

Mathematik I (i933)

JENS FLEMMING (HTW Dresden)

16.09.2026

Inhaltsverzeichnis

0	Vorwort	1
1	Motivation	2
1.1	Mathematik vs. Rechnen	2
1.2	Mathematik als Sprache	2
1.3	Computer	3
1.4	“KI”	3
1.5	Mittel gegen Komplexität	4
1.6	“Soft-Skills”	5
1.7	Mathematik im Studium	6
2	Schulmathematik	9
2.1	Online-Kurs OMB+	9
2.2	Themen, die als bekannt vorausgesetzt werden	9
2.3	Themen, die wiederholt und vertieft werden	10
2.4	Nicht behandelte Themen	10
3	Computereinsatz	11
3.1	Überblick	12
3.2	Python	16
3.3	SymPy	24
3.4	NumPy	27
4	Elementare Begriffe	32
4.1	Aussagen	33
4.2	Abbildungen	40
4.3	Zahlenfolgen	46
4.4	Funktionen	54
5	Vektoren und Matrizen	80
5.1	Formalismus	81
5.2	Vektorgeometrie	93
5.3	Lineare Abbildungen	101
5.4	Lineare Unterräume und Untermannigfaltigkeiten	107
5.5	Lineare Gleichungssysteme	115
5.6	Koordinatentransformationen	130
5.7	Ausblick: Eigenwerte und Eigenvektoren	146
6	Differentialrechnung	151
6.1	Idee und Begriffe	152
6.2	Vorgehen	163
6.3	Taylor-Polynome	167

	6.4 Fehlerrechnung	172
7	Integralrechnung	177
	7.1 Flächeninhalt	178
	7.2 Normalbereiche	183
	7.3 Praktische Integration	193
8	Wahrscheinlichkeitsrechnung	200
9	Statistik	201
10	Übung 1 (Grundlegendes)	202
	10.1 Logik, Argumentation, Begriffsbildung	202
	10.2 Wiederholung Schulmathematik	202
11	Übung 2 (Python und Jupyter)	203
	11.1 Jupyter	203
	11.2 Python	203
12	Übung 3 (Aussagen, Abbildungen, Folgen)	205
	12.1 Aussagen	205
	12.2 Abbildungen	205
	12.3 Zahlenfolgen	205
13	Übung 4 (Funktionen)	208
	13.1 Definitionsbereich	208
	13.2 Visualisierung	208
	13.3 Grenzwerte und Stetigkeit	209
14	Übung 5 (Matrizen, Vektorgeometrie)	211
	14.1 Rechnen mit Matrizen	211
	14.2 Rechnen mit Vektoren	212
	14.3 Skalarprodukt und orthogonale Projektion	212
15	Übung 6 (lineare Abbildungen, Basen)	214
	15.1 Lineare Abbildungen	214
	15.2 Unterräume und Basen	214
16	Übung 7 (lineare Gleichungssysteme)	216
	16.1 Formulieren von LGS	216
	16.2 Lineare Unabhängigkeit	217
	16.3 Geraden und Ebenen	218
17	Übung 8 (Koordinatentransformationen)	219
	17.1 Aufgabe 8.1 (KTA)	219
	17.2 Aufgabe 8.2 (KTA)	219
	17.3 Aufgabe 8.3 (KTA)	219
18	Übung 9 (Differentialrechnung)	220
	18.1 Aufgabe 9.1 (KTA)	220
	18.2 Aufgabe 9.2	220
	18.3 Aufgabe 9.3 (KTA)	220
	18.4 Aufgabe 9.4 (KTA)	220
	18.5 Aufgabe 9.5 (KTA)	220
19	Übung 10 (Integralrechnung)	222
	19.1 Stammfunktionen	222

	19.2	Flächeninhalte	223
	19.3	Aufgabe 10.6	223
20		Download	225
21		Organisatorisches	226
	21.1	Ablauf der Veranstaltung	226
	21.2	Hinweise zum Skript	227
	21.3	Prüfung	228
	21.4	Arbeitsplan	228
	21.5	Änderungshistorie	230

0 Vorwort

Dieses Vorlesungsskript ist die Arbeitsgrundlage für die Veranstaltung “Mathematik I” (Modul i933) für den Fernstudiengang Vermessung und Geoinformatik an der HTW Dresden ab Wintersemester 2026/2027. Hinweise zum Ablauf der Veranstaltung, insbesondere zum Umgang mit diesem Skript, finden Sie unter Organisatorisches.

Technischer Unterbau ist MyST Markdown mit einem eigenen, auf dem Standard-Book-Theme aufbauenden Theme. Das MyST-Projekt ist noch recht jung und etliche Features sind im Beta-Stadium. Hier und da ist also mit Ecken und Kanten zu rechnen. Fragen und Anregungen zur Technik gern an Jens Flemming.

1 Motivation

Wozu heute noch Mathematik? Das macht doch alles der Computer oder “die KI”.

1.1 Mathematik vs. Rechnen

Mathematik wird oft mit Rechnen verwechselt oder gleichgesetzt. Richtig ist: **Mathematik ist die Kunst das Rechnen zu vermeiden!** Mathematik strebt einerseits nach möglichst einfachen Lösungswegen für praktische Aufgabenstellungen und organisiert andererseits die trotzdem oft notwendige Rechenarbeit heute so, dass sie **mit dem Computer** erledigt werden kann.

Manuelles Rechnen ist für Ingenieure heute kaum mehr relevant. Lediglich für den Aufbau eines gewissen Grundverständnisses zu einem zu lösenden Problem kann es sinnvoll sein, einige einfache Spezialfälle per Hand auszurechnen. Alles, was einigermaßen praxisnah ist, ist zu komplex zum manuellen Rechnen und wird dem Computer überlassen.

Beispiel

Theoretisch kann man die zu bewegenden Erdmassen beim Bau von Verkehrswegen manuell durch Volumenberechnungen ermitteln. Schon aus der Tatsache, dass dazu Geländeoberfläche und geplanter Trassenverlauf als mathematische Formel gegeben sein müssten, folgt, dass dieser Gedanke praktisch nicht umsetzbar ist. Vielmehr werden Geländeoberfläche und Trassenverlauf nur als CAD-Modell bestehend aus Splines, Bézier-Kurven, NURBS usw. im Computer existieren, sodass auch die hinreichend genaue Volumenbestimmung nur einer Computerrechnung zugänglich ist.

Aus dem Beispiel wird klar: Manuelles Rechnen müssen nur noch wenige Ingenieure beherrschen (Grundlagenarbeit, Software-Entwicklung,..). **Solide Fähigkeiten beim Lösen von Aufgabenstellungen mit dem Computer sind hingegen zwingend für jeden Ingenieur.**

1.2 Mathematik als Sprache

Die Mathematik dient ganz Wesentlich als Sprache in Technik, Naturwissenschaften und Wirtschaftswissenschaften; und das in zweierlei Hinsicht:

1. Mit den Mitteln der Mathematik werden **Modelle der realen Welt** formuliert. Solche mathematischen Modelle ermöglichen einerseits eine genaue Untersuchung des betrachteten Problems. Andererseits ist die mathematische Formulierung die Grundlage für den Einsatz von ebenfalls mit mathematischen Mitteln formulierten Lösungsverfahren.

2. Soll eine Aufgabenstellung mit dem Computer gelöst werden, so muss die Aufgabe in eine **für Computer verständliche Form** gebracht werden. Computer verstehen nur ganz klare und präzise Anweisungen, sodass man üblicherweise die mathematische Sprache als gemeinsames Ausdrucksmittel von Mensch und Computer wählt.

1.3 Computer

Landläufig herrscht die Meinung vor, dass Computer immer richtig rechnen. Das stimmt leider nicht, wie wir im Laufe der Veranstaltung gelegentlich sehen werden. Wichtige Fähigkeit von Ingenieuren ist heute, zu **erkennen, in welchen Situationen der Computer vertrauenswürdige Ergebnisse liefert** und wann nicht! Diese Problematik werden wir in dieser einführenden Veranstaltung nur am Rande streifen; für eine Vertiefung sei auf das Teilgebiet der numerischen Mathematik verwiesen.

Computer sind grundsätzlich als “dumm” zu betrachten. Sie tun exakt das, was man ihnen sagt (und nicht dass, was man meint, ihnen gesagt zu haben). Ob der Computer das Richtige tut und ob wir den Ergebnissen trauen können, liegt allein in unserer Hand. Der Mensch muss die Prozesse und Hintergründe verstehen, muss einordnen und bewerten. **Der Computer übernimmt lediglich das Rechnen, nicht das Denken!**

1.4 “KI”

Der Begriff “künstliche Intelligenz” ist irreführend. Tatsächlich handelt es sich bei den heute unter diesem Begriff verstandenen Techniken um **simulierte Intelligenz**. Diese System können **nicht denken, keine zielgerichtete Kreativität** entfalten. Sie interpolieren und (in weit geringerem Maß) extrapolieren lediglich vorhandenes Faktenwissen. Aus dieser Feststellung folgen zwei wichtige Schlüsse für das Dasein als Ingenieur:

1. Wer nur Standardaufgaben lösen muss/möchte, kommt mit KI oft ans Ziel. Wer wirklich Neues schaffen möchte, muss denken und verstehen.
2. Von KI gelieferte Problemlösungen können beliebig grobe oder auch sehr subtile Fehler enthalten. Gerade im oft sicherheitskritischen Ingenieurbereich müssen aber nachweisbar richtige Lösungen her. Mindestens müssen die Lösungen nach bestem Wissen und Gewissen korrekt sein; ein Anspruch, den KI-basierte Lösungen ohne detaillierte Prüfung nicht erfüllen.

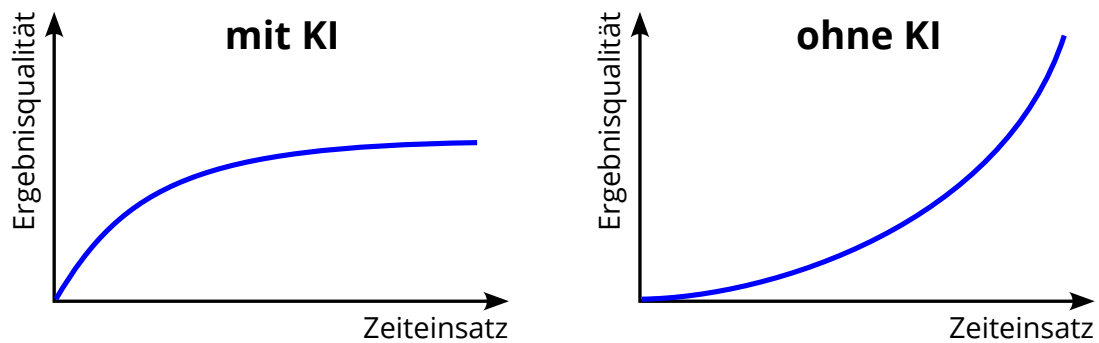
Wichtig: KI arbeitet nicht deterministisch, also nicht Schritt für Schritt nachvollziehbar, sondern statistisch, indem sie eine mit hoher Wahrscheinlichkeit passende Antwort “würfelt”. Die Wahrscheinlichkeiten basieren dabei auf der der KI beim Training zugänglichen Wissensbasis. Insbesondere Nischenthemen werden durch Mainstream-Themen überstrahlt. KI-Antworten dazu gleichen einer zufälligen, meist sinnlosen oder ganz falschen Vermischung des eigentlichen Themas mit den Mainstream-Themen im Umfeld.

