(0) = c_1$ folgt, dass alle $c_k = 0$ sind für ungerade k . u hat also die Form $u(x) = \sum_{j=0}^{\infty} c_{2j} x^{2j}$. Genauer können (und brauchen) wir die Lösung nicht angeben. Wir müssen uns noch überlegen, für welche x die Reihe überhaupt konvergiert. Da die Koeffizienten rekursiv definiert sind, tun wir dies mit dem Quotientenkriterium:

$$\left| \frac{c_{2j+2} x^{2j+2}}{c_{2j} x^{2j}} \right| = x^2 \frac{2j(2j-1)+1}{(2j+1)(2j+2)} \rightarrow x^2, \quad j \rightarrow \infty.$$

Daher liegt Konvergenz für $|x| < 1$ vor. Die Funktion u hat also die merkwürdige Form

$$u(x) = 1 - \frac{1}{2!} x^2 + \frac{1+1 \cdot 2}{4!} x^4 - \frac{(1+1 \cdot 2)(1+3 \cdot 4)}{6!} x^6 + - \dots, \quad |x| < 1.$$

□

Wir betrachten jetzt noch eine weitere Differentialgleichung mit variablen Koeffizienten a_k , die **Besselsche Differentialgleichung**

$$x^2 u''(x) + x u'(x) + (x^2 - n^2) u(x) = 0 \quad \text{oder} \quad \frac{1}{x} (x u'(x))' + \left(1 - \frac{n^2}{x^2}\right) u(x) = 0.$$

Hier ist $n \in \mathbb{N}_0$ ein fester Parameter.

Zur Erinnerung: Eine Eulersche Differentialgleichung entsteht, wenn statt $x^2 - n^2$ der Koeffizient von u eine Konstante wäre. In diesem Fall würde der Ansatz $u(x) = x^\lambda$ zum Erfolg führen. Dies motiviert den **verallgemeinerten Potenzreihenansatz**

$$u(x) = x^\lambda \sum_{j=0}^{\infty} c_j x^j = \sum_{j=0}^{\infty} c_j x^{j+\lambda}, \quad x > 0,$$

für reelle oder komplexe noch zu bestimmende λ, c_j . Ohne Einschränkung setzen wir $c_0 \neq 0$ voraus (sonst könnten wir x ausklammern und in den x^λ -Term stecken). Mit diesem Ansatz ist

$$\begin{aligned} x u'(x) &= \sum_{j=0}^{\infty} c_j (j + \lambda) x^{j+\lambda}, \\ x^2 u''(x) &= \sum_{j=0}^{\infty} c_j (j + \lambda)(j + \lambda - 1) x^{j+\lambda}, \\ x^2 u(x) &= \sum_{j=0}^{\infty} c_j x^{j+\lambda+2} = \sum_{j=0}^{\infty} c_{j-2} x^{j+\lambda}, \end{aligned}$$

wenn wir für die letzte Summe $c_{-1} = c_{-2} = 0$ vereinbaren. Einsetzen in die Differentialgleichung liefert

$$\sum_{j=0}^{\infty} x^{j+\lambda} [c_j (j + \lambda)(j + \lambda - 1) + c_j (j + \lambda) + c_{j-2} - n^2 c_j] = 0, \quad x > 0.$$

Der Koeffizientenvergleich führt auf die Rekursionsformel

$$c_j [(j + \lambda)^2 - n^2] + c_{j-2} = 0, \quad j = 0, 1, 2, \dots$$

Für $j = 0$ bekommen wir unter Beachtung von $c_0 \neq 0$ und $c_{-2} = 0$ die Bedingung $\lambda^2 = n^2$, also $\lambda = \pm n$. Wir untersuchen beide Fälle:

1. Fall: $\lambda = n$. Damit vereinfacht sich die Rekursionsformel zu

$$c_j [j(j + 2n)] + c_{j-2} = 0, \quad \text{d.h.} \quad c_j = -\frac{1}{j(j + 2n)} c_{j-2}, \quad j = 1, 2, 3, \dots,$$

da der Nenner positiv ist. Für $j = 1$ ergibt sich $c_1 = 0$ wegen $c_{-1} = 0$. Weiter folgern wir, dass $c_j = 0$ für ungerades j und

$$c_{2k} = \frac{n! (-1)^k}{k! (k + n)!} \left(\frac{1}{2}\right)^{2k} c_0, \quad k = 0, 1, 2, \dots$$

gilt (vollständige Induktion!). Daher erhalten wir als Lösung die Potenzreihe

$$u(x) = n! c_0 2^n \sum_{k=0}^{\infty} \frac{(-1)^k}{k! (k + n)!} \left(\frac{x}{2}\right)^{2k+n}.$$

Die Konvergenz dieser Reihe haben wir schon in Beispiel 4.20 des ersten Semesters untersucht. Nach geeigneter Normierung, d.h. Wahl von c_0 , erhalten wir:

Definition 4.23 (*Besselfunktionen*)

Für $n \in \mathbb{N}_0$ heißt

$$J_n(x) = \sum_{k=0}^{\infty} \frac{(-1)^k}{k! (k + n)!} \left(\frac{x}{2}\right)^{2k+n}, \quad x \in \mathbb{R},$$

die **Besselsche Funktion** der Ordnung n . Diese Potenzreihe konvergiert für alle $x \in \mathbb{C}$.

2. Fall: $\lambda = -n$ und $n \in \mathbb{N}$. Dann vereinfacht sich die Rekursionsformel zu

$$c_j [j(j-2n)] + c_{j-2} = 0, \quad j = 1, 2, \dots$$

Wieder erhalten wir $c_1 = 0$ und hieraus $c_j = 0$ für alle ungeraden j , da für ungerade j der Ausdruck $j - 2n$ niemals Null wird. Für gerade $j = 2k$ ergibt sich die Rekursion

$$c_{2k} 4k(k-n) + c_{2(k-1)} = 0, \quad k = 1, 2, \dots$$

Da jedoch $c_0 \neq 0$, erhalten wir $c_{2k} \neq 0$ für $k = 0, \dots, n-1$. Für $k = n$ führt dies zu einem Widerspruch, denn die linke Seite der Rekursionsformel verschwindet für $k = n$, was $c_{2(n-1)} = 0$ ergeben würde. Daher liefert der Fall $\lambda = -n$ über diesen Ansatz keine weitere Lösung der Differentialgleichung. Eine zweite, von der Besselfunktion linear unabhängige Lösung der Besselschen Differentialgleichung ist die sogenannte Neumannfunktion, oder Besselfunktion zweiter Art, die bei $x = 0$ eine Singularität hat.

Wir wollen noch **einige Eigenschaften** der Besselfunktionen anmerken.

- Ist n gerade (bzw. ungerade), so ist J_n eine gerade (bzw. ungerade) Funktion, d.h. es treten nur gerade (bzw. ungerade) Potenzen von x auf, d.h. es gilt $J_n(-x) = J_n(x)$ (bzw. $J_n(-x) = -J_n(x)$) für alle $x \in \mathbb{R}$.
- Es ist $J_0(0) = 1$ und $J_n(0) = 0$ für $n \geq 1$.
- Die Funktionen J_n , $n = 0, 1, 2, 3$, sind in Abbildung 1 dargestellt.

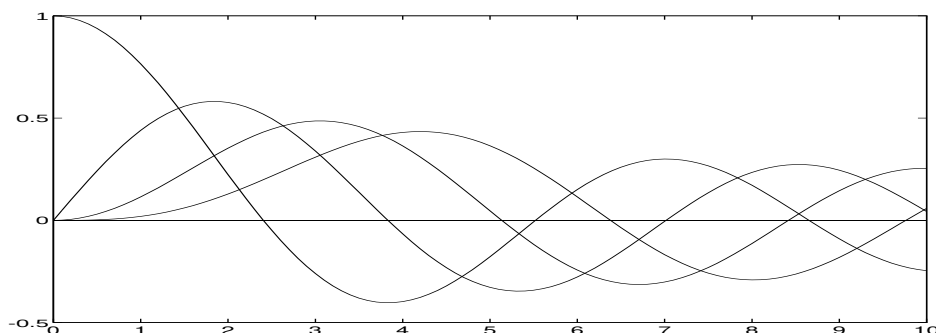


Abbildung 2: Die Besselfunktionen J_n , $n = 0, 1, 2, 3$

Satz 4.24 (Integraldarstellungen)

Für $n = 0, 1, 2, \dots$ gilt

$$\begin{aligned} J_n(x) &= \frac{1}{2\pi} \int_0^{2\pi} e^{i(x \sin t - nt)} dt, \quad x \in \mathbb{R} \quad (\text{komplexe Darstellung}), \\ &= \frac{1}{2\pi} \int_0^{2\pi} \cos(x \sin t - nt) dt, \quad x \in \mathbb{R} \quad (\text{reelle Darstellung}). \end{aligned}$$

Beweis: Wir betrachten $x \in \mathbb{R}$ als festen Parameter und die Funktionenreihe

$$e^{ix \sin t} = \sum_{m=0}^{\infty} \frac{1}{m!} (i \sin t)^m x^m.$$

Mit $e^{|x|} = \sum_{n=0}^{\infty} \frac{1}{n!} |x|^n$ ist eine integrierbare Majorante für alle t gegeben. Daher dürfen wir gliedweise integrieren und erhalten

$$\frac{1}{2\pi} \int_0^{2\pi} e^{ix \sin t} e^{-int} dt = \sum_{m=0}^{\infty} c_m x^m$$

mit

$$\begin{aligned} c_m &= \frac{1}{m! 2\pi} \int_0^{2\pi} (i \sin t)^m e^{-int} dt = \frac{1}{m! 2\pi} \int_0^{2\pi} \left(\frac{e^{it} - e^{-it}}{2} \right)^m e^{-int} dt \\ &= \frac{1}{m! 2\pi 2^m} \int_0^{2\pi} \sum_{j=0}^m \binom{m}{j} e^{i(m-j)t} (-1)^j e^{-int} e^{-ijt} dt \\ &= \frac{1}{2\pi 2^m} \sum_{j=0}^m \frac{(-1)^j}{j! (m-j)!} \int_0^{2\pi} e^{-i(n-m+2j)t} dt. \end{aligned}$$

Das Integral ist nur für $n - m + 2j = 0$ von 0 verschieden und hat dann den Wert 2π . Wegen $j \geq 0$ ist in diesem Fall $m \geq n$ und $m - n$ gerade. Schreiben wir $m = n + 2\mu$, so ist das Integral nur für $j = \mu$ von 0 verschieden und wir erhalten

$$c_{n+2\mu} = \frac{1}{2^{n+2\mu}} \frac{(-1)^\mu}{\mu! (n+\mu)!}, \quad \text{d.h.} \quad \frac{1}{2\pi} \int_0^{2\pi} e^{i(x \sin t - nt)} dt = \sum_{\mu=0}^{\infty} \frac{(-1)^\mu}{\mu! (n+\mu)!} \left(\frac{x}{2} \right)^{n+2\mu},$$

und dies war zu zeigen.