1 Motivation

Dabei sind die Antworten sauber und überzeugend formuliert, halten einer inhaltlichen Detailprüfung jedoch nicht stand.

Zur Ideenfindung ist KI durchaus geeignet. Die Abschließende Problemlösung muss jedoch vom Ingenieur vollständig durchdacht und verstanden sein.

Der Einsatz von KI liefert schnelle Erfolge, die langfristig aber hinter den durch Denken und Verstehen erzielbaren Resultaten zurückbleiben.



KI liefert schnell akzeptable Ergebnisse auf mittelmäßigem Niveau. Selbst denken und verstehen benötigt mehr Zeit, liefert aber langfristig deutlich hochwertigere Resultate.

1.5 Mittel gegen Komplexität

Neben der direkten Funktion als Sprache dient Mathematik auch dem Strukturieren von Denkprozessen. Die Mathematik schafft vereinfachte Formulierungen der realen Welt und abstrahiert komplexe Zusammenhänge und Vorgänge in klar definierten Begriffen und Strukturen. Durch die Gewöhnung an abstrakte Begriffe (als Ersatz für die viel zu komplexe Realität) werden Denkprozesse vereinfacht oder überhaupt erst möglich.

Beispiel (reelle Zahlen)

Reelle Zahlen wie $\sqrt{2}$ oder die Kreiszahl π erscheinen uns nach jahrelangem Training in der Schule ganz natürlich und wir arbeiten täglich mit ihnen. In der Realität haben sie aber im Gegensatz zu rationalen Zahlen keine Entsprechung; im Gegenteil, sie sind ein sehr abstraktes Konzept, welches die Formulierung vieler technischer Zusammenhänge deutlich vereinfacht (Video 130, Notizen dazu).

Abstraktion findet auch außerhalb der Mathematik statt, ist aber in der Mathematik zentrales Arbeitsmittel mit Wirkung in andere Wissenschaftsgebiete hinein.

Beispiel (Farben)

Der Begriff “Farbe” ist so alltäglich wie abstrakt. Was bedeutet “die Kirsche ist rot”? Wie erklärt man jemandem eine konkrete Farbe ohne auf Gegenstände dieser Farbe als Beispiele zu zeigen?

Von klein auf haben wir nur durch ständiges Zeigen auf Gegenstände und das Sagen des dazugehörigen Farbwortes ein Verständnis für den Begriff “Farbe” entwickelt. Nie haben wir eine konkrete Definition gehört. Erst in der Physik kommt man einer konkreten Definition nahe (Farbe = Wellenlänge des Lichts), die aber nicht die menschliche Farbwahrnehmung erklären kann: Warum wird eine Mischung aus rotem und grünem Licht in gleicher Weise wahrgenommen wie rein gelbes Licht?

Aus mathematischer Sicht ist “rot” eine Äquivalenzklasse bezüglich einer Äquivalenzrelation: zwei Gegenstände *haben die gleiche Farbe* (Definition!), wenn deren Farbwahrnehmung im Gehirn identisch ist. Mit welchem Wort (“rot”, “grün”, “gelb” usw.) man die einzelnen Gruppen von Gegenständen gleicher Farbwahrnehmung dann bezeichnet, ist eine andere Frage.



Das mathematische Konzept der Äquivalenzklassen ermöglicht völlig analog zum kleinkindlichen Lernprozess implizite Begriffsdefinitionen. Man kann so auch mit komplexen Situationen umgehen, die keine (hinreichend einfache) explizite Begriffsdefinition zulassen. Wir sehen hier: **Scheinbar abstrakte mathematische Konzepte sind oftmals direkte Formalisierungen alltäglichen Denkens!**

1.6 “Soft-Skills”

Beschäftigt man sich mit Mathematik, erlernt man ganz nebenbei viele Fähigkeiten, die im Privaten wie Beruflichen sehr nützlich sind:

1 Motivation

- genaues Analysieren,
- streng logisches Argumentieren,
- exakte Ausdrucksweise,
- scheinbar offensichtliches immer wieder hinterfragen,
- das eigene Denken und Verstehen im Grenzbereich testen.

Das Üben dieser Fähigkeiten in der klaren, strukturierten Welt der Mathematik rüstet für den Weg in diffusere Gebiete.

Beispiel (logisches Argumentieren)

Hase und Kröte liefern sich ein Wettrennen über 100 Meter. Der Hase sieht die Sache entspannt und lässt die Schildkröte 10 Meter weiter vorn starten. Beide laufen gleichzeitig los. Wer wird gewinnen?

Dazu zwei Argumentationen:

1. Der Hase wird gewinnen, da er viel schneller läuft als die Schildkröte.
2. Die Schildkröte wird gewinnen. Wenn der Hase bei der Startposition der Schildkröte ankommt, ist diese bereits ein Stück gelaufen, liegt also vorn. Kommt der Hase dann wiederum an dieser Position an, ist die Schildkröte wieder ein Stück weiter. Jedesmal, wenn der Hase an der vorherigen Position der Schildkröte ankommt, ist die Schildkröte schon wieder ein Stück weiter. Die Schildkröte liegt immer vorn, wird also gewinnen.

Wer hat recht? Wo ist der Fehler in der anderen Argumentation? Siehe Übung 1.

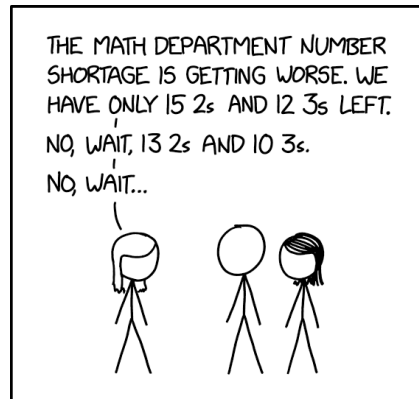
Beispiel (scheinbar offensichtliches hinterfragen)

Was heißt eigentlich "gleich viel"? Hat ein gesunder Mensch gleich viele Finger wie Fußzehen? Gibt es gleich viele gerade wie ungerade Zahlen? Gibt es gleich viele Brüche (d.h. rational Zahlen) wie es natürliche Zahlen gibt? Siehe Übung 1.

1.7 Mathematik im Studium

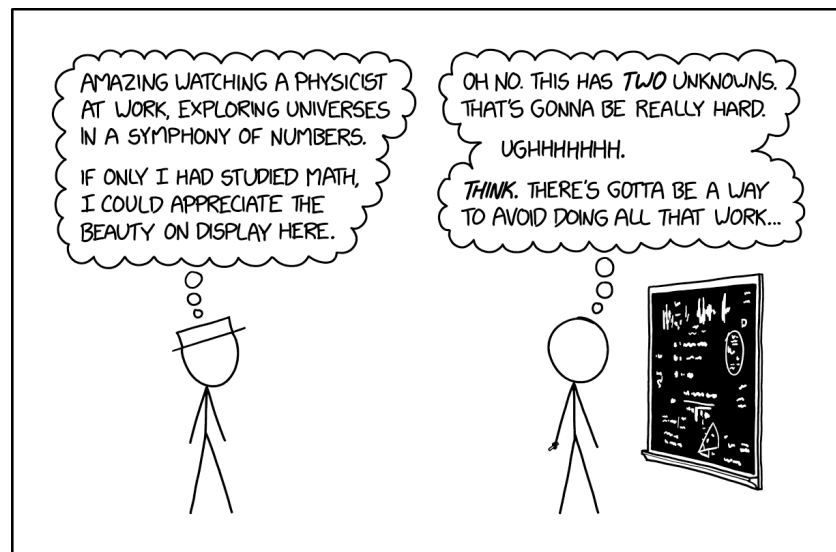
Im weiteren Studienverlauf wird die Mathematik in vielen Modulen eine tragende Rolle spielen. Einige Module aus den nächsten Semestern und die dort benötigten Themen aus Mathematik I und II sind:

- Statistik und Ausgleichungsrechnung (Grundlagen Wahrscheinlichkeitsrechnung und Statistik, Anwendungen der Differentialrechnung, Matrizen und Vektoren, lineare Gleichungssysteme).



"10 minutes ago we were down to only 2 0s!" "How many do we have now?" "I ... don't know!!" Quelle: Randall Munroe, xkcd.com/3009

- Landesvermessung und Satellitengeodäsie (Matrizen und Vektoren, Differentialrechnung für Funktionen mehrerer Veränderlicher).
- Remote Sensing (Matrizen und Vektoren).
- Positionierung und Navigation (Matrizen und Vektoren, Anwendungen der Differentialrechnung).
- KI und Geodata-Science (Matrizen und Vektoren, Grundlagen der Statistik, Differentialrechnung für Funktionen mehrerer Veränderlicher).
- Geodätische Auswertetechniken (Grundlagen der Statistik).
- Laserscanning und Punktwolkenverarbeitung (Grundlagen der Statistik).



I could type this into a solver, which MIGHT help, but would also mean I have to get a lot of parentheses right... Quelle: Randall Munroe, xkcd.com/2207

2 Schulmathematik

In diesem Kapitel erhalten Sie ein Überblick darüber, welche Themen der Schulmathematik in dieser Veranstaltung nochmals behandelt werden und welche eigenständig zu wiederholen sind, sofern Sie die Notwendigkeit dazu verspüren.

2.1 Online-Kurs OMB+

Die HTW Dresden unterstützt und nutzt den von zahlreichen Bildungseinrichtungen gemeinsam getragenen kostenfreien Online-Kurs OMB+ zur Auffrischung des Schulwissens im Bereich Mathematik. Dort gibt es einerseits das notwendige Faktenwissen, andererseits aber auch Aufgaben zum Test der eigenen Fähigkeiten. Große Teile des Materials sind ohne Anmeldung nutzbar. Einige interaktive Aufgaben werden erst nach kostenloser Registrierung nutzbar.

Nutzen Sie dieses Angebot, um parallel zur Veranstaltung Rückstände und Vergessenes aufzuarbeiten! In diesem Kapitel finden Sie jeweils Links zu konkreten Abschnitten des Online-Kurses.

2.2 Themen, die als bekannt vorausgesetzt werden

Folgende Themen werden nicht in der Lehrveranstaltung behandelt, sondern als aus der Schulmathematik bekannt vorausgesetzt:

- Kapitel IA in OMB+ (Mengen, Zahlen)
- Kapitel IB in OMB+ (Potenzen, Proportionalität)
- Kapitel II in OMB+ (Gleichungen in einer Unbekannten)
- Kapitel III in OMB+ (Ungleichungen in einer Variablen)
- Kapitel V in OMB+ (Geometrie)
- Kapitel VI in OMB+ (elementare Funktionen) (einige Teilthemen werden in der Veranstaltung wiederholt)
- Kapitel IX in OMB+ (Vektorgeometrie) (einige Teilthemen werden in der Veranstaltung wiederholt)
- Kapitel XII.4 in OMB+ (Logik und Mengenlehre - Mengen)
- Kapitel XII.6 in OMB+ (Logik und Mengenlehre - Kartesische Produkte)

Nebenbemerkung

Die Darstellung der Themen in OMB+ ist recht detailliert und teils umfangreicher

als in der Veranstaltung benötigt. Verschaffen Sie sich deshalb zunächst nur einen Überblick über die genannten Themen. Sollten Sie im Verlauf der Veranstaltung Nachholbedarf feststellen, sollten Sie die entsprechenden Abschnitte von OMB+ genauer durcharbeiten.

2.3 Themen, die wiederholt und vertieft werden

Folgende Themen werden in der Veranstaltung nochmals vollständig eingeführt, behandelt und vertieft. In einigen Ländern und Bundesländern gehören diese zum Schulstoff, sodass unter Umständen schon Vorwissen vorhanden ist.

- Kapitel IV in OMB+ (lineare Gleichungssysteme)
- Kapitel VII in OMB+ (Differentialrechnung)
- Kapitel VIII in OMB+ (Integralrechnung)
- Kapitel X in OMB+ (Stochastik)
- Kapitel XII in OMB+ (Logik und Mengenlehre), sofern benötigt (außer XII.4, XII.6),
- Kapitel XIII in OMB+ (Wahrscheinlichkeitstheorie)

2.4 Nicht behandelte Themen

OMB+ enthält auch Themen, die in der Veranstaltung nicht von Bedeutung sind:

- Kapitel IX in OMB+

3 Computereinsatz

Grundlagen der Computernutzung in der Mathematikausbildung.

3.1 Überblick

Der Einsatz des Computers zum Lösen von in der Sprache der Mathematik formulierten Aufgabenstellungen kann auf unterschiedliche Weise erfolgen. Dabei spielt einerseits die Auswahl geeigneter Software eine Rolle; andererseits muss man sich im Klaren sein über die Fähigkeiten und Unfähigkeiten des Computers.

Einschränkungen bei der Arbeit mit dem Computer

Viele mathematische Probleme können nicht exakt mit dem Computer gelöst werden, sondern nur näherungsweise. Die Gründe dafür liegen in der Arbeitsweise von Computern:

- Jede Eingabe kann nur aus endlich vielen Zahlen bestehen. Computer können also beispielsweise nicht direkt mit stetigen Funktionen (Potenzfunktionen, Sinus,...) oder unendlichen Zahlenfolgen arbeiten.
- Auch Ausgaben/Ergebnisse können nur aus endlich vielen Zahlen bestehen.
- Es sind stets nur endlich viele Rechenschritte möglich. Beispielsweise können im Allgemeinen keine unendlichen Summen mit dem Computer berechnet werden.

Ein weiterer zu beachtender Punkt ist, dass Computer nur einen sehr begrenzten Zahlenvorrat zur Verfügung haben. Üblicherweise kennen sie nur $2^{64} \approx 2 \cdot 10^{18}$ verschiedene Zahlen. Deshalb werden die meisten Zahlen durch den Computer auf eine vom Computer verarbeitbare Zahl gerundet, ohne dass dies für den Nutzer sofort ersichtlich ist. Je nach Kontext müssen die mitunter seltsam anmutenden Konsequenzen mit bedacht werden:

- Beispielsweise kann die Zahl 0.1 nicht direkt in Berechnungen verwendet werden. Stattdessen rechnet der Computer mit

$$0.1000000000000000055511151231257827021181583404541015625. \quad (3.1)$$

Dies kann zu unerwarteten Ergebnissen führen, wie das folgende kleine Python-Programm zeigt (die Programmiersprache Python wird im nächsten Unterkapitel behandelt):

```
print(0.1 * 0.1 - 0.01)  # sollte exakt 0 ergeben
```

```
1.734723475976807e -18
```

- Gelegentlich erhält man deshalb völlig falsche Ergebnisse:

```
print(1 + 10.5 ** 20 - 10.5 ** 20)  # sollte exakt 1 ergeben
```

```
0.0
```

Nebenbemerkung

Die beiden Beispiele können Sie leicht selbst ausprobieren. Öffnen Sie in Ihrem Webbrowser die JavaScript-Konsole (in Firefox z.B. durch Drücken von ++, in Chromium durch ++) und geben Sie dort $0.1 * 0.1 - 0.01$ bzw. $1 + 10.5 ** 20 - 10.5 ** 20$ ein. Die Ergebnisse sollte aussehen wie hier im Skript.

Die Auswirkungen und der Umgang mit solchen Effekten sind Gegenstand der numerischen Mathematik. In unserer Veranstaltung gehen wir auf dieses wichtige Thema jedoch nicht näher ein, sondern belassen es bei gelegentlichen Hinweisen auf in dieser Hinsicht problematische Situationen.

Symbolisches vs. numerisches Rechnen

Für das Ausführen von Berechnungen mit dem Computer gibt es zwei grundlegend verschiedene Herangehensweisen:

- Beim **symbolischen Rechnen** werden mathematische Formeln durch den Computer in der gleichen Weise umgeformt, kombiniert und gelöst wie man es beim manuellen Rechnen tun würde. Dadurch erhält man exakte Ergebnisse, die wiederum als Formeln vorliegen. Bei der verwendeten Software spricht man auch von **Computer-Algebra-Systemen (CAS)**.
- Beim **numerischen Rechnen** werden nur Grundrechenarten verwendet. An vielen Stellen des Rechenweges werden gerundete oder nur näherungsweise ermittelte Zwischenergebnisse verwendet. Als Endergebnis stehen nur Zahlen zur Verfügung, bei denen nicht mehr erkennbar ist, wie diese entstanden sind (z.B. 1.41421356 statt $\sqrt{2}$).

Im technischen und wissenschaftlichen Bereich wird praktisch nur numerisch gerechnet. Gründe:

- Numerische Berechnungen sind viel schneller, da Computer für diese Art von Berechnungen gebaut sind. Üblich sind heute mehrere Milliarden Grundrechenoperationen pro Sekunde.
- Symbolisches Rechnen liefert bei komplexen Aufgabenstellungen meist keine Lösung, da die Lösungen nicht explizit als Formeln darstellbar sind.

Gelegentlich verwendet man CAS für kleine, klar abgegrenzte Teilschritte, bei denen das Vorliegen einer expliziten Formel in den Folgeschritten (Geschwindigkeits-)Vorteile bringt.

Beispiel (automatisches Differenzieren)

Grundlage von “KI” sind künstliche neuronale Netze. Das sind aus Tausenden oder Millionen einfacher Grundfunktionen zusammengesetzte nichtlineare Funktionen, welche Millionen durch ein Optimierungsverfahren zu bestimmende Parameter enthalten. Für die Durchführung der Optimierung werden die Ableitungen des neuronalen Netzes bezüglich der Parameter benötigt (den Begriff “Ableitung” führen wir im Laufe der Veranstaltung ein). Diese können gut symbolisch berechnet werden, was Näherungsfehler vermeidet. Außerdem beschleunigt das Vorliegen einer konkreten Formel die Auswertung der Ableitungen während der Optimierung. Verfahrensbedingt müssen diese Ableitungen millionenfach ausgewertet werden.

Das symbolische Berechnen der Ableitungen von Funktionen, die aus einfachen Grundfunktionen zusammengesetzt sind, nennt man auch **automatisches Differenzieren**. Die eigentliche Optimierung erfolgt dann numerisch, da das Problem aufgrund seiner Komplexität nicht symbolisch lösbar ist.

Programmcode vs. grafische Oberfläche

Für praktisch alle Standardaufgaben im Ingenieurwesen gibt es Software mit grafischer Benutzeroberfläche (kurz: GUI für “graphical user interface”). Diese haben jedoch einige Nachteile:

- Möglich ist nur, was die GUI vorsieht. Erweiterungen oder Änderungen sind für Nutzer nur schwer oder gar nicht umsetzbar.
- Abläufe lassen sich schlecht automatisieren, da diese aus Folgen von Mausklicks bestehen.
- Für Aufgabenstellungen, die erstmalig oder nur selten auftreten, existiert keine Software mit GUI (oder ist extrem teuer).

Aus diesen Gründen ist der Einsatz einer geeigneten (!) Programmiersprache auch im Ingenieurbereich für zahlreiche Aufgaben deutlich vorteilhafter. Insbesondere

- sind die Anwendungsmöglichkeiten unbeschränkt,
- lassen sich alle Abläufe ohne große Anpassungen automatisieren,
- kann bei Bedarf im Nachgang eine GUI hinzugefügt werden.

Hinzu kommt, dass sich Programmcode einfacher versionieren lässt und gemeinsames Bearbeiten (auch in Echtzeit) problemlos möglich ist.

Üblich im Ingenieurbereich sind:

- Python/Jupyter (wird in dieser Veranstaltung verwendet),
- Octave oder Matlab (wird zunehmend durch Python verdrängt).

Nebenbemerkung

Die Arbeit mit Programmcode statt grafischer Oberfläche ist keineswegs komplizierter. Lediglich der Funktionsumfang ist größer. Aufgrund der Vielzahl von Möglichkeiten ist es nicht zweckmäßig für jede Funktion einen Button oder einen Menüeintrag zu bieten. Stattdessen werden alle Aktionen durch Textbefehle ausgelöst.

3.2 Python

In dieser Veranstaltung nutzen wir die Programmiersprache Python um dem Computer seine Aufgaben mitzuteilen. Werkzeuge aus dem Jupyter-Ökosystem ermöglichen die Python-Programmierung in einer sehr komfortablen grafischen Umgebung.

Warum Python?

Python ist eine moderne, kostenfreie und quelloffene (“open source”) Programmiersprache. Sie entstand in den 1990er Jahren. Die erste offizielle Version wurde 1994 veröffentlicht. Autor und BDFL (“benevolent dictator for life”) ist Guido van Rossum.

Python-Code ist sehr gut lesbar und enthält im Gegensatz zu den meisten anderen Programmiersprachen nur wenige “kryptische” Zeichen und Zeichenketten. Viele technische Aspekte des Programmierens werden von Python erledigt ohne den Nutzer damit zu konfrontieren. Mit Python kann man Software aus allen Bereichen entwickeln; von einfachen Skripten für die tägliche Arbeit bis zu vollwertigen Desktop- oder Webanwendungen. Tausende Erweiterungen ermöglichen zügiges und effizientes Arbeiten.

Als freie Software verfügt Python über eine riesige Online-Community, die mit Rat und Tat zur Seite steht. Kaum eine Frage, die nicht schon irgendwo beantwortet wurde.

Python ist für alle Plattformen verfügbar: Linux, macOS, Windows und viele andere. Python steht beispielsweise hinter großen kommerziellen Projekten wie dem Youtube-Player, läuft eventuell aber auch in der Steuereinheit Ihrer Waschmaschine (vgl. Micro-Python).