Die reelle Darstellung folgt einfach durch Gleichsetzen der Realteile der rechten und linken Seiten, da die linke Seite reell ist. $\square$

4.6 Numerische Lösungsverfahren

Auch wenn im allgemeinen keine geschlossene Lösung bestimmt werden kann, so gibt es doch generelle numerische Zugänge zur Anfangswertaufgabe,

$$u'(x) = f(x, u(x)), \quad x \in [x_0, x_0 + a], \quad u(x_0) = u^{(0)}.$$

Die Voraussetzungen des Satzes von Picard-Lindelöf seien erfüllt, d.h. es existiere eine eindeutig bestimmte Lösung u der (AWA). Es ist unser Ziel, an äquidistanten Stützstellen $x_k = x_0 + k h$, $k = 0, 1, 2, \dots, N$ (wobei $h \leq a/N$ ist) Näherungswerte $u_k \in \mathbb{R}$ an die Werte $u(x_k)$ der exakten Lösung u zu bestimmen. Zwei Möglichkeiten seien kurz dargestellt, wobei wir u skalar- und reellwertig annehmen. Die Methoden lassen sich aber auf Systeme von Differentialgleichungen übertragen. Sei also u im Folgenden immer die exakte Lösung der (AWA). Nach dem Mittelwertsatz gilt

$$u(x+h) = u(x) + u'(z) h \quad \text{für ein } z \in (x, x+h).$$

Wir approximieren nun die Ableitung $u'(z) \approx u'(x)$ und erhalten mit der Differentialgleichung

$$u(x+h) \approx u(x) + u'(x)h = u(x) + f(x, u(x))h.$$

Damit erhalten wir ein Verfahren, **das Eulersche Polygonzugverfahren**, zur Näherung der Lösung u , indem wir ausgehend vom Anfangswert $u_0 = u(x_0)$ die Folge

$$u_{k+1} := u_k + f(x_k, u_k)h \approx u(x_{k+1})$$

definieren. Die u_k sind dann Näherungen an die Funktionswerte $u(x_k)$.

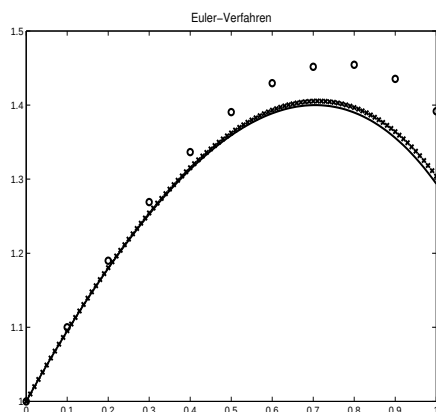
Dies ist ein einfaches numerisches Verfahren, d.h. einfach zu verstehen. Es ist jedoch numerisch sehr langsam und wird deswegen üblicherweise durch bessere Verfahren ersetzt.

Beispiel 4.25

(a) Gegeben sei die Anfangswertaufgabe

$$v'(x) = (1 - x^2) - \frac{x}{v(x)}, \quad v(0) = 1.$$

Dieses Problem lässt sich nicht geschlossen lösen. Wir sind also auf ein numerisches Verfahren angewiesen, um eine Vorstellung von der Lösung zu bekommen. Das Bild zeigt die Ergebnisse des Euler-Verfahrens mit $h = 0.1$ (o) und $h = 0.01$ (x) verglichen mit einer sehr viel genaueren Approximation der Lösung (durchgezogene Linie)



(b) Die Differentialgleichung $u'(x) = u(x)$, $x \geq 0$, $u(0) = 1$, wird durch $u(x) = \exp(x)$ gelöst. Das Eulerverfahren liefert $u_{k+1} = u_k + hu_k = (1+h)u_k$, also $u_k = (1+h)^k$ für $k = 0, 1, \dots$. Dieses u_k ist wirklich eine Näherung an $u(x_k) = u(kh) = \exp(kh)$. Dies sehen wir, indem wir $x > 0$ festhalten und $h = x/N$ für $N \in \mathbb{N}$, sowie $x_k^{(N)} = kh^{(N)}$, $k = 0, 1, \dots, N$ setzen. Dann ist $x_N^{(N)} = x$, und für dieses Beispiel ist

$$u_N = [1+h]^N = \left[1 + \frac{x}{N}\right]^N,$$

In der Tat konvergiert u_N gegen $u(x_N^{(N)}) = u(x) = \exp(x)$ für $N \rightarrow \infty$. □

Allgemein können wir für das Euler-Verfahren folgende Konvergenzaussage zeigen.

Satz 4.26 Sei $I = [x_0, x_0 + a]$, $f : I \times D \rightarrow \mathbb{R}$ stetig, lipschitzstetig auf D , d.h. $|f(x, u) - f(x, v)| \leq L|u - v|$, und die Funktionen $x \mapsto f(x, v(x))$ stetig differenzierbar für alle $v \in C^1(I)$ mit $v(I) \subset D$. Dann gilt für das Eulersche Polygonzugverfahren mit $x_k = x_0 + k \frac{a}{N}$, $k = 0, 1, \dots, N$:

$$|u_k - u(x_k)| \leq \frac{aM}{2LN} \left(e^{aL} - 1 \right), \quad k = 0, \dots, N,$$

wobei $M = \max\{|u''(x)| : x \in I\}$ ist. Natürlich müssten wir genauer $x_k^{(N)}$ und $u_k^{(N)}$ schreiben, um die Abhängigkeit von N anzudeuten, und es ist

$$\max_{k=0, \dots, N} |u_k^{(N)} - u(x_k^{(N)})| \leq \frac{c}{N}$$

für eine von N unabhängige Konstante $c > 0$. In diesem Sinn konvergiert für $N \rightarrow \infty$ der Fehler $|u_k^{(N)} - u(x_k^{(N)})|$ so schnell gegen 0 wie $1/N$. Aus diesem Grund nennt man es ein Verfahren von **erster Ordnung**.

Beweis: Wegen der Voraussetzung an f können wir die Differentialgleichung differenzieren und erhalten, dass die Lösung u der (AWA) nicht nur einmal, sondern sogar zweimal stetig differenzierbar ist. Also ist die Definition der Konstante M und die Voraussetzung an u'' sinnvoll.

Wir schätzen den Fehler $|u(x_k) - u_k|$ induktiv ab und schreiben zunächst mit $h = a/N$:

$$|u(x_{k+1}) - u_{k+1}| \leq \left| [u(x_{k+1}) - u(x_k)] + [u(x_k) - u_k] - hf(x_k, u_k) \right|.$$

Den ersten Teil schätzen wir mit der Taylorformel ab:

$$u(x_{k+1}) - u(x_k) = hu'(x_k) + \frac{h^2}{2} u''(z_k) = hf(x_k, u(x_k)) + \frac{h^2}{2} u''(z_k),$$

mit $z_k \in (x_k, x_{k+1})$. Weiter folgt mit der Dreiecksungleichung und der Lipschitzbedingung,

$$\begin{aligned} |u(x_{k+1}) - u_{k+1}| &\leq |u(x_k) - u_k| + h |f(x_k, u(x_k)) - f(x_k, u_k)| + \frac{h^2 M}{2} \\ &\leq (1 + hL) |u(x_k) - u_k| + \frac{1}{2} h^2 M. \end{aligned}$$

Dies gilt für alle $k = 0, 1, \dots$. Vollständige Induktion ergibt die explizite Darstellung,

$$|u(x_k) - u_k| \leq \frac{1}{2} h^2 M \sum_{\ell=0}^{k-1} (1 + hL)^\ell, \quad k = 1, 2, \dots$$

Mit der geometrischen Summenformel folgt

$$|u(x_k) - u_k| \leq \frac{1}{2} h^2 M \frac{(1 + hL)^k - 1}{hL} = \frac{aM}{2LN} \left[\left(1 + \frac{aL}{N} \right)^k - 1 \right] \leq \frac{aM}{2LN} \left(e^{aL} - 1 \right).$$

Damit ist der Satz bewiesen. □

Wir wollen noch ein in der Praxis häufig verwendetes Verfahren, das **Runge-Kutta-Verfahren**, beschreiben, das erheblich schnellere Konvergenz erlaubt. Wir gehen wieder vom Mittelwertsatz

aus und approximieren

$$\begin{aligned}
u(x_k + h) &= u(x_k) + u'(z) h \approx u(x_k) + \frac{h}{6} \left(u'(x_k) + 4u'(x_k + \frac{h}{2}) + u'(x_k + h) \right) \\
&= u(x_k) + \frac{h}{6} \left(\underbrace{u'(x_k)}_{f(x_k, u(x_k))} + 2 \underbrace{u'(x_k + \frac{h}{2})}_{f(x_k + \frac{h}{2}, u(x_k + \frac{h}{2}))} + 2 \underbrace{u'(x_k + \frac{h}{2})}_{f(x_k + \frac{h}{2}, u(x_k + \frac{h}{2}))} + \underbrace{u'(x_k + h)}_{f(x_k + h, u(x_k + h))} \right) \\
&\approx u(x_k) + \frac{h}{6} (K_1 + 2K_2 + 2K_3 + K_4) =: u_{k+1}
\end{aligned}$$

Die Wahl von

$$K_1 := f(x_k, u_k)$$

ist einfach, da die Ableitung $u'(x_k)$ wie beim Euler-Verfahren durch die Differentialgleichung gegeben ist, wenn wir eine Approximation u_k bestimmt haben. Für K_2 muss $u'(x_k + \frac{h}{2}) = f(x_k + \frac{h}{2}, u(x_k + \frac{h}{2}))$ approximiert werden. Die Idee ist es, den unbekannten Funktionswert $u(x_k + \frac{h}{2})$ durch einen Schritt des Euler-Verfahrens mit Schrittweite $\frac{h}{2}$ anzunähern. Dann ergibt sich

$$K_2 = f\left(x_k + \frac{h}{2}, u_k + \frac{h}{2}K_1\right)$$

Für K_3 wird K_2 anstelle von K_1 als „Prädiktor“ gewählt für den unbekannten Funktionswert $u(x_k + \frac{h}{2})$ in $f(x_k + \frac{h}{2}, u(x_k + \frac{h}{2}))$, sodass

$$K_3 = f(x_k + \frac{h}{2}, u_k + \frac{h}{2}K_2)$$

gesetzt wird. Schließlich wird für K_4 wieder ein Euler-Schritt mit Schrittweite h und dem Prädiktor K_3 für die Ableitung festgelegt, d.h.

$$K_4 = f(x_k + h, u_k + h K_3).$$

Es lässt sich zeigen (nicht mehr ganz so einfach), dass

$$|u_k^{(N)} - u(x_k^{(N)})| \leq c \frac{1}{N^4}, \quad j = 0, \dots, N,$$

gilt, wenn f viermal stetig differenzierbar ist

Beispiel 4.27

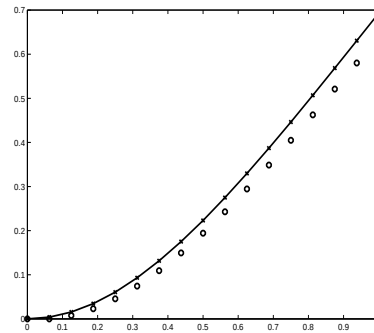
Wir betrachten die (AWA)

$$u'(x) = \frac{2x}{(1+x^2)^2} e^{-u(x)}, \quad x \geq 0, \quad u(0) = 0.$$

Dies ist eine Differentialgleichung mit getrennten Variablen und wird gelöst durch $u(x) = \ln[2 - 1/(1+x^2)]$, $x \geq 0$ (bitte herleiten und nachrechnen).