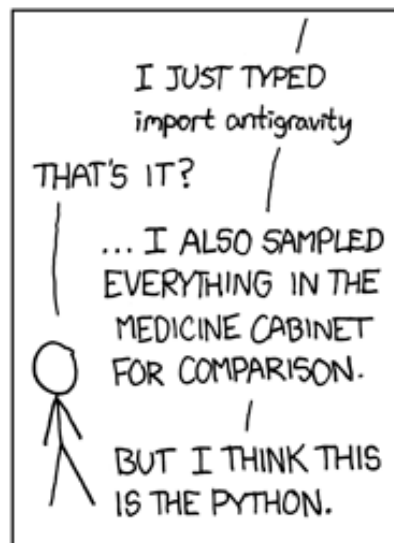
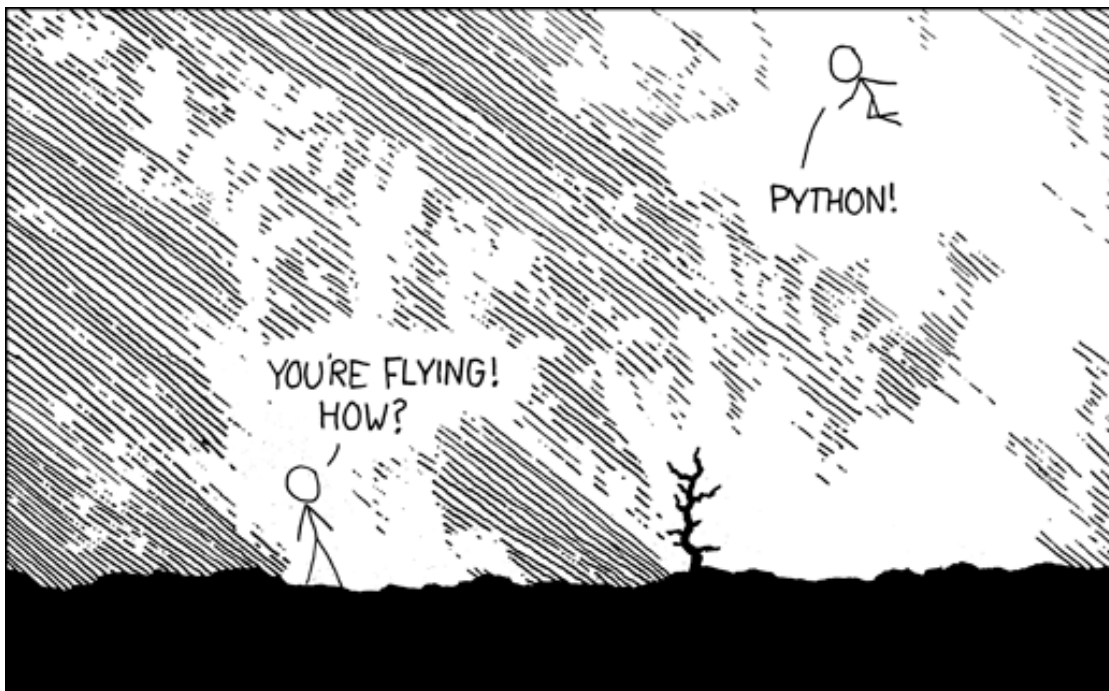
Nebenbemerkung

Es gibt zwei Varianten von Python: Python 2 und Python 3. Entsprechender Quellcode ist nicht kompatibel, d.h. für Python 2 geschriebene Programm funktionieren nicht mit Python 3. Wir verwenden hier ausschließlich Python 3. Python 2 wird von den Entwicklern seit Januar 2020 als veraltet eingestuft und nicht mehr weiterentwickelt.

Jupyter

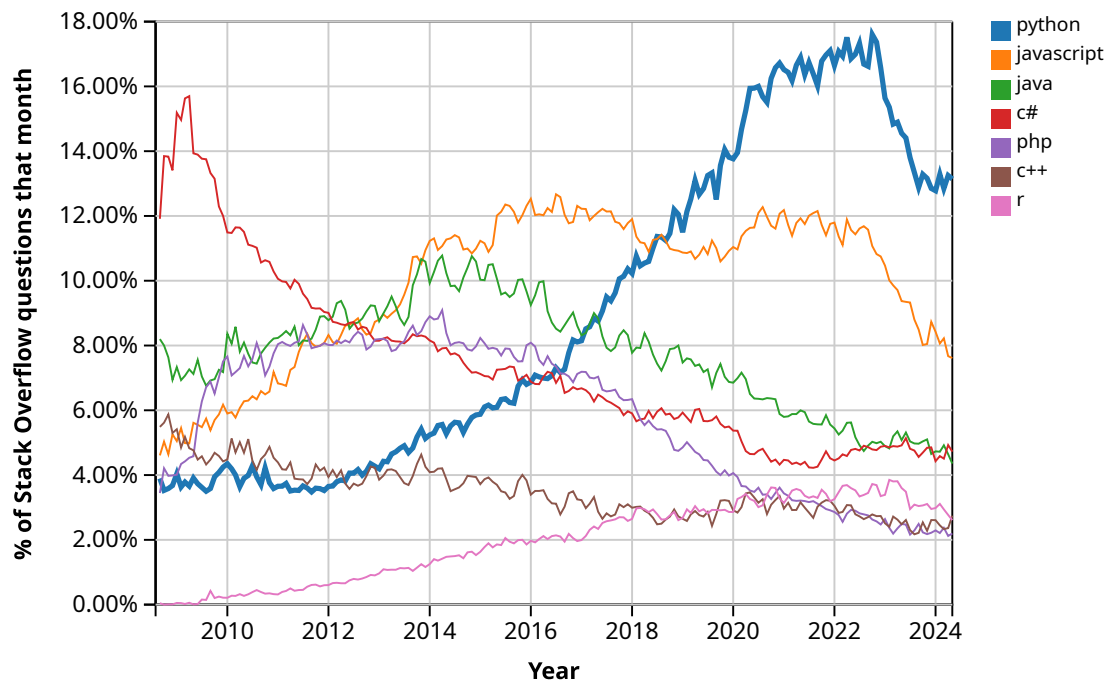
Das Jupyter-Ökosystem ist eine Sammlung von Werkzeugen für die Python-Programmierung mit dem Schwerpunkt auf Nutzerfreundlichkeit und leichte Bedienbarkeit für Anfänger. Jupyter erlaubt das Programmieren mit Python im Webbrowser. Ausgaben, wie beispielsweise komplexe und interaktive Visualisierungen, können direkt zwischen dem zugehörigen Quellcode platziert werden. Alles befindet sich in einem einzigen Dokument: Code, Ausgaben, erklärender Text, Bilder, Formeln,...

JupyterLab steht in vielen Formen zur Verfügung (vgl. Übung 2):

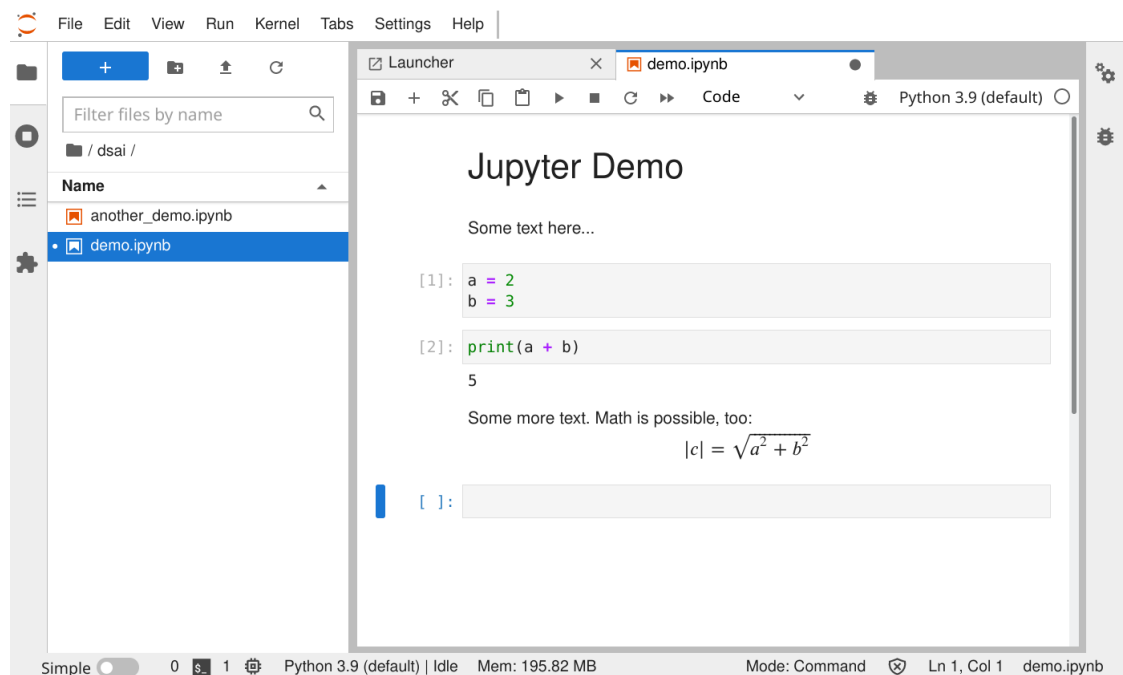


I wrote 20 short programs in Python yesterday. It was wonderful. Perl, I'm leaving you. Source: Randall Munroe, xkcd.com/353

3 Computereinsatz



Fragen zu Programmiersprachen auf Stack Overflow. Quelle: Stack Overflow Trends (vom Autor bearbeitet)



JupyterLab ist das am weitesten verbreitete Mitglied des Jupyter-Ökosystems. Es bringt Python in den Webbrowser.

- als rein browser-basierte Anwendung ohne Installation (Funktionsumfang geringfügig eingeschränkt),
- als lokale Installation auf dem eigenen Computer,
- als Installation auf Servern der Hochschule,
- auf Servern anderer Organisationen.

Welche Variante genutzt wird, ist für die Veranstaltung nebensächlich. Außerdem kann Python auch ganz ohne Jupyter genutzt werden.

Python-Crashkurs

Erwähnen hier als ersten Einstieg in Python beispielhaft einige Grundfunktionen. Weitere Möglichkeiten folgen im Laufe der Veranstaltung nach Bedarf. Für das weiterführende Selbststudium sind das Kapitel Python Crash Course sowie die darauf folgenden Kapitel in Data Science and Artificial Intelligence for Undergraduates geeignet.