In der folgenden Tabelle sind die Fehler $\max_{j=0, \dots, n} |u(x_j) - u_k|$ für das Euler- und das Runge-Kutta-Verfahren an der Stelle $x = 1$ und verschiedene Werte von n angegeben. Das Bild zeigt die beiden Approximationen (Euler-Verfahren $\circ$, Runge-Kutta-Verfahren $\times$) für $n = 16$.

n	Euler	Runge–Kutta
2	$.18E + 00$	$.13E - 02$
4	$.72E - 01$	$.56E - 04$
8	$.32E - 01$	$.31E - 05$
16	$.16E - 01$	$.18E - 06$
32	$.76E - 02$	$.11E - 07$
64	$.37E - 02$	$.68E - 09$
128	$.19E - 02$	$.43E - 10$
256	$.93E - 03$	$.27E - 11$



Hier erkennen wir deutlich die Überlegenheit des Runge–Kutta-Verfahrens. Beachte, dass sich bei Verdoppelung von N die Fehler jeweils halbieren (Euler-Verfahren) bzw. auf ein Sechzehntel gehen (Runge-Kutta-Verfahren). $\square$

5 Die Laplacetransformation

5.1 Parameterintegrale

Die Gammafunktion ist für $t > 0$ definiert durch

$$\Gamma(t) = \int_0^\infty e^{-x} x^{t-1} dx.$$

Es handelt sich also um ein Integral, das von einem Parameter t abhängt. Bezüglich des Parameters können wir den Ausdruck wieder als Funktion auffassen, was oft nützlich sein kann, wie wir später in diesem Abschnitt sehen werden, wenn wir die Laplacetransformationen untersuchen. Um Eigenschaften dieser Funktion wie Stetigkeit oder Differenzierbarkeit zu untersuchen, müssen Grenzprozesse vertauscht werden. Im allgemeinen ist dies **nicht** erlaubt, wie wir aus dem folgenden einfachen Beispiel sehen:

Beispiel 5.1 Für $k \in \mathbb{N}$ seien die Funktionen $f_k : \mathbb{R} \rightarrow \mathbb{R}$ definiert durch

$$f_k(x) := \begin{cases} x + 1 - k, & k - 1 \leq x < k, \\ k + 1 - x, & k \leq x < k + 1, \\ 0, & x \notin [k - 1, k + 1). \end{cases}$$

(Die Funktion f_k ist ein „Hut“ mit Spitze bei $x = k$ und 0 für $|x - k| \geq 1$.) Die Folge (f_k) konvergiert punktweise gegen die Nullfunktion, d.h. $\lim_{k \rightarrow \infty} f_k(x) = 0$ für jedes $x \in \mathbb{R}$, aber $\int_{-\infty}^{\infty} f_k(x) dx = 1$ für alle $k \in \mathbb{N}$. □

Unter welchen Voraussetzungen das Vertauschen erlaubt ist, besagt der folgende wichtige Satz.

Satz 5.2 (Lebesguescher Konvergenzsatz) Sei $I \subset \mathbb{R}$ ein Intervall und $f_n \in L(I)$ eine Folge von Funktionen, die punktweise fast überall gegen $f : I \rightarrow \mathbb{R}$ konvergiert. Wenn eine Funktion $g \in L(I)$ existiert mit $|f_n(x)| \leq g(x)$ fast überall, dann ist $f \in L(I)$ und

$$\lim_{n \rightarrow \infty} \int_I f_n(x) dx = \int_I f(x) dx.$$

Für diesen Beweis müssen wir auf die Literatur verweisen. Aber zwei wichtige Konsequenzen aus dem Satz wollen wir herausstellen.

Satz 5.3 Sei auf einem Intervall $I \subseteq \mathbb{R}$ und einem kompakten Intervall $[a, b] \subseteq \mathbb{R}$ eine Funktion $f : [a, b] \times I \rightarrow \mathbb{R}$ gegeben. Für jedes feste $t \in [a, b]$ sei $f(t, \cdot)$ Lebesgue-integrierbar, d.h. $f(t, \cdot) \in L(I)$. Dann ist durch

$$G(t) := \int_I f(t, x) dx, \quad t \in [a, b]$$

eine Funktion auf $[a, b]$ definiert.

- (a) G ist **stetig** auf $[a, b]$, wenn $f(\cdot, x)$ stetige Funktionen sind für fast alle $x \in I$, und es eine Funktion $g \in L(I)$ gibt, sodass für alle $t \in [a, b]$ gilt $|f(t, x)| \leq g(x)$ für fast alle $x \in I$.
- (b) Wenn $f(\cdot, x)$ differenzierbar ist auf $[a, b]$ für fast alle $x \in I$ und es eine Funktion $g \in L(I)$ gibt, sodass die Ableitung nach t (x fest) majorisiert werden kann, also $\left| \frac{\partial f}{\partial t}(t, x) \right| \leq g(x)$ für alle $t \in [a, b]$ und fast alle x , dann ist G **differenzierbar** und es gilt

$$G'(t) = \int_I \frac{\partial f}{\partial t}(t, x) dx, \quad t \in [a, b].$$

Beweis: (a) Sei (t_n) eine Folge aus $[a, b]$, die gegen $\hat{t} \in [a, b]$ konvergiert. Zu jedem t_n definieren wir eine Funktion $h_n \in L(I)$ durch $h_n(x) = f(t_n, x)$. Da $f(\cdot, x) \in C([a, b])$ gilt, konvergiert h_n punktweise für fast alle $x \in I$ gegen $h(x) := f(\hat{t}, x)$ für $n \rightarrow \infty$. Außerdem ist

$$|h_n(x)| = |f(t_n, x)| \leq g(x).$$

Also ist g integrierbare Majorante zu (h_n) und der Lebesguesche Konvergenzsatz liefert

$$\lim_{n \rightarrow \infty} G(t_n) = \lim_{n \rightarrow \infty} \int_I f(t_n, x) dx = \int_I \lim_{n \rightarrow \infty} f(t_n, x) dx = \int_I f(\hat{t}, x) dx = G(\hat{t}).$$

Da diese Identität für beliebige konvergente Folgen gilt, ist $G : [a, b] \rightarrow \mathbb{R}$ stetig.

(b) Analog folgt die zweite Aussage mit

$$h_n(x) := \frac{f(t_n, x) - f(\hat{t}, x)}{t_n - \hat{t}} \quad \text{und} \quad h(x) = \frac{\partial f}{\partial t}(\hat{t}, x).$$

h_n wird durch g majorisiert, da zu $x \in I$ nach dem Mittelwertsatz der Differentialrechnung ein $s_n \in (t_n, \hat{t})$ bzw. $s_n \in (\hat{t}, t_n)$ existiert mit

$$|h_n(x)| = \left| \frac{\partial f}{\partial t}(s_n, x) \right| \leq g(x).$$

□

Beispiel 5.4 (a) Für jedes Intervall $[a, b] \subseteq \mathbb{R}$ sind die Voraussetzungen des Satzes erfüllt für die Funktion

$$G(t) = \int_0^1 \frac{e^{-(1+x^2)t^2}}{1+x^2} dx, \quad t \in [a, b].$$

Also ist G auf $\mathbb{R}$ differenzierbar und wir erhalten mit der Substitution $u = tx$

$$G'(t) = -2 \int_0^1 \frac{(1+x^2)te^{-(1+x^2)t^2}}{1+x^2} dx = -2e^{-t^2} \underbrace{\int_0^t e^{-u^2} du}_{= w(t)}.$$

Wir definieren w als die Stammfunktion zu $w'(t) = e^{-t^2}$ (die wir nicht geschlossen ausdrücken können). Der zweite Hauptsatz der Integralrechnung liefert

$$\begin{aligned} G(t) - G(0) &= \int_0^t G'(\tau) d\tau = -2 \int_0^t \left(e^{-\tau^2} \int_0^\tau e^{-x^2} dx \right) d\tau \\ &= -2 \int_0^t w'(\tau) w(\tau) d\tau \\ &= -2 (w(\tau))^2 \Big|_0^t + 2 \int_0^t w(\tau) w'(\tau) d\tau \\ &= -2 \left(\int_0^t e^{-x^2} dx \right)^2 - G(t) + G(0). \end{aligned}$$

Da

$$G(0) = \int_0^1 \frac{1}{1+x^2} dx = \arctan x \Big|_0^1 = \frac{\pi}{4}$$

ist, ergibt sich insgesamt

$$\left(\int_0^t e^{-x^2} dx \right)^2 = \frac{\pi}{4} - G(t).$$

Nun können wir noch den Grenzwert

$$\lim_{t \rightarrow \infty} G(t) = \lim_{t \rightarrow \infty} \int_0^1 \frac{e^{-(1+x^2)t^2}}{1+x^2} dx = 0$$

unter Verwendung des Lebesgueschen Konvergenzsatzes bestimmen und erhalten

$$\int_0^\infty e^{-x^2} dx = \frac{1}{2} \sqrt{\pi}.$$

□

(b) Die Voraussetzung einer Beschränkung durch eine Funktion $g \in L(I)$ ist erforderlich. Ein Gegenbeispiel liefert die Funktion $G : \mathbb{R} \rightarrow \mathbb{R}$ mit

$$G(t) = \int_{-1}^1 \frac{t}{t^2 + x^2} dx = 2 \arctan \frac{1}{t}.$$

Die Funktion ist unstetig in $t = 0$; denn $G(0) = 0$, aber

$$\lim_{\substack{t \rightarrow 0 \\ t > 0}} G(t) = \pi \quad \text{und} \quad \lim_{\substack{t \rightarrow 0 \\ t < 0}} G(t) = -\pi.$$

5.2 Definition und Grundlegende Eigenschaften der Laplace-Transformation

Zu einer beliebigen Funktion $f : [0, \infty) \rightarrow \mathbb{C}$ ist die **Laplace transformation** definiert als die Funktion $\mathcal{L}f : I \subseteq \mathbb{R}_{\geq 0} \rightarrow \mathbb{C}$, die durch das Parameterintegral

$$\mathcal{L}f(s) = \int_0^\infty f(t) e^{-st} dt, \quad s \in I,$$

gegeben ist, wenn das Integral existiert (im Sinne des Lebesgue-Integrals). Unser Ziel ist es, Eigenschaften von $\mathcal{L}f$ zu nutzen, um Differentialgleichungen zu lösen. Zunächst klären wir, für welche Funktionen wir die Laplacetransformation definieren.

Lemma 5.5 Ist $f : [0, \infty) \rightarrow \mathbb{C}$ auf jedem kompakten Intervall $[0, b] \rightarrow \mathbb{C}$ integrierbar und von **exponentiellem Typ**, d.h. es gibt Konstanten $c > 0$ und $a \in \mathbb{R}$ mit

$$|f(t)| \leq c e^{at} \quad \text{für fast alle } t \geq 0,$$

so existiert die Laplacetransformierte $\mathcal{L}f : (\alpha, \infty) \rightarrow \mathbb{C}$ für $\alpha = \max\{0, a\}$.

Beweis: Es gilt

$$\int_0^T |f(t)| e^{-st} dt \leq c \int_0^T e^{(a-s)t} dt = \frac{c}{a-s} (e^{(a-s)T} - 1) \leq \frac{c}{s-a}.$$

Also existiert die Laplacetransformierte nach Satz 6.32 im HM I Skript für alle $s > a$ und es gilt

$$\mathcal{L}f(s) = \lim_{T \rightarrow \infty} \int_0^T f(t) e^{-st} dt.$$

□

Die Menge

$$\mathcal{E} = \{f : [0, \infty) \rightarrow \mathbb{C} : f \text{ genügt den Bed. in Lemma 5.5}\}$$

bildet einen Vektorraum und beinhaltet eine große Klasse von Funktionen. Etwa sind Polynome $f(x) = \sum_{j=0}^n a_j x^j$ in $\mathcal{E}$, da sie von exponentiellem Typ sind, denn sie wachsen nicht schneller als $\exp(x)$. Auch stückweise stetige Funktionen von exponentiellem Typ liegen im Vektorraum $\mathcal{E}$. Die Funktion $f(x) = e^{x^2}$ ist **nicht** vom exponentiellen Typ.

Beispiele 5.6

(a) Für $f(t) = t^n$ mit $n \in \mathbb{N}_0$, $t \geq 0$, gilt

$$(\mathcal{L}f)(s) = \frac{n!}{s^{n+1}}, \quad s > 0.$$

Dies lässt sich mit vollständiger Induktion nach n beweisen.

(b) Für $f(t) = e^{at}$, $t \geq 0$, mit beliebiger Konstante $a \in \mathbb{C}$ ist $|f(t)| = \exp(t \operatorname{Re} a)$, d.h. es ist $f \in \mathcal{E}$. Für $s > \alpha := \max\{0, \operatorname{Re} a\}$ folgt

$$(\mathcal{L}f)(s) = \int_0^\infty e^{(a-s)t} dt = \frac{1}{a-s} e^{-(s-a)t} \Big|_{t=0}^{t \rightarrow \infty} = \frac{1}{s-a},$$

denn $|\exp(-(s-a)t)| = \exp(-(s - \operatorname{Re} a)t) \rightarrow 0$, $t \rightarrow \infty$.

(c) Für $f(t) = \begin{cases} 1 & , \quad 0 \leq t \leq 1, \\ 0 & , \quad t > 1, \end{cases}$ ergibt sich

$$(\mathcal{L}f)(s) = \int_0^1 e^{-st} dt = \frac{1}{s} (1 - e^{-s}), \quad s > 0.$$

Wir sehen an den Beispielen, dass $\mathcal{L}f$ stetig ist. Einige Rechenregeln fassen wir im folgenden Satz zusammen.

Satz 5.7 (Rechenregeln)

Sei $f \in \mathcal{E}$ und $\mathcal{L}f : (\alpha, \infty) \rightarrow \mathbb{C}$ die Laplacetransformierte von f . Dann gilt:

- (a) **Linearität:** Ist $g \in \mathcal{E}$ eine zweite Funktion mit $\mathcal{L}g : (\beta, \infty) \rightarrow \mathbb{C}$ und $\lambda, \mu \in \mathbb{C}$ Konstanten, so gilt

$$\mathcal{L}(\lambda f + \mu g)(s) = \lambda (\mathcal{L}f)(s) + \mu (\mathcal{L}g)(s), \quad s > \max\{\alpha, \beta\}.$$

- (b) **Ähnlichkeitssatz:** Sei $a > 0$. Dann ist

$$\mathcal{L}(f(at))(s) = \frac{1}{a} \mathcal{L}f\left(\frac{s}{a}\right), \quad s > a\alpha.$$

- (c) **Dämpfungssatz:** Sei $a \in \mathbb{R}$. Dann ist

$$\mathcal{L}(e^{at}f(t))(s) = \mathcal{L}f(s-a), \quad s > \max\{0, a+\alpha\}.$$

- (d) **Verschiebungssatz:** Sei $a > 0$ und

$$g(t) = \begin{cases} 0 & , \quad 0 \leq t < a, \\ f(t-a) & , \quad t \geq a. \end{cases}$$

Dann ist

$$\mathcal{L}g(s) = e^{-as} \mathcal{L}f(s), \quad s > \alpha.$$

Bemerkung: Wir haben in der Formulierung des Satzes eine bequeme, aber gefährliche Notation benutzt: Bei (b) schreiben wir z.B. $\mathcal{L}(f(at))(s)$ und meinen $\mathcal{L}g(s)$, wobei g die Funktion $g(t) = f(at)$, $t \geq 0$, ist. $\mathcal{L}$ hat als Argument eine **Funktion**, bildet also Funktionen (die wir als Funktionen mit Argument t schreiben) in Funktionen mit Argument s ab. $f(at)$ ist aber eigentlich ein **Funktionswert** (nämlich eine reelle oder komplexe Zahl), wir meinen aber die Funktion $t \mapsto f(at)$.

Beweis: (a) ergibt sich direkt aus der Definition.

(b) Mit der Substitution $\tau = at$ (also $d\tau = a dt$) folgt

$$\mathcal{L}(f(at))(s) = \int_0^\infty f(at) e^{-st} dt = \int_0^\infty f(\tau) e^{-\frac{s}{a}\tau} \frac{1}{a} d\tau = \frac{1}{a} (\mathcal{L}f)\left(\frac{s}{a}\right).$$

(c) Einsetzen in die Definition liefert

$$\mathcal{L}f(s-a) = \int_0^\infty f(t) e^{-(s-a)t} dt = \int_0^\infty [e^{at}f(t)] e^{-st} dt = \mathcal{L}(e^{at}f(t))(s).$$

(d) Die Substitution $\tau = t - a$ (also $d\tau = dt$) führt auf

$$\mathcal{L}g(s) = \int_a^\infty f(t-a) e^{-st} dt = \int_0^\infty f(\tau) e^{-s(\tau+a)} d\tau = e^{-sa} \mathcal{L}f(s).$$

Die Integrationsbereiche ergeben sich aus der Forderung, dass alle auftretenden Integrale existieren müssen. $\square$

Beispiele 5.8 Wegen $\cos \omega t = \frac{1}{2}(e^{i\omega t} + e^{-i\omega t})$ und Beispiel 5.6 ist

$$\begin{aligned} \mathcal{L}(\cos \omega t)(s) &= \frac{1}{2} [\mathcal{L}(e^{i\omega t})(s) + \mathcal{L}(e^{-i\omega t})(s)] \\ &= \frac{1}{2} \left[\frac{1}{s - i\omega} + \frac{1}{s + i\omega} \right] = \frac{s}{s^2 + \omega^2}, \quad s > 0, \omega \in \mathbb{R}, \end{aligned}$$

und analog

$$\begin{aligned} \mathcal{L}(\sin \omega t)(s) &= \frac{1}{2i} [\mathcal{L}(e^{i\omega t})(s) - \mathcal{L}(e^{-i\omega t})(s)] \\ &= \frac{1}{2i} \left[\frac{1}{s - i\omega} - \frac{1}{s + i\omega} \right] = \frac{\omega}{s^2 + \omega^2}, \quad s > 0, \omega \in \mathbb{R}. \end{aligned}$$

$\square$

Die Laplacetransformierte einer Funktion $f \in \mathcal{E}$ ist beliebig oft differenzierbar. Diese wichtige Eigenschaft wollen wir in den nächsten beiden Sätzen festhalten.

Satz 5.9

Für jedes $f \in \mathcal{E}$ ist die Laplacetransformierte $\mathcal{L}f : (\alpha, \infty) \rightarrow \mathbb{C}$ stetig, und es gilt $\lim_{s \rightarrow \infty} \mathcal{L}f(s) = 0$.

Beweis: Mit $[a, b] \subseteq (\alpha, \infty)$ gilt $|f(t)e^{-st}| \leq c e^{-(s-\alpha)t} \leq c e^{-(a-\alpha)t}$ für fast alle $t \in [0, \infty)$. Also lässt sich Satz 5.3 anwenden. Also ist $\mathcal{L}f$ stetig auf $[a, b]$ und somit in jeder Stelle $s > \alpha$. Mit der Abschätzung

$$\left| \int_0^\infty f(t) e^{-st} dt \right| \leq c \int_0^\infty e^{-(s-\alpha)t} dt = \frac{c}{s - \alpha} \rightarrow 0, \quad s \rightarrow \infty$$

erhalten wir den angegebenen Grenzwert. $\square$

Satz 5.10 (Differentiation im Bildraum)

Die Laplacetransformierte $\mathcal{L}f$ zu $f \in \mathcal{E}$ ist beliebig oft differenzierbar und es gilt:

$$\frac{d^n}{ds^n} \mathcal{L}f(s) = (-1)^n \mathcal{L}(t^n f(t))(s), \quad s > \alpha, \quad n \in \mathbb{N},$$

d.h.

$$\frac{d^n}{ds^n} \int_0^\infty f(t) e^{-st} dt = \int_0^\infty \frac{\partial^n}{\partial s^n} (f(t) e^{-st}) dt,$$

also dürfen Differentiation und Integration hier vertauscht werden.

Beweis: Wir können Satz 5.3 (Differentiation von Parameterintegralen) anwenden für $n = 1$, da wegen $f \in \mathcal{E}$ mit $g(t) = cte^{-\epsilon t}$ für $s > \alpha + \epsilon$ mit $\epsilon > 0$ eine integrierbare Majorante für $\frac{d}{ds} f(t) e^{-st}$ gegeben ist. Für $n \geq 2$ folgt das Resultat sofort durch iterative Anwendung dieses Resultats. $\square$

Beispiele 5.11

Die Differentiation im Bildraum liefert die Laplacetransformierten

$$\begin{aligned} \mathcal{L}(t \cos \omega t)(s) &= -\frac{d}{ds} \mathcal{L}(\cos \omega t)(s) \\ &= -\frac{d}{ds} \left(\frac{s}{s^2 + \omega^2} \right) = \frac{s^2 - \omega^2}{(s^2 + \omega^2)^2}, \quad s > 0, \\ \mathcal{L}(t \sin \omega t)(s) &= -\frac{d}{ds} \left(\frac{\omega}{s^2 + \omega^2} \right) = \frac{2\omega s}{(s^2 + \omega^2)^2}, \quad s > 0. \end{aligned}$$

$\square$

5.3 Anwendung bei Anfangswertaufgaben

Die wichtigste Eigenschaft der Laplacetransformierten zur Nutzung bei Differentialgleichungen ist die Differentiation im Urbildraum.

Satz 5.12 (Differentiation im Urbildraum)

Sei f n -mal stetig differenzierbar auf $[0, \infty)$ und alle $f^{(j)} \in \mathcal{E}$, $j = 0, \dots, n$. Dann gilt:

$$\mathcal{L}(f^{(n)})(s) = s^n \mathcal{L}f(s) - \sum_{j=0}^{n-1} f^{(j)}(0) s^{n-j-1}, \quad s > \alpha,$$

für jedes $n \in \mathbb{N}$. Hier bedeute (α, ∞) das gemeinsame Intervall, auf dem die Laplacetransformierte für alle $f^{(j)}$ existiert.

Beweis: Wir zeigen dies nur für $n = 1$. Mit partieller Integration ist

$$\mathcal{L}(f')(s) = \int_0^\infty f'(t) e^{-st} dt = f(t) e^{-st} \Big|_{t=0}^{t \rightarrow \infty} + s \int_0^\infty f(t) e^{-st} dt.$$

Wir merken an, dass $\lim_{t \rightarrow \infty} (f(t)e^{-st}) = 0$ für jedes $s > \alpha$, da $\int_0^\infty |f(t)| e^{-st} dt \leq c \int_0^\infty e^{(\alpha-s)t} dt = c/(s-\alpha)$ existiert. Dies liefert

$$\mathcal{L}(f')(s) = -f(0) + s \mathcal{L}f(s).$$

Der Beweis für $n \geq 2$ wird mit vollständiger Induktion geführt. $\square$

Betrachten wir nun eine lineare Differentialgleichung n -ter Ordnung mit konstanten Koeffizienten

$$a_n u^{(n)}(t) + \cdots + a_1 u'(t) + a_0 u(t) = b(t), \quad t \geq 0,$$

und Anfangsbedingungen $u^{(j)}(0) = u_j$ für $j = 0, 1, \dots, n-1$. Für eine stetige Funktion $b : [0, \infty) \rightarrow \mathbb{C}$ existiert genau eine Lösung $u \in C^n[0, \infty)$ im Raum der n -mal stetig differenzierbaren Funktionen. Mit der Differentiation im Urbildraum, also

$$\begin{aligned} (\mathcal{L}u^{(j)})(s) &= s^j \mathcal{L}u(s) - \sum_{l=0}^{j-1} u^{(l)}(0) s^{j-l-1} \\ &= s^j \mathcal{L}u(s) - \sum_{l=0}^{j-1} u_l s^{j-l-1}, \quad j = 1, 2, \dots, n, \end{aligned}$$

lässt sich nun die Differentialgleichung in den „Laplaceraum transformieren“ und wir erhalten

$$\begin{aligned} \mathcal{L}u(s) [a_n s^n + \cdots + a_1 s + a_0] \\ = \mathcal{L}b(s) + \underbrace{a_n \sum_{l=0}^{n-1} u_l s^{n-l-1} + a_{n-1} \sum_{l=0}^{n-2} u_l s^{n-l-2} + \cdots + a_1 u_0}_{=: q(s)}. \end{aligned}$$

Der Ausdruck $a_n s^n + \cdots + a_1 s + a_0 =: p(s)$ auf der linken Seite ist genau das **charakteristische Polynom** der Differentialgleichung. Im Laplaceraum haben wir also die Lösung:

$$\mathcal{L}u(s) = \frac{\mathcal{L}b(s)}{p(s)} + \frac{q(s)}{p(s)} \quad \text{für } s > \alpha,$$

wobei α größer als die größte Nullstelle von p ist. Hier ist q ebenfalls ein Polynom vom Grad kleiner oder gleich $n-1$. Damit ist das allgemeine Vorgehen offensichtlich. Wenn wir die rechte Seite als Laplacetransformierte einer Funktion schreiben könnten, d.h. $\frac{\mathcal{L}b(s)}{p(s)} + \frac{q(s)}{p(s)} = \mathcal{L}g(s)$, so ergibt sich direkt die Lösung $u = g$.