Taschenrechner

```
# Punkt vor Strich
1 + 2 * 3
```

7

```
# Potenzen
2 ** 8
```

256

```
# Wurzeln
2.5 ** 0.5
```

1.5811388300841898

```
# Division
5 / 2
```

2.5

```
# ganzzahlige Division
5 // 2
```

2

3 Computereinsatz

```
# Rest bei ganzzahliger Division  
5 % 2
```

```
1
```

Bessere Ausgaben

```
# einfacher Text  
print('irgendein Text')
```

```
irgendein Text
```

```
# Text mit Kommazahl  
print(f'Wurzel aus 2 ist {2.0 ** 0.5}.')
```

```
Wurzel aus 2 ist 1.4142135623730951.
```

Mehr Infos zu solchen Ausgaben: f-String und Minisprache für Stringformatierung.

```
# Text mit gerundeter Kommazahl  
print(f'Wurzel aus 2 ist {2.0 ** 0.5:.3f}.')
```

```
Wurzel aus 2 ist 1.414.
```

Variablen

```
zahl = 2.0  
wurzel = zahl ** 0.5  
print(f'Wurzel aus {zahl} ist {wurzel:.3f}.')
```

```
Wurzel aus 2.0 ist 1.414.
```

Es gibt verschiedene Variablentypen (Ganzzahl, Kommazahl, Zeichenkette,...).

```
print(type(2))  
print(type(2.0))  
print(type('Hallo!'))
```

```
<class 'int'>  
<class 'float'>  
<class 'str'>
```

Umwandlungen sind zwischen geeigneten Typen möglich.

```
zeichenkette1 = '123'
zeichenkette2 = '456'
print(zeichenkette1 + zeichenkette2)

zahl1 = int(zeichenkette1)
zahl2 = int(zeichenkette2)
print(f'und als Zahlen: {zahl1 + zahl2}')
```

```
123456
und als Zahlen: 579
```

Listen Listen sind Folgen von Variablen, deren Name einem einheitlichen Schema folgen (Listenname plus Index).

```
liste = ['rot', 'gelb', 'grün', 'blau']
print(liste[0])
print(liste[1])
print(liste[2])
print(liste[3])
print(liste[-1]) # letztes Element
print(liste[1:3]) # Elemente mit Index 1 (2. Element) und 2 (3.
→ Element)
```

```
rot
gelb
grün
blau
blau
['gelb', 'grün']
```

Bedingte Ausführung

```
a = 3
b = 2 # zwei Zahlen (Nutzereingaben oder aus Datei gelesen)

if b != 0:
    print(f'{a} / {b} = {a / b}')
else:
    print('Division durch Null!')
```

```
3 / 2 = 1.5
```

Wichtig: Test auf Gleichheit mit `==`. Das einfache `=` ist der Zuweisungsoperator und kann nicht für Vergleiche verwendet werden.

Wiederholte Ausführung

```
liste = ['rot', 'gelb', 'grün', 'blau']  
  
for farbe in liste:  
    print(farbe)
```

```
rot  
gelb  
grün  
blau
```

```
for i in range(0, 10, 2):  
    print(i)
```

```
0  
2  
4  
6  
8
```

Module und Pakete Die Funktionalität von Python kann durch Module und Pakete erweitert werden. Module sind Sammlungen von zusätzlichen Funktionen. Pakete sind Sammlungen von Modulen. Praktisch besteht bei der Verwendung kein Unterschied zwischen Modulen und Paketen.

```
import math  
  
math.cos(math.pi / 4)
```

```
0.7071067811865476
```

```
import math as m  
  
m.cos(m.pi / 4)
```

```
0.7071067811865476
```

```
from math import cos, pi  
  
cos(pi / 4)
```

```
0.7071067811865476
```

Python verfügt nur über wenige eingebaute Grundfunktionen. Man wird fast immer zusätzliche Module und Pakete benötigen, die ggf. erst installiert werden müssen.

3.3 SymPy

SymPy ist ein Python-Paket für symbolisches Rechnen (Ableitungen, Stammfunktionen, Formeln vereinfachen, Gleichungen lösen,...). Siehe auch Introductory Tutorial in der SymPy-Dokumentation.

Nebenbemerkung

Dieses Kapitel ist keine systematische Einführung, sondern liefert nur einige Beispiele. Alles weitere wird im Verlauf der Veranstaltung nach Bedarf eingeführt und vertieft.

```
import sympy
```

Variablen

SymPy verwendet ganz normale Python-Variablen, die aber einen SymPy-spezifischen Typ haben. Variablen, die später mit SymPy verwendet werden sollen, müssen deshalb zunächst durch SymPy erstellt werden.

```
x = sympy.symbols('x')  # zeige Python -Variable x in Formeln als x an
print(x)
print(type(x))
```

```
x
<class 'sympy.core.symbol.Symbol'>
```

Mehrere Variablen können gemeinsam erstellt werden:

```
x, y, z = sympy.symbols('x y z')

2 * x + y ** 2 + z
```

```
2*x + y**2 + z
```

Formeln vereinfachen

```
x = sympy.symbols('x')
a = (x + 1) ** 2
b = x ** 2 + 2 * x + 1
```

```
sympy.simplify(a - b)
```

```
0
```

```
x = sympy.symbols('x')
sympy.simplify(sympy.cos(x) ** 2 + sympy.sin(x) ** 2)
```

```
1
```

Nullstellen von Polynomen

SymPy liefert Polynomnullstellen samt Vielfachheiten (der Begriff “Vielfachheit” wird im Laufe der Veranstaltung erklärt):

```
x = sympy.symbols('x')
sympy.roots(x ** 3 - 4 * x ** 2 + x + 6, x)
```

```
{3: 1, 2: 1, -1: 1}
```

```
x = sympy.symbols('x')
sympy.roots(x ** 3 - 3 * x ** 2 + 3 * x - 1, x)
```

```
{1: 3}
```

Gleichungen lösen

```
x = sympy.symbols('x')
sympy.solve(x ** 2 - 2, x)  # löst x ** 2 - 2 = 0
```

```
[-sqrt(2), sqrt(2)]
```

Das Ergebnis ist eine Liste (kein Intervall!) von Lösungen.

```
x = sympy.symbols('x')
formel = x ** 4 - 4 * x ** 3 + 7 * x ** 2 - 16 * x + 12
```


3.4 NumPy

NumPy ist ein Paket, welches Python um Funktionen für das schnelle Rechnen mit großen Zahlenmengen (Matrizen, Vektoren) erweitert. NumPy ist im Python-Ökosystem der De-Facto-Standard für numerische Berechnungen wie beispielsweise das Lösen großer linearer oder nichtlinearer Gleichungssysteme.

Nebenbemerkung

Dieses Kapitel ist keine systematische Einführung, sondern liefert nur einige Beispiele. Alles weitere wird im Verlauf der Veranstaltung nach Bedarf eingeführt und vertieft. Eine systematische Einführung in NumPy finden Sie beispielsweise im Kapitel Efficient Computations with NumPy von Data Science and Artificial Intelligence for Undergraduates.

```
import numpy as np
```

Grundlage für die Arbeit mit NumPy sind die NumPy-Arrays, welche zur Darstellung von Vektoren und Matrizen dienen.

NumPy-Arrays erstellen

```
np.array([4, 5, 2, 3]) # aus einer Liste
```

```
array([4, 5, 2, 3])
```

```
np.zeros(4) # eindimensionales Array mit Nullen
```

```
array([0., 0., 0., 0.])
```

```
np.ones((3, 2)) # zweidimensionales Array mit Einsen
```

```
array([[1., 1.],
       [1., 1.],
       [1., 1.]])
```

```
np.eye(5) # Einheitsmatrix
```

```
array([[1., 0., 0., 0., 0.],
       [0., 1., 0., 0., 0.],
       [0., 0., 1., 0., 0.],
       [0., 0., 0., 1., 0.]])
```

```
[0., 0., 0., 0., 1.]])
```

Zugriff auf Array-Elemente

```
a = np.array([[1, 2, 3], [4, 5, 6], [7, 8, 9]])  
a
```

```
array([[1, 2, 3],  
       [4, 5, 6],  
       [7, 8, 9]])
```

```
print(a[0, 0]) # 1. Zeile, 1. Spalte  
print(a[1, 0]) # 2. Zeile, 1. Spalte  
print(a[-1, -1]) # letzte Zeile, letzte Spalte
```

```
1  
4  
9
```

```
a[1:3, :2] # 2. und 3. Zeile, 1. und 2. Spalte
```

```
array([[4, 5],  
       [7, 8]])
```

```
a[0, :] # 1. Zeile, alle Spalten
```

```
array([1, 2, 3])
```

Eigenschaften von Arrays

```
a.shape # Größe (Zeilen, Spalten)
```

```
(3, 3)
```

```
a.dtype # Datentyp der Einträge
```

```
dtype('int64')
```

```
b = np.array([4.1, 6, 3, 5, 4, 3])  
print(b.shape)
```

```
print(b.dtype)
```

```
(6,)
float64
```

Rechnen mit Arrays

```
a = np.array([4.1, 6, 3])
b = np.array([5, 4, 3])
```

```
print(2 * a + b)
```

```
[13.2 16.   9. ]
```

```
print(a * b)  # elementweises Produkt
```

```
[20.5 24.   9. ]
```

```
A = np.array([[1, 2], [3, 4]])
print(A)
```

```
[[1 2]
 [3 4]]
```

```
print(A * A)  # elementweises Produkt
```

```
[[ 1  4]
 [ 9 16]]
```

```
print(np.sin(A))  # Sinus elementweise
```

```
[[ 0.84147098  0.90929743]
 [ 0.14112001 -0.7568025 ]]
```

Arrays vergleichen

```
a = np.array([1, 2, 3])
b = np.array([-1, 3, 3])
```

```
print(a > b)
```

```
[ True False False]
```

```
c = a > b  # Zuweisung (!) des Vergleichsergebnisses an c

print(c.any())  # mindestens ein Eintrag True?
print(c.all())  # alle Einträge True?
```

```
True
False
```

Weitere Möglichkeiten für den Elementzugriff

```
a = np.array([[1, 2, 3], [4, 5, 6], [7, 8, 9]])

a
```

```
array([[1, 2, 3],
       [4, 5, 6],
       [7, 8, 9]])
```

```
a[a > 3]
```

```
array([4, 5, 6, 7, 8, 9])
```

```
a[np.logical_and(a > 3, a < 6)]
```

```
array([4, 5])
```

```
a[a > 4] = 20

a
```

```
array([[ 1,  2,  3],
       [ 4, 20, 20],
       [20, 20, 20]])
```

Zufallszahlen

Für Tests und Simulationen sind Zufallszahlen sehr nützlich. Um Zufallszahlen zu erzeugen, muss zunächst der Zufallszahlengenerator von NumPy mit einem Startwert initialisiert werden. Der Startwert legt die Folge der generierten Zufallszahlen fest, sodass

bei späterer Wiederholung von Rechnungen sehr leicht wieder die gleichen Zufallszahlen verwendet werden können.

```
rng = np.random.default_rng(123)  # wie der Startwert gewählt wird, ist
↪ egal
```

```
rng.integers(23, 42, (4, 10))  # ganze Zahlen von 23 bis 41
```

```
array([[23, 35, 34, 24, 40, 27, 27, 26, 29, 26],
       [29, 38, 31, 40, 31, 28, 37, 38, 39, 39],
       [23, 32, 28, 27, 27, 38, 38, 27, 30, 37],
       [25, 34, 31, 40, 37, 27, 38, 38, 27, 32]])
```

```
rng.uniform(0, 2, (4, 4))  # gleichverteilt auf [0, 2]
```

```
array([[0.46311125, 0.33180799, 0.99557794, 1.16544928],
       [0.36867597, 0.02978983, 0.94226646, 1.45648666],
       [1.83720098, 1.25106801, 1.83424515, 1.7293805 ],
       [0.43628575, 1.73225486, 1.46150387, 0.55573058]])
```

4 Elementare Begriffe

Wiederholung und Präzisierung einiger Begriffe, die in allen Bereichen der Mathematik anzutreffen sind.