Bevor wir den Identitätssatz formulieren, der den letzten Schluss, **die Rücktransformation** begründet, betrachten wir ein Beispiel für dieses Vorgehen. Oft ist $\mathcal{L}b$ eine rationale Funktion, so dass auf der rechten Seite ebenfalls eine rationale Funktion steht. Diese können wir mit der **Partialbruchzerlegung** in einfache Summanden zerlegen, von denen wir schon berechnet haben, dass sie Laplacetransformierte zu bestimmten Funktionen sind.

Beispiel 5.13

(a) Betrachte die Anfangswertaufgabe

$$u'''(t) - 3u''(t) + 3u'(t) - u(t) = t^2 e^t, \quad t \geq 0,$$

$$u(0) = 1, \quad u'(0) = 0, \quad u''(0) = -2.$$

Dann ist

$$\begin{aligned}\mathcal{L}u'(s) &= s \mathcal{L}u(s) - u(0) = s \mathcal{L}u(s) - 1, \\ \mathcal{L}u''(s) &= s^2 \mathcal{L}u(s) - s u(0) - u'(0) = s^2 \mathcal{L}u(s) - s, \\ \mathcal{L}u'''(s) &= s^3 \mathcal{L}u(s) - s^2 u(0) - s u'(0) - u''(0) = s^3 \mathcal{L}u(s) - s^2 + 2.\end{aligned}$$

Die rechte Seite transformiert sich zu

$$\mathcal{L}(x^2 e^x)(s) = \frac{d^2}{ds^2} \mathcal{L}(e^x) = \frac{d^2}{ds^2} \frac{1}{s-1} = \frac{2}{(s-1)^3} \quad s > 1.$$

Also ist eingesetzt:

$$(s^3 \mathcal{L}u(s) - s^2 + 2) - 3(s^2 \mathcal{L}u(s) - s) + 3(s \mathcal{L}u(s) - 1) - \mathcal{L}u(s) = \frac{2}{(s-1)^3},$$

d.h.

$$\mathcal{L}u(s) \underbrace{[s^3 - 3s^2 + 3s - 1]}_{=(s-1)^3} = \frac{2}{(s-1)^3} + (s^2 - 3s + 1) \quad \text{oder}$$

$$\mathcal{L}u(s) = \frac{2}{(s-1)^6} + \frac{s^2 - 3s + 1}{(s-1)^3}, \quad s > 1.$$

Der erste Term auf der rechten Seite liegt schon in „Partialbruchform“ vor. Für den zweiten machen wir den Ansatz

$$\frac{s^2 - 3s + 1}{(s-1)^3} = \frac{a}{s-1} + \frac{b}{(s-1)^2} + \frac{c}{(s-1)^3}.$$

Multiplikation mit $(s-1)^3$ und $s \rightarrow 1$ liefert $c = -1$. Multiplikation mit s und $s \rightarrow \infty$ liefert $a = 1$. Setzt man a und c ein, so ist

$$\frac{b}{(s-1)^2} - \frac{1}{(s-1)^3} = \frac{s^2 - 3s + 1}{(s-1)^3} - \frac{1}{s-1} = -\frac{s}{(s-1)^3}.$$

Multiplikation mit $(s-1)^2$ und $s \rightarrow \infty$ liefert schließlich $b = -1$. Daher ist

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

Mit dem Dämpfungssatz ergibt sich $\mathcal{L}(t^{n-1}e^t)(s) = \frac{(n-1)!}{(s-1)^n}$ und deshalb folgt aus

$$\begin{aligned}\mathcal{L}u(s) &= \frac{2}{(s-1)^6} + \frac{1}{s-1} - \frac{1}{(s-1)^2} - \frac{1}{(s-1)^3} \\ &= \frac{2}{5!} \mathcal{L}(t^5 e^t)(s) + \mathcal{L}(e^t)(s) - \mathcal{L}(t e^t)(s) - \frac{1}{2} \mathcal{L}(t^2 e^t)(s), \quad s > 1,\end{aligned}$$

die Lösung

$$u(t) = \frac{2}{5!} t^5 e^t + e^t - t e^t - \frac{1}{2} t^2 e^t = \left[\frac{1}{60} t^5 - \frac{1}{2} t^2 - t + 1 \right] e^t.$$

Die Lösung u im Beispiel müssten wir nun eigentlich noch verifizieren, indem wir sie in die Differentialgleichung einsetzen. Der folgende wichtige Identitätssatz besagt aber, dass u die Lösung sein muss, und erspart uns deswegen diese Kontrolle.

(b) Als weiteres Beispiel betrachten wir ein Anfangswertproblem

$$y'' + 2y' + 2y = g, \quad y(0) = 0, \quad y'(0) = 0$$

mit der unstetigen Inhomogenität

$$g(t) = \begin{cases} 1, & 0 \leq t < \pi \\ 0, & t \geq \pi. \end{cases}$$

Zum Beispiel könnte dieses Modell das Abschalten einer externen Spannungsquelle bei einem Stromkreis beschreiben. Anwenden der Laplacetransformation führt auf

$$\begin{aligned} s^2 \mathcal{L}y(s) - s y(0) - y'(0) + 2(s \mathcal{L}y(s) - y(0)) + 2 \mathcal{L}y(s) \\ = \mathcal{L}g(s) = \frac{1 - e^{-\pi s}}{s} \end{aligned}$$

mit dem Verschiebungssatz. Also ist

$$\mathcal{L}y(s) = \frac{1 - e^{-\pi s}}{s(s^2 + 2s + 2)}.$$

Definieren wir

$$F(s) = \frac{1}{s(s^2 + 2s + 2)} = \frac{1}{2s} - \frac{1}{2} \frac{(s+1) + 1}{(s+1)^2 + 1},$$

so ergibt sich $\mathcal{L}f(s) = F(s)$ für $f(t) = \frac{1}{2}(1 - e^{-t} \cos t - e^{-t} \sin t)$. Mit dem Verschiebungssatz folgt die Lösung

$$y(t) = \begin{cases} f(t) & = \frac{1}{2}(1 - e^{-t} \cos t - e^{-t} \sin t), & \text{für } 0 \leq t < \pi \\ f(t) - f(t - \pi) & = -\frac{1}{2}(1 + e^{\pi}) e^{-t} (\cos t + \sin t), & \text{für } t \geq \pi. \end{cases}$$

□

Satz 5.14 („Injektivität“ der Laplacetransformation)

Seien $f, g \in \mathcal{E}$ mit $\mathcal{L}f(s) = \mathcal{L}g(s)$ für $s > \alpha$. Dann ist $f(t) = g(t)$ für fast alle $t \geq 0$.

Beweis: Wir skizzieren den Beweis für den Fall $h = f - g \in C[0, \infty) \cap \mathcal{E}$. (Der allgemeine Beweis erfordert einige zusätzliche Überlegungen - an welcher Stelle?). Sei also

$$0 = \mathcal{L}h(s) = \int_0^\infty h(t) e^{-st} dt = \int_0^\infty h(t) e^{-\alpha t} e^{-(s-\alpha)t} dt \quad \text{für } s > \alpha,$$

d.h., wenn wir $\tilde{h}(t) = h(t) e^{-\alpha t}$ setzen, gilt

$$0 = \int_0^\infty \tilde{h}(t) e^{-\sigma t} dt \quad \text{für alle } \sigma > 0.$$

Wir zeigen nun, dass $\tilde{h} = 0$ und somit $h = 0$ ist.

Für $\sigma = n \in \mathbb{N}_{\geq 1}$ führen wir die Substitutionen $t = -\ln x$ und $x = \cos \tau$ durch und erhalten

$$\begin{aligned} 0 = \mathcal{L}\tilde{h}(n) &= \int_0^\infty \tilde{h}(t)e^{-nt} dt = \int_0^1 \tilde{h}(-\ln x) x^{n-1} dx \\ &= \int_0^{\frac{\pi}{2}} \tilde{h}(-\ln(\cos \tau)) \cos^{n-1} \tau \sin \tau d\tau. \end{aligned}$$

Wir definieren nun $\psi(\tau) = \tilde{h}(-\ln(\cos \tau)) \sin \tau$ und $\psi(\pi/2 + \tau) = \psi(\pi/2 - \tau)$ für $\tau \in [0, \pi/2]$. Weiter setzen wir ψ zu einer geraden 2π -periodischen Funktion fort. Dann haben wir

$$\int_{-\pi}^{\pi} \psi(\tau) \cos^j \tau d\tau = 0$$

für alle $j = 0, 1, 2, 3, \dots$.

Aus der Eulerschen Formel und der Binomischen Formel folgt

$$(\cos nt + i \sin nt) = e^{int} = (\cos t + i \sin t)^n = \sum_{j=0}^n \binom{n}{j} (\cos^{n-j} t) i^j (\sin^j t).$$

Also gilt, wenn wir nur den Realteil dieser Identität betrachten und $\sin^2 t = 1 - \cos^2 t$ ausnutzen, eine Gleichung von der Form

$$\cos nt = \sum_{j=0}^n d_{j,n} \cos^j t$$

mit Zahlen $d_{j,n} \in \mathbb{R}$. Wegen

$$\int_{-\pi}^{\pi} \psi(t) \cos^j t dt = 0$$

für alle $j = 0, 1, 2, 3, \dots$, kann man durch Induktion nachweisen, dass auch alle Fourier-Koeffizienten $\frac{1}{2\pi} \int_{-\pi}^{\pi} \psi(t) \cos(nt) dt$ von ψ verschwinden. Die Parsevalsche Gleichung (s. Satz 3.13) impliziert $\psi = 0$ bzw. $\tilde{h} = 0$. □

Auch Systeme von linearen Differentialgleichungen können wir mit der Laplacetransformation lösen.

Beispiel 5.15

Bestimme die Funktionen u und v mit

$$\begin{aligned} u'(t) - v'(t) - 2u(t) + 2v(t) &= \sin t, \\ u''(t) + 2v'(t) + u(t) &= 0, \end{aligned}$$

und den Anfangsbedingungen $u(0) = u'(0) = 0$, $v(0) = 0$.

Wegen der homogenen Randbedingung ist $\mathcal{L}(u'') = s^2 \mathcal{L}u$, $\mathcal{L}(u') = s \mathcal{L}u$, $\mathcal{L}(v') = s \mathcal{L}v$. Wir setzen zur Abkürzung $U := \mathcal{L}u$ und $V := \mathcal{L}v$. Damit wird aus dem System von Differentialgleichungen

für u, v ein lineares Gleichungssystem für U und V :

$$\begin{aligned} sU(s) - sV(s) - 2U(s) + 2V(s) &= \frac{1}{1+s^2}, \\ s^2U(s) + 2sV(s) + U(s) &= 0, \quad \text{also} \\ U(s)(s-2) + V(s)(2-s) &= \frac{1}{1+s^2}, \\ U(s)(s^2+1) + 2sV(s) &= 0. \end{aligned}$$

Multiplikation der ersten Gleichung mit $(1+s^2)$, der zweiten mit $(s-2)$ und Subtraktion liefert

$$\begin{aligned} V(s) &= -\frac{1}{(s-2)(1+s^2+2s)} = -\frac{1}{(s-2)(s+1)^2}, \quad s > 2, \\ U(s) &= -\frac{2s}{1+s^2}V(s) = \frac{2s}{(s+i)(s-i)(s-2)(s+1)^2}, \quad s > 2. \end{aligned}$$

Wir führen die Partialbruchzerlegung von V und U durch: Der Ansatz

$$V(s) = \frac{\alpha}{s-2} + \frac{\beta}{(s+1)^2} + \frac{\gamma}{s+1} = -\frac{1}{(s-2)(s+1)^2}, \quad s > 2,$$

liefert:

- (i) $\alpha + (s-2)\left(\frac{\beta}{(s+1)^2} + \frac{\gamma}{s+1}\right) = -\frac{1}{(s+1)^2}$ Für $s \rightarrow 2$ folgt $\alpha = -1/9$.
- (ii) $\beta + (s+1)^2\frac{\alpha}{s-2} + (s+1)\gamma = -\frac{1}{s-2}$. Also folgt für $s \rightarrow -1$, dass $\beta = 1/3$.
- (iii) $\gamma + \frac{s+1}{s-2}\alpha + \frac{\beta}{s+1} = -\frac{1}{(s-2)(s+1)}$. Für $s \rightarrow +\infty$ liefert dies $\gamma + \alpha = 0$, also $\gamma = -\alpha = \frac{1}{9}$.