4.1 Aussagen

Die Mathematik tritt uns als Ansammlung von Begriffsdefinitionen und Aussagen zu den Beziehungen zwischen diesen Begriffen entgegen. Aussagen werden auch gern als *Satz* oder *Theorem* bezeichnet, wobei je nach Wichtigkeit der Aussage und ihrer Rolle in der Behandlung eines Themas unterschieden wird in Lemma (Hilfssatz), Proposition (Satz), Theorem (Hauptsatz), Korollar (Folgerung).

In dieser Veranstaltung werden wir die traditionelle Trennung der Inhalte in Definition, Satz, Beweis nicht vornehmen, sondern im Sinne der besseren Zugänglichkeit eine weniger strenge Form wählen. Dennoch sollten wir uns mit dem mathematischen Begriff der Aussage näher befassen, da Aussagen stets und ständig eine Rolle spielen werden, bemerkt und unbemerkt.

Was ist eine Aussage?

Merke

Eine **Aussage** ist ein Satz, der entweder wahr oder falsch ist. Ein solcher Satz kann sowohl in Alltagssprache als auch mit Hilfe mathematischer Begriffe und Symbole formuliert sein.

Diese Definition des Begriffs “Aussage” klingt banal, ist sie aber keinesweges. Offensichtlich ist

Dresden ist die Hauptstadt von Sachsen.

eine Aussage, sogar eine wahre. Auch

Leipzig ist die Hauptstadt von Sachsen.

ist eine Aussage, jedoch eine falsch oder unwahre.

Wie sieht es aber mit

In 50 Jahren wird Leipzig die Hauptstadt Sachsens sein.

aus? Hier können wir kein klares “wahr” oder “falsch” zuordnen. Wir könnten uns auf den Standpunkt zurückziehen, das die Aussage entweder richtig oder falsch ist und wir nur nicht wissen, welche der beiden Möglichkeiten zutrifft. Tatsächlich kann heute aber niemand wissen, welche Rolle Leipzig in 50 Jahren spielen wird, egal wie viele Informationen zur Verfügung stehen. Man könnte als möglichen Ausweg die Formulierung abschwächen:

In 50 Jahren wird Leipzig wahrscheinlich die Hauptstadt Sachsens sein.

Aber auch hier ist keine klare Zuordnung von “richtig” oder “falsch” möglich. Es handelt sich mehr um eine Meinung als um eine Aussage. Und Meinungen sind bekanntlich nicht

klar mit “richtig” oder “falsch” bewertbar.

Anderes Beispiel:

Am 1. Januar 1999 fuhren 23 rote Autos am Hauptgebäude der Hochschule vorbei.

Sehr wahrscheinlich hat an diesem Tag niemand Autos vor der Hochschule gezählt. Es ist also nicht möglich, diesem Satz ein klares “richtig” oder “falsch” zuzuordnen. Die dafür benötigten Informationen liegen objektiv nicht vor.

Beispiel (Bäcker-Paradoxon)

Wir definieren hiermit den Begriff des Bäckers als eine Person, die genau für all die Menschen Brötchen bäckt, die nicht selbst für sich Brötchen backen. Ist dann folgender Satz eine Aussage?

Der Bäcker bäckt auch Brötchen für sich selbst.

Wäre dies eine Aussage, so müsste sie entweder wahr oder falsch sein. Beide Varianten führen aber zum Widerspruch.

- Wenn der Satz wahr wäre, der Bäcker also auch für sich bäckt, dann folgt aus der Bäcker-Definition, dass er nicht für sich bäckt.
- Wenn der Satz falsch wäre, der Bäcker also nicht auch für sich bäckt, dann folgt aus der Bäcker-Definition, dass er für sich bäckt.

Ursprünglich wurde dieses Beispiel 1918 als Barbier-Paradoxon veröffentlicht.

Nebenbemerkung

Bei der alltäglichen Arbeit mit den mathematischen Werkzeugen werden wir praktisch nie auf solch problematische Situationen treffen wie in den obigen Beispielen zu Nicht-Aussagen. Man könnte daher zu der Meinung gelangen, dass derartige Probleme nur bei der Verwendung von Alltagssprache auftreten und somit in einer potentiell unpräzisen Formulierung begründet sind. Wenn man nur mathematische Begriffe und Symbole verwendet, geht bestimmt alles seinen geordneten Gang...

Leider falsch! Selbst in klar abgegrenzten, wohl durchdachten und fehlerfrei konstruierten mathematischen Strukturen, kann man stets “Aussagen” formulieren, von denen man mathematisch exakt beweisen kann, dass man nicht entscheiden kann, ob sie wahr oder falsch sind. In unserem Sinne handelt es sich dann also nicht um Aussagen. Diese für die bis dahin sehr heile Mathematikerwelt erschütternde Erkenntnis wurde 1931 von Kurt Gödel veröffentlicht und ist als Gödelscher

Unvollständigkeitssatz bekannt.

Tatsächlich sind unter den unentscheidbaren “Aussagen” auch leicht verständliche zu finden. Zum Beispiel ist folgende Frage zur Mächtigkeit von Mengen unentscheidbar (Kontinuumshypothese):

Es gibt eine Menge, die mächtiger ist als die natürlichen Zahlen, aber weniger mächtig als die reellen Zahlen.

Zum Begriff der Mächtigkeit vergleiche auch Aufgabe 1.2 in Übung 1.

Den Ansatz, dass Aussagen entweder wahr oder falsch sind, bezeichnet man als zweiwertige Logik. Weniger verbreitete alternative Ansätze sind die dreiwertige Logik, die Fuzzy-Logik und die intuitionistische Logik.

Äquivalenz von Aussagen

Wir werden bei der Beschäftigung mit Mathematik häufig Aussagen umformen bzw. umformulieren. Ziel ist, diese zu vereinfachen ohne dabei deren Sinngehalt zu verändern. Sind zwei Aussagen entweder beide gleichzeitig wahr oder beide gleichzeitig falsch, so heißen die Aussagen **äquivalent**. Man sagt auch: die eine Aussage gilt **genau dann, wenn** die andere Aussage gilt. Als Symbol verwendet man den **Äquivalenzpfeil** $\Leftrightarrow$ (ein zweiseitiger Pfeil mit doppelter Linie):

$$\text{Aussage 1} \quad \Leftrightarrow \quad \text{Aussage 2.} \quad (4.1)$$

Gilt nur die eine Richtung, d.h. folgt aus der Gültigkeit von Aussage 1, dass auch Aussage 2 wahr ist, so schreibt man

$$\text{Aussage 1} \quad \Rightarrow \quad \text{Aussage 2.} \quad (4.2)$$

Sprechweise: “Aussage 1 **impliziert** Aussage 2” oder “aus Aussage 1 **folgt** Aussage 2”. Der Pfeil $\Rightarrow$ heißt **Implikationspfeil** oder **Folgerungspfeil**. Alternative Schreibweise:

$$\text{Aussage 2} \quad \Leftarrow \quad \text{Aussage 1.} \quad (4.3)$$

Beispiel (Umformen von Gleichungen)

Äquivalenzpfeile können und sollten genutzt werden um die einzelnen Schritte beim Umformen von mathematischen Ausdrücken (z.B. Gleichungen oder Ungleichungen) miteinander in Beziehung zu setzen. Soll die Gleichung

$$4x - 5 = 19 \quad (4.4)$$

gelöst werden, so kann der Rechenweg als Kette äquivalenter Aussagen aufge-

geschrieben werden:

$$4x - 5 = 19 \Leftrightarrow 4x = 24 \Leftrightarrow x = 6. \quad (4.5)$$

Zeilenumbrüche sind natürlich möglich und meist auch zweckmäßig im Sinne der Übersichtlichkeit:

$$\begin{aligned} 4x - 5 &= 19 \\ \Leftrightarrow 4x &= 24 \\ \Leftrightarrow x &= 6. \end{aligned} \quad (4.6)$$

Beispiel (nicht äquivalente Umformungen)

Bei Lösen von Gleichungen und Ungleichungen muss beachtet werden, dass nicht alle Arten von Umformungsschritten äquivalente Aussagen liefern. Gesucht seien beispielsweise alle $x \in \mathbb{R}$, die $x + 2 = \sqrt{x + 4}$ erfüllen. Wir schreiben zunächst nur Implikationspfeile in “Rechenrichtung” und prüfen anschließend, ob jeweils auch die Gegenrichtung gilt:

$$\begin{aligned} x + 2 &= \sqrt{x + 4} \\ \Rightarrow (x + 2)^2 &= x + 4 && \text{(Quadrieren)} \\ \Rightarrow x^2 + 4x + 4 &= x + 4 && \text{(Ausmultiplizieren)} \\ \Rightarrow x^2 + 3x &= 0 && \text{(Zusammenfassen)} \\ \Rightarrow x(x + 3) &= 0 && (x \text{ ausklammern}) \\ \Rightarrow x = 0 \text{ oder } x = -3 && \text{(Faktoren Null setzen).} \end{aligned} \quad (4.7)$$

Setzen wir nun $x = 0$ und $x = -3$ in die ursprüngliche Gleichung ein, so stellen wir fest, dass nur $x = 0$ eine Lösung ist. Für $x = -3$ erhalten wir hingegen die falsche Aussage $-1 = 1$. Der Rechenweg liefert zwar

$$x + 2 = \sqrt{x + 4} \quad \Rightarrow \quad x = 0 \text{ oder } x = -3; \quad (4.8)$$

die Probe hingegen ergibt

$$x + 2 = \sqrt{x + 4} \quad \nRightarrow \quad x = 0 \text{ oder } x = -3, \quad (4.9)$$

wobei $\nRightarrow$ für “die Implikation gilt nicht” steht.

Mindestens einer der durchgeführten Rechenschritte ist demnach nicht umkehrbar. Geht man die Schritte rückwärts durch, so stellt man fest, dass der erste Rechenschritt (Quadrieren) nicht umkehrbar ist, denn das Wurzelziehen als Umkehrung des Quadrierens liefert die Implikation

$$(x + 2)^2 = x + 4 \quad \Rightarrow \quad |x + 2| = \sqrt{x + 4}. \quad (4.10)$$

Die Betragsstriche tauchen jedoch in der ursprünglichen Gleichung nicht auf. Reduziert auf das Wesentliche gilt für eine beliebige reelle Zahl $a \in \mathbb{R}$

$$\sqrt{a^2} = |a|. \quad (4.11)$$

Klar wird dies am Beispiel $a = -2$. Dann ist

$$\sqrt{a^2} = \sqrt{(-2)^2} = \sqrt{4} = 2 = |-2| = |a| \neq a. \quad (4.12)$$

Merke: Im Zweifelsfall sollte beim Umformen von Gleichungen und Ungleichungen die Gültigkeit der Implikation in beide Richtungen separat geprüft werden! Es ist streng zu unterscheiden zwischen äquivalenten Umformungen und einfachen Folgerungen, die eine allgemeinere, nicht äquivalente Aussage liefern.

Beschreiben von Mengen durch Aussagen

Ein ganz grundlegendes sprachliches Mittel in der Mathematik sind Mengen. Wir werden häufig mit Mengen von Objekten (Zahlen, Vektoren,...) mit gewissen Eigenschaften arbeiten. Welche Eigenschaften gegeben sein müssen, damit ein Objekt zur betrachteten Menge gehört, kann man mittels Aussagen über das Objekt formulieren. Die betrachtete Menge wird dann festgelegt als die Menge aller Objekte, für die eine gegebene Aussage über das einzelne Objekt wahr ist. Schreibweise:

$$M := \{x \in \text{Grundmenge} : \text{Aussage über } x\}. \quad (4.13)$$

Dabei sind

- M das Symbol für die zu beschreibende Menge,
- $:=$ der Zuweisungsoperator ("gib dem rechten Objekt das linke Symbol als Name"),
- "Grundmenge" die Menge aller zunächst betrachteten Objekte.