Daher ist

$$V(s) = \mathcal{L}v(s) = -\frac{1}{9}\frac{1}{s-2} + \frac{1}{3}\frac{1}{(s+1)^2} + \frac{1}{9}\frac{1}{s+1}, \quad s > 2,$$

und die Rücktransformation liefert

$$v(t) = -\frac{1}{9}e^{2t} + \frac{1}{3}te^{-t} + \frac{1}{9}e^{-t}.$$

Analog liefert der Ansatz für die komplexe Partialbruchzerlegung

$$\begin{aligned} U(s) &= \frac{\alpha}{s+i} + \frac{\beta}{s-i} + \frac{\gamma}{s-2} + \frac{\delta}{(s+1)^2} + \frac{\eta}{s+1} \\ &= \frac{2s}{(s+i)(s-i)(s-2)(s+1)^2} \end{aligned}$$

die Konstanten

$$\begin{aligned}
\alpha &= \lim_{s \rightarrow -i} [(s+i)U(s)] = \frac{-2i}{(-2i)(-i-2)(1-i)^2} = \dots = -\frac{1+2i}{10}, \\
\beta &= \lim_{s \rightarrow +i} [(s-i)U(s)] = \frac{2i}{(2i)(i-2)(1+i)} = -\frac{1-2i}{10}, \\
\gamma &= \lim_{s \rightarrow 2} [(s-2)U(s)] = \frac{4}{5 \cdot 9} = \frac{4}{45}, \\
\delta &= \lim_{s \rightarrow -1} [(s+1)^2 U(s)] = \frac{-2}{2 \cdot (-3)} = \frac{1}{3}, \\
\alpha + \beta + \gamma + \eta &= \lim_{s \rightarrow \infty} [s U(s)] = 0, \quad \text{also} \\
\eta &= -(\alpha + \beta) - \gamma = \frac{1}{5} - \frac{4}{45} = \frac{1}{9}.
\end{aligned}$$

Also ist

$$\begin{aligned}
u(t) &= -\frac{1+2i}{10} e^{-it} - \frac{1-2i}{10} e^{it} + \frac{4}{45} e^{2t} + \frac{1}{3} t e^{-t} + \frac{1}{9} e^{-t} \\
&= -\frac{1}{10} (2 \cos t + 4 \sin t) + \frac{4}{45} e^{2t} + \frac{1}{3} t e^{-t} + \frac{1}{9} e^{-t} \\
&= -\frac{1}{5} \cos t - \frac{2}{5} \sin t + \frac{4}{45} e^{2t} + \frac{1}{3} t e^{-t} + \frac{1}{9} e^{-t}.
\end{aligned}$$

Für U haben wir hier die komplexe Partialbruchzerlegung in Linearfaktoren durchgeführt. Alternativ dazu können wir auch die reelle Partialbruchzerlegung vornehmen mit dem Ansatz

$$U(s) = \frac{\alpha s + \beta}{s^2 + 1} + \frac{\gamma}{s - 2} + \frac{\delta}{(s + 1)^2} + \frac{\eta}{s + 1} = \frac{2s}{(s^2 + 1)(s - 2)(s + 1)^2}.$$

Wie oben erhalten wir $\gamma = \frac{4}{45}$, $\delta = \frac{1}{3}$, sowie $\alpha + \gamma + \eta = \lim_{s \rightarrow \infty} [s U(s)] = 0$. Für $s = 0$ erhalten wir $\beta - \frac{\gamma}{2} + \delta + \eta = U(0) = 0$. Schließlich ist

$$\begin{aligned}
\eta &= \lim_{s \rightarrow -1} \left[(s+1) U(s) - \frac{\delta}{s+1} \right] \\
&= \lim_{s \rightarrow -1} \left[\frac{2s}{(s^2 + 1)(s - 2)(s + 1)} - \frac{1}{3(s + 1)} \right] \\
&= \lim_{s \rightarrow -1} \left[\frac{-s^3 + 2s^2 + 5s + 2}{3(s^2 + 1)(s - 2)(s + 1)} \right] = \lim_{s \rightarrow -1} \frac{-s^2 + 3s + 2}{3(s^2 + 1)(s - 2)} = \frac{1}{9}.
\end{aligned}$$

Also $\alpha = -\gamma - \eta = -\frac{4}{45} - \frac{1}{9} = -\frac{9}{45} = -\frac{1}{5}$ und $\beta = \frac{\gamma}{2} - \delta - \eta = \frac{2}{45} - \frac{1}{3} - \frac{1}{9} = -\frac{2}{5}$. Daher ist

$$U(s) = -\frac{1}{5} \frac{s+2}{(s^2+1)} + \frac{4}{45} \frac{1}{s-2} + \frac{1}{3} \frac{1}{(s+1)^2} + \frac{1}{9} \frac{1}{s+1}, \quad s > 2.$$

Aus den am Ende des Kapitels aufgelisteten Formeln entnehmen wir

$$\frac{s+2}{s^2+1} = \mathcal{L}(\cos t + 2 \sin t)(s).$$

Damit folgt schließlich

$$u(t) = -\frac{1}{5}(\cos t + 2 \sin t) + \frac{4}{45}e^{2t} + \frac{1}{3}te^{-t} + \frac{1}{9}e^{-t},$$

und dies stimmt mit der komplexen Herleitung überein. $\square$

5.4 Faltungsprodukt und Faltungssatz

Eine generelle Möglichkeit, Lösungen zu inhomogenen linearen Differentialgleichungen zu bestimmen, ergibt sich durch das Faltungsprodukt.

Definition 5.16 (Faltung)

Zu $f, g : [0, \infty) \rightarrow \mathbb{C}$ ist das **Faltungsprodukt** oder die **Faltung** $f * g : [0, \infty) \rightarrow \mathbb{C}$ definiert durch

$$(f * g)(t) := \int_0^t f(t - \tau) g(\tau) d\tau, \quad t \geq 0,$$

wenn das Integral für alle $t \in \mathbb{R}_{\geq 0}$ existiert.

Die folgenden Gesetze folgen unmittelbar aus der Definition:

- (a) **Kommutativität:** Mit der Substitution $s = t - \tau$ erhalten wir

$$(f * g)(t) = \int_0^t f(s) g(t - s) ds = (g * f)(t).$$

- (b) **Assoziativität:** Mit der Substitution $\tau = t - s - \sigma$ und Vertauschen der Integrationsreihenfolge erhalten wir

$$\begin{aligned} (f * (g * h))(t) &= \int_0^t \int_0^{t-s} f(s) g(\sigma) h(t - s - \sigma) d\sigma ds \\ &= - \int_0^t \int_{t-s}^0 f(s) g(t - s - \tau) h(\tau) d\tau ds \\ &= \int_0^t \left(\int_0^{t-\tau} f(s) g(t - \tau - s) ds \right) h(\tau) d\tau = ((f * g) * h)(t). \end{aligned}$$

- (c) **Distributivgesetze:** Es gilt $f * (g + h) = f * g + f * h$ und $(f + g) * h = f * h + g * h$.

- (d) **Homogenitätsgesetze:** Mit $\lambda \in \mathbb{C}$ gilt $f * (\lambda g) = \lambda(f * g) = (\lambda f) * g$.

Beispiel 5.17

Sei $f_n(t) = \frac{t^n}{n!}$, $t \geq 0$, $n \in \mathbb{N}_0$. Wir zeigen $f_n * f_m = f_{n+m+1}$ für alle $n, m \in \mathbb{N}_0$ durch vollständige Induktion nach m .

Induktionsanfang: Sei $m = 0$, $n \in \mathbb{N}_0$ beliebig. Dann ist

$$(f_n * f_0)(t) = \int_0^t \frac{s^n}{n!} ds = \frac{t^{n+1}}{(n+1)!} = f_{n+1}(t).$$

Induktionsschritt: Sei die Behauptung richtig für ein $m \in \mathbb{N}_0$ und alle $n \in \mathbb{N}_0$. Dann folgt mit den Eigenschaften (a), (b) und dieser Induktionsvoraussetzung

$$(f_n * f_{m+1})(t) = (f_n * (f_0 * f_m))(t) = ((f_n * f_0) * f_m)(t) = f_{n+1} * f_m = f_{n+2+m}.$$

□

Der folgende Satz besagt, dass die Faltung zweier Funktionen f und g so glatt ist wie die glattere der beiden Funktionen:

Satz 5.18

- (a) Ist $f : [0, \infty) \rightarrow \mathbb{C}$ auf jedem beschränkten Intervall $I \subset [0, \infty)$ integrierbar und $g : [0, \infty) \rightarrow \mathbb{C}$ stetig, dann ist auch $f * g : [0, \infty) \rightarrow \mathbb{C}$ stetig.
- (b) Ist $f : [0, \infty) \rightarrow \mathbb{C}$ stetig und $g : [0, \infty) \rightarrow \mathbb{C}$ stetig differenzierbar, so ist auch $f * g$ stetig differenzierbar und es gilt

$$(f * g)'(t) = (f * g')(t) + g(0)f(t).$$

Beweis: (a) Sei $t \in [0, c]$, $h \in \mathbb{R}$ mit $t+h \in [0, c]$. Da g auf $[0, c]$ stetig ist, gibt es (in Abhängigkeit von $c > 0$) eine Konstante k mit $|g(t)| \leq k$ für alle $t \in [0, c]$. Es ist

$$\begin{aligned} (f * g)(t+h) - (f * g)(t) &= \int_0^t f(\tau) [g(t+h-\tau) - g(t-\tau)] d\tau \\ &\quad + \int_t^{t+h} f(\tau) g(t+h-\tau) d\tau. \end{aligned}$$

Da f integrierbar ist, ist durch $2k|f(\tau)|$ eine integrierbare Majorante für das erste Integral gegeben. Also konvergiert das erste Integral nach dem Lebesgueschen Konvergenzsatz gegen 0 für $h \rightarrow 0$. Das zweite Integral konvergiert auch gegen Null für $h \rightarrow 0$, da es sich durch $k \int_t^{t+h} |f(t)| dt \rightarrow 0$ abschätzen lässt.

(b) Der Differenzenquotient ist

$$\begin{aligned} &\frac{1}{h} \left(\int_0^{t+h} g(t+h-\tau) f(\tau) d\tau - \int_0^t g(t-\tau) f(\tau) d\tau \right) \\ &= \frac{1}{h} \int_t^{t+h} g(t+h-\tau) f(\tau) d\tau + \frac{1}{h} \int_0^t (g(t+h-\tau) - g(t-\tau)) f(\tau) d\tau. \end{aligned}$$

Nach dem Mittelwertsatz der Integralrechnung gibt es zu h ein $z \in (t, t+h)$, sodass für das erste Integral

$$\frac{1}{h} \int_t^{t+h} g(t+h-\tau) f(\tau) d\tau = g(t+h-z) f(z) \rightarrow g(0) f(t)$$

gilt. Beim zweiten Integral dürfen wir nach dem Lebesgueschen Konvergenzsatz den Grenzwert $h \rightarrow 0$ mit dem Integral vertauschen. Also erhalten wir insgesamt, dass die Ableitung existiert und es ist

$$(f * g)'(t) = \frac{d}{dt} \int_0^t g(t-\tau) f(\tau) d\tau = g(0) f(t) + \int_0^t g'(t-\tau) f(\tau) d\tau = g(0) f(t) + (f * g')(t) .$$

□

Die für uns wichtigste Aussage in Bezug auf Faltungen ist der Faltungssatz, der uns die Möglichkeit gibt, die Rücktransformation von einem Produkt von Funktionen in Form eines Faltungsproduktes anzugeben.

Satz 5.19 (*Faltungssatz*)

Seien $f, g \in \mathcal{E}$ und $\mathcal{L}f : (\alpha, \infty) \rightarrow \mathbb{C}$, $\mathcal{L}g : (\beta, \infty) \rightarrow \mathbb{C}$. Dann ist auch $f * g \in \mathcal{E}$ und es gilt:

$$\mathcal{L}(f * g)(s) = (\mathcal{L}f)(s) \cdot (\mathcal{L}g)(s) \quad \text{für } s > \gamma := \max\{\alpha, \beta\}.$$

Beweisidee: Zunächst überlegen wir uns, dass $f * g$ von exponentiellem Typ ist, da

$$\begin{aligned} |(f * g)(t)| &\leq \int_0^t |f(t-\tau)| |g(\tau)| d\tau \\ &\leq c \int_0^t e^{\alpha(t-\tau)} e^{\beta\tau} d\tau \leq c \int_0^t e^{\gamma(t-\tau)+\gamma\tau} d\tau \leq ct e^{\gamma t}. \end{aligned}$$

Für den weiteren Beweis müssen wir Integrationsreihenfolgen vertauschen. Dies deuten wir hier nur an, da wir uns erst im nächsten Semester mit solchen Integralen beschäftigen werden. Es gilt

$$\begin{aligned} \mathcal{L}(f * g)(s) &= \int_0^\infty \int_0^t f(t-\tau) g(\tau) d\tau e^{-st} dt \\ &= \int_0^\infty \int_\tau^\infty f(t-\tau) g(\tau) e^{-s(t-\tau+\tau)} dt d\tau \\ &= \int_0^\infty \int_0^\infty f(u) e^{-su} du g(\tau) e^{-s\tau} d\tau \quad (\text{mit } u = t - \tau) \\ &= \mathcal{L}f(s) \int_0^\infty g(\tau) e^{-s\tau} d\tau = \mathcal{L}f(s) \cdot \mathcal{L}g(s), \end{aligned}$$

und somit wäre der Satz bewiesen.

□

Beispiel 5.20

(a) Betrachte Anfangswertprobleme von der Form

$$y''(t) + 4y(t) = g(t), \quad y(0) = 3, \quad y'(0) = -1$$

mit einer stetigen rechten Seite g von exponentiellem Typ. Transformieren wir diese Gleichung, so ergibt sich

$$s^2 \mathcal{L}y(s) - 3s + 1 + 4\mathcal{L}y(s) = \mathcal{L}g(s).$$

Also ist

$$\mathcal{L}y(s) = 3 \frac{s}{s^2 + 4} - \frac{1}{s^2 + 4} + \frac{1}{s^2 + 4} \mathcal{L}g(s),$$

und mit dem Faltungssatz folgt die Lösung

$$y(t) = 3 \cos 2t - \frac{1}{2} \sin 2t + \frac{1}{2} \int_0^t \sin 2(t - \tau) g(\tau) d\tau.$$

Beachte die Struktur: das Faltungsintegral ist eine partikuläre Lösung der inhomogenen Gleichung, und mit $\{\cos 2t, \sin 2t\}$ ist ein Fundamentalsystem zur homogenen Differentialgleichung gegeben.

Als Fazit lässt sich also festhalten, dass sich in Form eines Faltungsintegrals die Lösung zu jeder beliebigen Inhomogenität bei linearen Differentialgleichungen mit konstanten Koeffizienten angeben lässt.

(b) Mit Hilfe des Faltungssatzes und der Idee, eine Gleichung zu transformieren, lassen sich auch allgemeinere sogenannte Integrodifferentialgleichungen in manchen Fällen lösen. Gesucht ist etwa eine Lösung u der Gleichung

$$u'(t) - \int_0^t (t - s)u(s) ds = 1.$$

Anwendung der Laplacetransformation führt mit dem Faltungssatz auf die Gleichung

$$s\mathcal{L}u(s) - u(0) - \mathcal{L}g(s) \mathcal{L}u(s) = (\mathcal{L}(1))(s)$$

mit $g(t) = t$. Für $u(0) = 0$ erhalten wir aus der Tabelle der Laplacetransformationen und mit Hilfe anschließender Partialbruchzerlegung

$$\mathcal{L}u(s) = \frac{s}{s^3 - 1} = \frac{1}{3} \frac{1}{s - 1} - \frac{1}{3} \frac{s + \frac{1}{2}}{(s + \frac{1}{2})^2 + \frac{3}{4}} + \frac{1}{2} \frac{1}{(s + \frac{1}{2})^2 + \frac{3}{4}}.$$

Die Rücktransformation liefert die gesuchte Lösung

$$u(t) = \frac{1}{3}e^t - \frac{1}{3}e^{-\frac{1}{2}t} \cos \frac{\sqrt{3}}{2}t + \frac{1}{\sqrt{3}}e^{-\frac{1}{2}t} \sin \frac{\sqrt{3}}{2}t.$$

□

5.5 Die Delta-Distribution

Häufig treten in der Praxis anregende Kräfte in Form von Impulsen auf. Wir nehmen an, eine Schwingung, die durch die Differentialgleichung

$$u''(t) + u(t) = g(t), \quad t > 0,$$

beschrieben wird, werde durch einen solchen Impuls angeregt, d.h. die Inhomogenität g ist nur in einem extrem kleinen Zeitintervall von Null verschieden. Um die Situation zu modellieren, betrachte zunächst

$$g_\varepsilon(t) = \begin{cases} \frac{1}{2\varepsilon}, & t \in (t_0 - \varepsilon, t_0 + \varepsilon) \\ 0, & \text{sonst} \end{cases}$$

Die Funktionen g_ε sind so definiert, dass

$$\int_{-\infty}^{\infty} g_\varepsilon(t) dt = \int_{t_0-\varepsilon}^{t_0+\varepsilon} \frac{1}{2\varepsilon} dt = 1$$

gilt, also der zur Verfügung stehende „Impuls“ unabhängig von ε konstant ist. Was passiert für $\varepsilon \rightarrow 0$? Es gilt

$$\lim_{\varepsilon \rightarrow 0} g_\varepsilon(t) = 0 \quad \text{für } t \neq t_0 \quad \text{und} \quad \lim_{\varepsilon \rightarrow 0} \int_{-\infty}^{\infty} g_\varepsilon(t) dt = 1.$$

Also kann die abstrakte Vorstellung von einem Impuls zum Zeitpunkt $t = t_0$ nicht durch eine Funktion beschrieben werden. Allgemeiner erhalten wir für stetige Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ mit dem Mittelwertsatz

$$\int_{-\infty}^{\infty} g_\varepsilon(t) f(t) dt = \frac{1}{2\varepsilon} \int_{t_0-\varepsilon}^{t_0+\varepsilon} f(t) dt = f(z)$$

für eine Zwischenstelle $z \in (t_0 - \varepsilon, t_0 + \varepsilon)$. Also ist im Grenzfall

$$\lim_{\varepsilon \rightarrow 0} \int_{-\infty}^{\infty} g_\varepsilon(t) f(t) dt = f(t_0).$$

Diese Beobachtungen führen auf die folgende Definition:

Definition 5.21 (*Delta-Distribution*)

(a) Eine lineare Abbildung $\mathcal{D} : C_0^\infty(\mathbb{R}) \rightarrow \mathbb{C}$ heißt **Distribution** (oder **verallgemeinerte Funktion**), wobei

$$C_0^\infty(\mathbb{R}) = \{\varphi \in C^\infty(\mathbb{R}) : \text{es gibt } b > 0 \text{ mit } \varphi(x) = 0 \text{ für } |x| > b\}.$$

(b) Die **Delta-Distribution** $\mathcal{D}_\delta$ ist die Distribution, die durch

$$\mathcal{D}_\delta \varphi = \varphi(0), \quad \text{für alle } \varphi \in C_0^\infty(\mathbb{R})$$

definiert ist.

Beispiel: Für $f \in L^1(\mathbb{R})$ ist durch

$$\mathcal{D}_f \varphi = \int_{\mathbb{R}} f(x) \varphi(x) dx$$

eine Distribution definiert. Setzen wir etwa $f = g_\varepsilon$ mit $t_0 = 0$ (s.o.), so folgt

$$\lim_{\varepsilon \rightarrow 0} \mathcal{D}_{g_\varepsilon} \varphi = \lim_{\varepsilon \rightarrow 0} \int_{\mathbb{R}} g_\varepsilon(x) \varphi(x) dx = \varphi(0) = \mathcal{D}_\delta \varphi$$

Also konvergiert $\mathcal{D}_{g_\varepsilon}$ im Sinne der Distributionen gegen die Delta-Distribution.

Dieses Beispiel erläutert die in den Anwendungen übliche Notation:

$$\mathcal{D}_\delta \varphi = \int_{\mathbb{R}} \delta(x) \varphi(x) dx,$$

wobei das Symbol $\delta(x)$ stets in diesem Sinne einer verallgemeinerten Funktion zu verstehen ist.

Lemma 5.22 Seien $f, g \in C(\mathbb{R})$ stetige Funktionen und für die zugeordneten Distributionen gelte $\mathcal{D}_f = \mathcal{D}_g$, d.h.

$$\mathcal{D}_f \varphi = \int_{\mathbb{R}} f(x) \varphi(x) dx = \int_{\mathbb{R}} g(x) \varphi(x) dx = \mathcal{D}_g \varphi$$

für jede Funktion $\varphi \in C_0^\infty(\mathbb{R})$, dann ist $f = g$.

Beweis: Wir zeigen einen Widerspruch und nehmen dazu an, dass $f(x) \neq g(x)$ für ein x , ohne Beschränkung $f(x) \geq g(x)$. Wegen der Stetigkeit von f und h folgt $f \geq g$ auf einem ganzen Intervall $I = (a, b) \subseteq \mathbb{R}$. Nun konstruieren wir eine „Testfunktion“, indem wir zunächst

$$\psi(t) = \begin{cases} e^{-\frac{1}{t^2}}, & \text{für } t > 0 \\ 0, & \text{für } t \leq 0 \end{cases}$$

festlegen und dann

$$\varphi(t) = \begin{cases} 0, & \text{für } t \leq a \\ \psi\left(\frac{t-a}{\frac{a+b}{2}-t}\right) & \text{für } a < t \leq \frac{a+b}{2} \\ \psi\left(\frac{b-t}{t-\frac{a+b}{2}}\right) & \text{für } \frac{a+b}{2} < t < b \\ 0, & \text{für } t \geq b \end{cases}$$

definieren. Es gilt $\varphi \in C_0^\infty(\mathbb{R})$ (s. HMI, Beispiel 5.36) und wir erhalten

$$0 = (F_f - F_g)(\varphi) = \int_a^b (f(x) - g(x)) \varphi(x) dx.$$

Mit den Eigenschaften integrierbarer Funktionen (s. Satz 6.8(e) in HM I) folgt $(f - g) \varphi = 0$ f.ü. auf I , also $f = g$ auf I , ein Widerspruch. $\square$

Aufgrund des Lemmas ist eine eindeutige Zuordnung $f \leftrightarrow F_f$ gegeben. Also können wir die stetigen Funktionen $C(\mathbb{R})$ als Teilmenge der Menge aller Distributionen auffassen. Das Lemma lässt sich auch etwa für integrierbare Funktionen verallgemeinern, ist aber aufwendiger, so dass wir auch $L^1(\mathbb{R})$ als Teilmenge der Distributionen auffassen können.