Beispiel (Menge der geraden Zahlen)

Die Menge der geraden Zahlen kann als

$$M := \{x \in \mathbb{N} : x \text{ ist durch } 2 \text{ teilbar}\} \quad (4.14)$$

geschrieben werden. Statt "8 ist eine gerade Zahl" können wir dann $8 \in M$ schreiben, was viel kürzer ist.

In der Mathematik sind solche die sprachliche Formulierung abkürzenden Mittel wesentliches Werkzeug um auch komplexe Zusammenhänge auf wenig Raum und damit besser überschaubar darzustellen. So kann beispielsweise die Aussage

Die Summe zweier Quadrate gerader Zahlen ist wieder eine gerade

Zahl.

nun als

$$x, y \in M \Rightarrow x^2 + y^2 \in M \quad (4.15)$$

geschrieben werden.

Merke: Die vielen Symbole und Formeln in der Mathematik dienen nicht dem Quälen von Studierenden, sondern sind die **Arbeitgrundlage für präzises und überschaubares Formulieren komplexer Zusammenhänge**.

Logik vs. Mengenoperationen

Zwei (oder mehr) gegebene Aussagen können durch sogenannte logische Operation zu neuen Aussagen verknüpft werden. So ist die Formulierung

Es gilt Aussage 1 oder Aussage 2.

selbst wieder eine Aussage. Analog gibt es die Und-Verknüpfung, die Entweder-Oder-Verknüpfung und die Nicht-Operation. Letztere ordnet einer Aussage genau die gegen-teilige Aussage zu.

Logische Operationen mit Aussagen übertragen sich eins-zu-eins auf Mengen. Dort entstehen dann z.B. Vereinigungen (Oder-Verknüpfung) und Durchschnitte (Und-Verknüpfungen):

$$\begin{aligned} &\{x \in \mathbb{N} : x \text{ durch 2 teilbar } \mathbf{oder} \text{ durch 3 teilbar}\} \\ &= \{x \in \mathbb{N} : x \text{ durch 2 teilbar}\} \cup \{x \in \mathbb{N} : x \text{ durch 3 teilbar}\} \\ &\{x \in \mathbb{N} : x \text{ durch 2 teilbar } \mathbf{und} \text{ durch 3 teilbar}\} \\ &= \{x \in \mathbb{N} : x \text{ durch 2 teilbar}\} \cap \{x \in \mathbb{N} : x \text{ durch 3 teilbar}\} \end{aligned} \quad (4.16)$$

Nebenbemerkung

Die Und-Verknüpfung mehrerer Aussagen schreibt man gern auch als durch Komma getrennte Auflistung der Einzelaussagen. Beispielsweise sind die Schreibweisen

$$\{x \in \mathbb{N} : x \text{ durch 2 teilbar und durch 3 teilbar}\} \quad (4.17)$$

und

$$\{x \in \mathbb{N} : x \text{ durch 2 teilbar}, x \text{ durch 3 teilbar}\} \quad (4.18)$$

völlig gleichbedeutend. Bei Oder-Verknüpfungen wird das “oder” hingegeben stets explizit angegeben.

Nebenbemerkung

Gelegentlich sieht man statt “und” und “oder” zur Verknüpfung von Aussagen die Symbole $\wedge$ und $\vee$. Für “nicht” wird gelegentlich $\neg$ verwendet. Diese Symbolik erschwert die Lesbarkeit und wird deshalb hier vermieden. Bei manchen Autoren im Bereich Ingenieurmathematik hat man den Eindruck, dass das Verwenden dieser Symbole die “mathematische Bildung” des Autors unterstreichen soll. In Mathematikbüchern für Mathematiker spielen diese Symbole bis auf wenige Spezialgebiete hingegen keine Rolle.

Beispiel

Das Zusammenspiel von Logik und Mengenlehre lässt sich gut an der Beschreibung geometrischer Objekte nachvollziehen. Betrachten dazu in der Ebenen die drei Teilmengen

- Kreuz:

$$\{(x, y) \in \mathbb{R}^2 : x \in [1, 3] \text{ oder } y \in [0, 1]\}, \quad (4.19)$$

- Rechteck:

$$\{(x, y) \in \mathbb{R}^2 : x \in [1, 3], y \in [0, 1]\}, \quad (4.20)$$

- Kreisring:

$$\{(x, y) \in \mathbb{R}^2 : 1 \leq x^2 + y^2 \leq 4\} \quad (4.21)$$

(Video 410, Notizen dazu).

4.2 Abbildungen

Abbildungen sind neben Mengen das wichtigste sprachliche Mittel in der Mathematik. Oft interessieren weniger die Objekte (z.B. Zahlen) selbst, sondern deren Beziehungen zueinander. Eine Analogie im täglichen Leben ist die Bedeutung von Daten in "sozialen" Netzwerken und Messenger-Diensten. Es ist kaum von Interesse was jemand schreibt oder sagt, sondern vor allem wer mit wem kommuniziert. Entsprechend stehen nicht die (meist zumindest oberflächlich verschlüsselten) Inhalte im Fokus von Überwachungsorganisationen und Datenhändlern, sondern die Metadaten, also wer wann mit wem kommuniziert hat.

Begriffe und Beispiele

Merke

Eine **Abbildung** ist eine Zuordnung zwischen den Elementen zweier Mengen A und B , bei denen jedem Element der Menge A genau ein Element der Menge B zugeordnet wird. Die Menge A wird als **Definitionsbereich** der Abbildung bezeichnet. Die Menge aller Elemente von B , die einem Element von A zugeordnet sind, heißt **Wertebereich** der Abbildung. Die Menge B selbst wird als Bildbereich oder als Grundmenge des Wertebereichs bezeichnet.

Zur einfacheren Verwendung bezeichnet man Abbildungen meist mit einem Symbol, z.B. f , und schreibt $f : A \rightarrow B$, wenn man Definitionsbereich und Bildbereich gleich mit angeben möchte. Beachte, dass B üblicherweise nicht der Wertebereich selbst, sondern eine Obermenge des Wertebereichs ist. Das genaue Bestimmen des Wertebereichs erfordert meist etwas Mühe, die man sich auf diese Weise (zunächst) spart.

Der Definitionsbereich einer Abbildung f wird auch mit $D(f)$ bezeichnet, der Wertebereich mit $W(f)$. Diese Symbole sind allerdings in der Literatur nicht einheitlich, sodass stets genau zu prüfen ist, was mit welchem Symbol gemeint ist.

Die Festlegung der Zuordnungsvorschrift kann entweder in Alltagssprache oder in Form von mathematischen Ausdrücken erfolgen. Ist $f : \mathbb{R} \rightarrow \mathbb{R}$ beispielsweise eine Abbildung, die jeder reellen Zahl x wieder eine reelle Zahl zuordnet, so schreibt man $f(x) := x^2$ oder $f : x \mapsto x^2$ (beachte die spezielle Form des Pfeilanzugs!). Umgekehrt verwendet man bei gegebener Abbildung $f : X \rightarrow Y$ das Symbol $f(x)$ für das Element in Y , auf das ein gegebenes Element $x \in X$ durch f abgebildet wird.

Abbildungen treten in ganz verschiedenen Varianten auf. Weiter unten betrachten wir einige wichtige Spezialfälle; hier zunächst eine buntere Beispielsammlung.

Nebenbemerkung

Die Angabe des konkreten Definitionsbereichs ist nicht immer ganz einfach. Deshalb wird dieser manchmal zunächst weggelassen. Der Definitionsbereich soll dann aus all den Elementen einer angegebenen Grundmenge bestehen, für die die Abbildung sinnvolle Werte liefert.

Beispiel (Flächeninhalt)

Sei F die Menge aller Teilmengen von $\mathbb{R}^2$, für die der Begriff “Flächeninhalt” sinnvoll erklärt werden kann (F für “Flächen”). Dann kann man die Abbildung $A : F \rightarrow \mathbb{R}$ betrachten, die jeder Fläche $f \in F$ deren Flächeninhalt (reelle Zahl!) zuordnet. Die Abbildung A lässt sich offensichtlich nicht als Formel aufschreiben. Überhaupt müsste noch geklärt werden, wie der Definitionsbereich F genau aussieht (werden wir später im Semester klären).

Wichtig: Gewöhnen Sie sich von Anfang an an Flexibilität bei der Wahl von Symbolen. Eine Abbildung muss keineswegs immer f heißen!

Beispiel (Klausurnoten)

Sei N die Menge aller Noten, die bei der Bewertung der Klausuren eines bestimmten Studierendenjahrgangs auftraten. Die Abbildung S ordne jeder Note die Menge aller Studierenden zu, die diese Note erhalten haben. Beachte:

- Der Definitionsbereich von S ist auch dann die komplette Menge N , wenn nicht jede Note vergeben worden ist. Beispielsweise könnte die Note 5.0 nicht auftauchen (alle haben bestanden). Dann wäre $S(5.0) = \emptyset$ (leere Menge).
- Die genaue Formulierung ist hier wichtig. Falsch wäre

Jeder Note werden die Studierenden mit dieser Note zugeordnet.

Denn diese Formulierung widerspricht der Forderung, dass den Elementen des Definitionsbereichs jeweils genau ein Element des Wertebereichs zugeordnet wird. In der korrekten Formulierung oben wird jeder Note eine (!) Menge von Studierenden zugeordnet. Ob solch eine Menge dann kein oder ein oder mehrere Elemente enthält, ist unerheblich für die formale Korrektheit der Definition der Abbildung.

Beispiel (gerade Zahlen)

Definieren die Abbildung $f : \mathbb{N} \rightarrow \mathbb{N}$ durch $n \mapsto 2n$. Jeder natürlichen Zahl wird

also ihr Doppeltes zugeordnet.

Zur Beschreibung grundlegender Eigenschaften von Abbildungen führen wir folgende Begriffe ein.

Merke

Sei $f : A \rightarrow B$ eine Abbildung.

- Die Abbildung f heißt **injektiv**, wenn jedes Element von B höchstens einem Element von A zugeordnet wird. Kein Element von B wird also zwei verschiedenen Elementen von A zugeordnet. In Formeln:

$$\text{für alle } a_1, a_2 \in A \text{ gilt: } f(a_1) = f(a_2) \Rightarrow a_1 = a_2. \quad (4.22)$$

- Die Abbildung f heißt **surjektiv**, wenn jedes Element von B mindestens einem Element von A zugeordnet ist. In Formeln:

$$W(f) = B. \quad (4.23)$$

- Die Abbildung f heißt **bijektiv**, wenn sie gleichzeitig injektiv und surjektiv ist.

Beachte, dass Surjektivität von der Wahl der Grundmenge B abhängt. Bei gleichem Wertebereich kann B unterschiedlich aussehen. So ist beispielsweise die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ nicht surjektiv, die Abbildung $g : \mathbb{R} \rightarrow [0, \infty), x \mapsto x^2$ hingegen schon.