Für unser Vorhaben ist noch die Laplacetransformation der Delta-Distribution wichtig. Wir ver-

wenden die im Beispiel angegebene Approximation:

$$\begin{aligned}
\mathcal{L}(\delta(\cdot - t_0))(s) &= \lim_{\varepsilon \rightarrow 0} \mathcal{L}(g_\varepsilon(t - t_0))(s) \\
&= \lim_{\varepsilon \rightarrow 0} \int_0^\infty g_\varepsilon(t - t_0) e^{-st} dt \\
&= \lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \int_{t_0-\varepsilon}^{t_0+\varepsilon} e^{-st} dt \\
&= \lim_{\varepsilon \rightarrow 0} \frac{\sinh(\varepsilon s)}{\varepsilon s} e^{-t_0 s} = e^{-t_0 s} \quad \text{für } t_0 > 0
\end{aligned}$$

(Regel von de l'Hospital).

Definition 5.23 Für die Delta-Distribution δ ist

$$\mathcal{L}(\delta(\cdot - t_0))(s) = e^{-t_0 s}.$$

Somit lässt sich nun die abstrakte Vorstellung von einem Impuls im einführenden Beispiel durch die Differentialgleichung

$$u''(t) + u(t) = \delta(t - t_0), \quad t > 0,$$

im Sinne von Distributionen beschreiben. Nehmen wir noch die Anfangswerte $u(0) = u'(0) = 0$ an, so erhalten wir die transformierte Gleichung

$$\mathcal{L}u''(s) + \mathcal{L}u(s) = (s^2 + 1)\mathcal{L}u(s) = e^{-t_0 s}$$

bzw.

$$\mathcal{L}u(s) = \frac{1}{s^2 + 1} e^{-t_0 s}.$$

Der Verschiebungssatz liefert

$$u(t) = \begin{cases} 0, & t < t_0 \\ \sin(t - t_0), & t \geq t_0. \end{cases}$$

Also lässt sich mit Hilfe der Delta-Distribution eine allgemeinere Klasse von Differentialgleichungen bequem behandeln. Beachte, dass die Lösung der Differentialgleichung nicht stetig differenzierbar in $t = t_0$ ist, d.h. wir sind hier insbesondere einem erweiterten Begriff von einer „Lösung“ einer Differentialgleichung begegnet.

Zum Abschluss listen wir die wichtigsten Laplacetransformationen auf.

Originalfunktion $f(t)$	Bildfunktion $\mathcal{L}f(s) = F(s)$	Anmerkung
t^n	$\frac{n!}{s^{n+1}}$	$n = 0, 1, 2, \dots$
$e^{\lambda t}$	$\frac{1}{s - \lambda}$	$\lambda \in \mathbb{C}$
$\cos \omega t$	$\frac{s}{s^2 + \omega^2}$	$\omega \in \mathbb{R}$
$\sin \omega t$	$\frac{\omega}{s^2 + \omega^2}$	$\omega \in \mathbb{R}$
$\cosh \mu t$	$\frac{s}{s^2 - \mu^2}$	$\mu \in \mathbb{R}$
$\sinh \mu t$	$\frac{\mu}{s^2 - \mu^2}$	$\mu \in \mathbb{R}$
$e^{\mu t} \cos \omega t$	$\frac{s - \mu}{(s - \mu)^2 + \omega^2}$	$\omega, \mu \in \mathbb{R}$
$e^{\mu t} \sin \omega t$	$\frac{\omega}{(s - \mu)^2 + \omega^2}$	$\omega, \mu \in \mathbb{R}$
$t \cos \omega t$	$\frac{s^2 - \omega^2}{(s^2 + \omega^2)^2}$	$\omega \in \mathbb{R}$
$t \sin \omega t$	$\frac{2\omega s}{(s^2 + \omega^2)^2}$	$\omega \in \mathbb{R}$
$t \cosh \mu t$	$\frac{s^2 + \mu^2}{(s^2 - \mu^2)^2}$	$\mu \in \mathbb{R}$
$t \sinh \mu t$	$\frac{2\mu s}{(s^2 - \mu^2)^2}$	$\mu \in \mathbb{R}$
$t^n e^{\lambda t}$	$\frac{n!}{(s - \lambda)^{n+1}}$	$n = 0, 1, 2, \dots, \quad \lambda \in \mathbb{C}$
$\frac{\sin t}{t}$	$\operatorname{arccot} s$	
$J_0(at)$	$\frac{1}{\sqrt{s^2 + a^2}}$	$a \in (0, \infty)$

Index

- Ähnlichkeitssatz, 96
- Abbildung
 - linear, 25
- Abbildungsmatrix, 26
- abhängige Variablen, 10
- Abstand, 24
 - euklidischer, 18
- adjungierte Matrix, 29
- affiner Unterraum, 6
- Anfangswertproblem, 69
- Ansatz in der Form einer Potenzreihe, 82
- Ansatz vom Typ der rechten Seite, 79
- Basis, 14
- Besselsche Differentialgleichung, 84
- Besselsche Funktion, 85
- Besselsche Ungleichung, 54
- Betrag, 18
- Bild
 - einer Matrix, 32
- binomische Formel, 18
- Cauchy-Schwarzsche Ungleichung, 19
- charakteristisches Polynom, 44, 63
- Dämpfungssatz, 96
- Darstellungsmatrix, 26
- Delta-Distribution, 108
- Determinante, 37
 - Definition, 39
- Differentialgleichung
 - n -ter Ordnung, 70
 - linear, 70
 - mit konstanten Koeffizienten, 63
 - numerische Lösungsverfahren, 87
- Differentialgleichungssystem, 69
 - homogen, 71
 - inhomogen, 71
 - linear, 71
- Distribution, 109
 - Delta-Distribution, 109, 111
- Drehung, 30
- Dreiecksungleichung, 19
- dyadisches Produkt, 30
- Ebene, 6
 - Hesse'sche Normalenform, 23
 - Parameterform, 7
- Eigenvektor, 44
- Eigenwert, 44
- elektrischer Schwingkreis, 62
- Element, 26
 - neutrales, 4
 - Nullelement, 4
 - Pivotelement, 9
- elementare Umformungen, 9
- euklidischer Abstand, 18
- euklidischer Vektorraum, 20
- Eulersche Differentialgleichung, 67
- Eulersches Polygonzugverfahren, 88
- Exponentialansatz, 63
- Faltung, 105
- Faltungsprodukt, 105
- Fourierkoeffizient, 52
- Fourierpolynom, 52
- Fourierreihe, 60
- Fourierscher Entwicklungssatz, 60
- freie Variable, 10
- Fundamentalmatrix, 72
- Fundamentalsystem, 72
 - kanonisches, 72
- Funktion
 - gerade, 53
 - periodische, 48
 - stückweise stetig, 55
 - ungerade, 53
 - von exponentiellem Typ, 95
- Gammafunktion, 92
- Gauß'sches Eliminationsverfahren, 9
- Gerade, 6
- Gregory-Leibniz-Reihe, 59
- hermitesch, 19
- Hesse'sche Normalenform, 23
- homogenes Gleichungssystem, 8
- inhomogenes Gleichungssystem, 8
- inneres Produkt, 18
- Kern, 32
- Koeffizienten, 26
- Komposition, 27
- Konvergenz

- im quadratischen Mittel, 60
- Koordinateneinheitsvektor, 14
- Kreuzprodukt, 38
- Kroneckersymbol, 26
- Laplacetransformation, 94
 - Assoziativität der Faltung, 105
 - Differentiation im Bildraum, 98
 - Distributivgesetze der Faltung, 105
 - Eigenschaften der Faltung, 106
 - Faltung, 105
 - Faltungssatz, 107
 - Homogenitätsgesetze, 105
 - Injektivität, 101
 - Kommutativität der Faltung, 105
 - Rücktransformation, 99
 - Tabelle, 112
- Lebesguescher Konvergenzsatz, 92
- lineare Abbildung, 25
 - Abbildungsmatrix, 26
 - Darstellungsmatrix, 26
- Lineare Unabhängigkeit, 12
- linearer Unterraum, 6
- lineares Gleichungssystem, 8
 - Lösungsfälle, 33
- Linearität, 96
- Linearkombination, 12
- Lotfußpunkt, 22, 23
- Lotka - Volterra Modell, 69
- Matrix, 26
 - ähnlich, 37
 - adjungierte, 29
 - Bild, 32
 - Determinante, 37
 - inverse, 34
 - Regeln, 36
 - Kern, 32
 - orthogonale, 31
 - Rechenregeln, 28
 - reguläre oder nicht singuläre, 34
 - transponierte, 29
 - unitäre, 31
- Matrixprodukt, 27
- Maximumnorm, 70
- neutrales Element, 4
- nichttriviale Lösung, 33
- Norm, 18
- Normalform einer Geraden oder Ebene, 21
- normierter Vektorraum, 20
- Nullelement, 4
- Nullmatrix, 26
- orthogonal, 20
- orthogonale Matrix, 31
- orthogonale Projektion, 22, 23, 31
- Orthonormalbasis, 50
- Parallelogramm, 37
- Parameterform, 6, 7
- Parameterintegrale, 92
- Parsevalsche Gleichung, 60
- Partialbruchzerlegung, 99
- partikuläre Lösung, 32
- periodische Fortsetzung, 54
- Pivotelement, 9
- Polynom
 - trigonometrisches, 47
- Potenzansatz, 68
- Potenzreihenansatz
 - verallgemeinerter, 84
- Potenzreihenmethode, 82
- Projektion
 - orthogonale, 22, 23
- quadratisches Mittel, 49
- Räuber-Beute Modell, 68
- Rang, 9
- Rechtssystem, 39
- Reduktion der Ordnung, 66, 82
- reguläre Matrix, 34
- Resonanz, 79
- Riemann-Lebesgue-Lemma, 55
- Runge-Kutta-Verfahren, 89
- Satz
 - Differentiation im Bildraum
 - (Laplacetransformation), 98
 - Eigenschaften der Determinante, 42
 - Eigenschaften der Faltung, 106
 - Entwicklungsregel der Determinante, 41
 - Faltungssatz der
 - Laplacetransformation, 107
 - Fourier Entwicklungssatz, 60
 - Injektivität der Laplacetransformation, 101
 - Integraldarstellungen der Besselfunktionen, 86
 - Rechenregeln für Matrizen, 28

- Rechenregeln
 - der Laplacetransformation, 96
- Vektoroperationen, 38
- von Picard-Lindelöf, 70
- zu Nullstellen des charakteristischen Polynoms, 63, 67
- zum Fundamentalsystem, 74, 76
- zum Kreuzprodukt, 38
- zum Polynomzugverfahren, 89
- zur Wronskideterminante, 73
- Schnittwinkel, 24
- Schwingkreis
 - elektrischer, 62
 - gekoppelter, 77
- Schwingung
 - gedämpfte, 65
 - ungedämpfte, 65
- Schwingungsdifferentialgleichung, 64
- singuläre Matrix, 34
- Skalarprodukt, 18
- Skalarprodukt auf $\mathcal{T}_n$, 48
- Spaltenvektoren, 4
- Spiegelung, 31
- stückweise stetig, 55
- Teilraum, 6
- transponierte Matrix, 29
- unitäre Matrix, 31
- unitärer Vektorraum, 20
- Unterraum
 - affiner, 6
 - Basis, 14
 - Dimension, 16
 - linearer, 6
 - von A aufgespannt, 12
- Variation der Konstanten, 80
- Vektor, 3
 - Betrag, 18
 - Länge, 18
 - orthogonaler, 20
- Vektorprodukt, 38
- Vektorraum, 5
 - affiner Unterraum, 6
 - Basis, 14
 - Dimension, 16
 - euklidisch, 20
 - linearer Unterraum, 6
 - normiert, 20
 - unitär, 20
- Verkettung, 27
- Verschiebungssatz, 96
- Winkel, 20, 24
- Wronskideterminante, 73