Zu den obigen drei Beispielen:

- Die Zuordnung des Flächeninhalts ist nicht injektiv, da verschiedene Flächen den gleichen Inhalt haben können. Sie ist aber surjektiv, da wir zu jeder reellen Zahl eine Fläche mit diesem Inhalt finden können. Die Abbildung ist nicht bijektiv.
- Die Zuordnung der Studierenden zu den Klausurnoten ist injektiv, denn an die gleiche Gruppe von Studierenden können nicht zwei verschiedene Noten vergeben worden sein. Zur Klärung der Surjektivität müssen wir uns zunächst auf eine Grundmenge für den Wertebereich festlegen. Da prinzipiell jede Kombination von Studierenden bei diesem Abbildungstyp auftreten kann, ist die **Potenzmenge** der Studierenden (also die Menge aller möglichen Teilmengen von Studierenden) eine geeignete Wahl. Da in der Abbildung viel weniger Noten zur Verfügung stehen als es verschiedene Teilmengen von Studierenden gibt, ist die Abbildung offensichtlich nicht surjektiv. Somit ist sie auch nicht bijektiv.
- Die Zuordnung der geraden Zahlen zu den natürlichen Zahlen ist sowohl injektiv (keine gerade Zahl wird zwei verschiedenen Zahlen zugeordnet) als auch surjektiv.

(jede gerade Zahl tritt in der Zuordnung auf). Also ist die Abbildung bijektiv.

Die Frage nach Injektivität und Surjektivität kann auch anhand der Lösbarkeit von Gleichungen beschrieben werden: Ist $f : X \rightarrow Y$ eine beliebige Abbildung und ist $y \in Y$, so können wir nach Lösbarkeit und Lösungen $x \in X$ der Gleichung

$$f(x) = y \quad (4.24)$$

fragen. Offensichtlich besitzt die Gleichung genau dann für jedes $y \in Y$ eine Lösung, wenn f surjektiv ist. Außerdem ist für y aus dem Wertebereich von f die Lösung genau dann eindeutig, wenn f injektiv ist.

Folgender Begriff taucht gelegentlich auf, liefert aber inhaltlich keine neuen Aspekte. Er ist ein Beispiel für Begriffe, die rein zur sprachlichen Vereinfachung bzw. Fixierung dienen.

Merke

Sei $f : X \rightarrow Y$ eine Abbildung. Die Menge

$$\{(x, f(x)) \in X \times Y : x \in X\} \quad (4.25)$$

heißt **Graph** der Abbildung f .

In die gleiche Rubrik fällt der folgende Begriff:

Merke

Sind $f : A \rightarrow B$ und $g : B \rightarrow C$ zwei Abbildungen, so heißt die durch $a \mapsto g(f(a))$ gegebene Abbildung **Komposition** oder **Hintereinanderausführung** von f und g . Als Symbol für diese Abbildung verwenden wir $g \circ f$.

Spezialfall: Punkte und Vektoren

Sei $n \in \mathbb{N}$ eine feste gegebene Zahl (oft $n = 2$ oder $n = 3$). Ein **Punkt** ist eine n -Tupel aus n reellen Zahlen. 2-Tupel sind auch als Paare bekannt, 3-Tupel als Tripel, 4-Tupel als Quadrupel usw. Alternativ kann man einen Punkt als Abbildung der Menge $\{1, 2, \dots, n\}$ in die reellen Zahlen ansehen. Die n Bildelemente einer solchen Abbildung nennt man **Koordinaten** des Punktes. Statt der Kammerschreibweise für das Auswerten einer Abbildung nutzt man hier einfache Auflistungen der Koordinaten. So entspricht die Auflistung $(7, 4, 5)$ beispielsweise der Abbildung

$$1 \mapsto 7, \quad 2 \mapsto 4, \quad 3 \mapsto 5. \quad (4.26)$$

Die Menge aller Punkte zu einem festen n wird mit dem Symbol $\mathbb{R}^n$ bezeichnet. So umfasst $\mathbb{R}^2$ beispielsweise alle Punkte mit zwei Koordinaten. $\mathbb{R}^2$ wird auch als “die Ebene” und $\mathbb{R}^3$ als “der Raum” bezeichnet.

Das Wort **Vektor** ist aus technischer Sicht ein Synonym für Punkt, meint also ein Element aus $\mathbb{R}^n$. Man unterscheidet die Elemente von $\mathbb{R}^n$ jedoch gelegentlich nach deren Verwendung. Soll durch ein Element aus $\mathbb{R}^n$ ein geometrischer Ort beschrieben werden (“Punkt im Koordinatensystem”), so verwendet man den Begriff Punkt. Soll jedoch eine Richtung beschrieben werden, so spricht man von einem Vektor. Während Punkte, wie oben eingeführt, als richtungslose Auflistung $(a_1, \dots, a_n)$ ihrer Koordinaten geschrieben werden, werden Vektoren meist als vertikale Listen

$$\begin{bmatrix} a_1 \\ \vdots \\ a_n \end{bmatrix} \quad \text{oder} \quad \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} \quad (4.27)$$

ihrer Koordinaten dargestellt. Bei Vektoren spricht man statt von Koordinaten auch von den **Komponenten** des Vektors.

Die Arbeit mit Punkten und Vektoren wird später detaillierter betrachtet.

Spezialfall: Zahlenfolgen

Eine Abbildung, die jeder natürlichen Zahl eine reelle Zahl zuordnet, heißt **Zahlenfolge**. Man weicht hier aus Gewohnheit und aus Gründen der Platzeinsparung von der üblichen Schreibweise für das Auswerten von Abbildungen ab: Ist $a : \mathbb{N} \rightarrow \mathbb{R}$ eine Zahlenfolge, so setzt man

$$a_n := a(n) \quad \text{für alle } n \in \mathbb{N}. \quad (4.28)$$

Man schreibt auch gern $a = (a_1, a_2, \dots)$ im Sinne eines Punktes (oder Vektors) mit unendlich vielen Koordinaten oder $a = (a_n)_{n \in \mathbb{N}}$.

Gelegentlich treten in der Mathematik **doppelt unendliche Zahlenfolgen** auf, insbesondere in der Signalverarbeitung. Dies sind Abbildungen $a : \mathbb{Z} \rightarrow \mathbb{R}$, die auch negative Indizes nutzen, also nach beiden Richtungen laufende Folgen darstellen. In dieser Veranstaltung werden wir solchen Folgen jedoch nicht begegnen.

Zahlenfolgen werden wir in Kürze noch näher betrachten.

Spezialfall: Funktionen

Bezeichnen Abbildungen zwischen Teilmengen der reellen Zahlen oder Mengen von aus endlich vielen reellen Zahlen zusammengesetzten Objekten (Punkte, Vektoren, Matrizen,...) als **Funktion**. Dabei soll der Definitionsbereich nicht nur aus einzelnen Zahlen, Punkten usw. bestehen, sondern den kontinuierlichen Charakter der reellen Zahlen tragen. Beispielsweise sind Zahlenfolgen in diesem Sinne also keine Funktionen.

Gelegentlich werden die Begriffe Abbildung und Funktion synonym verwendet; verstehen in dieser Veranstaltung unter Funktionen jedoch nur zahlenbezogene und keine allgemeinen Abbildungen.

Funktionen spielen eine sehr wesentliche Rolle in der angewandten Mathematik, sodass wir ihnen ein eigenes Unterkapitel widmen werden.

Spezialfall: Allgemeine Folgen

Neben Zahlenfolgen ist gelegentlich auch die Betrachtung anderer Folgen zweckmäßig. Zum Beispiel werden wir im Laufe der Zeit mit Folgen von Punkten in der Ebene arbeiten (müssen). Dies sind dann Abbildungen von $\mathbb{N}$ nach $\mathbb{R}^2$. Auch Folgen von Funktionen spielen in der Mathematik und ihren Anwendungen eine wichtige Rolle; jedoch weniger in der Grundlagenausbildung. Wenn wir uns im Laufe des Semesters mit dem Begriff des Flächeninhalts befassen, werden wir auf Folgen von Mengen stoßen, wobei jede dieser Mengen wiederum eine Ansammlung von Mengen sein wird. Folgen, also Abbildung der natürlichen Zahlen in andere Mengen, sind überall...

4.3 Zahlenfolgen

Der Begriff der Zahlenfolge wurde bereits eingeführt (vgl. Section 4.2).

Nebenbemerkung (Folge vs. Menge)

Beachte, dass bei einer Zahlenfolge die Reihenfolge der Elemente ganz wesentlich ist. Hingegen ist bei der aus den Elementen der Zahlenfolge gebildeten Menge $\{a_1, a_2, \dots\}$ die Reihenfolge egal. Insbesondere kann die Menge viel weniger Elemente enthalten als die Zahlenfolge. Betrachtet man beispielsweise die konstante Zahlenfolge $a = (1, 1, \dots)$, so besteht die zugehörige Menge der Elemente aus nur einem Element: $\{a_1, a_2, \dots\} = \{1\}$, ganz im Gegensatz zur unendlich viele Elemente enthaltenden Zahlenfolge.

Visualisierung

Die Visualisierung von Zahlenfolgen kann mit dem Python-Paket Matplotlib (und vielen anderen) erfolgen. Zu beachten ist, dass die einzelnen Punkte nicht durch Linien verbunden werden, da Zahlenfolgen nur auf der Menge $\{1, 2, \dots\}$ definiert sind und zwischen diesen Werten nicht.

Beispiel

Die Zahlenfolge

$$a_n := \frac{1}{n}, \quad n \in \mathbb{N} \quad (4.29)$$

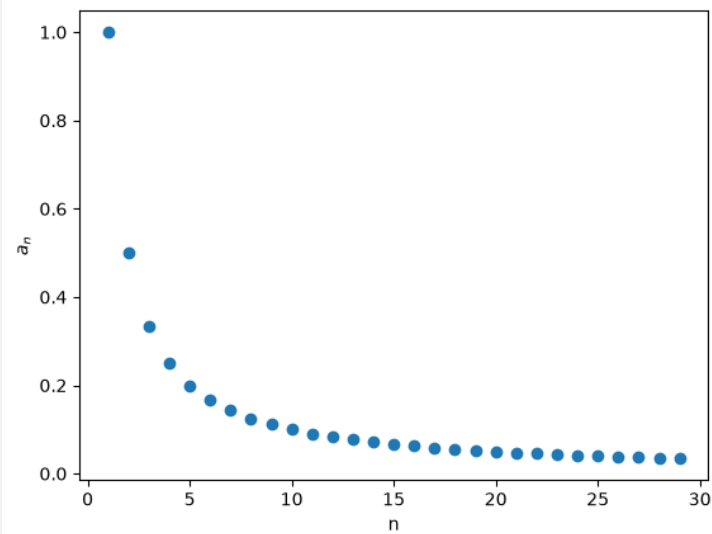
kann wie folgt visualisiert werden:

```
import numpy as np
import matplotlib.pyplot as plt

n = np.arange(1, 30) # Indizes der Zahlenfolge (natürlich nur
    ↳ endlich viele)
a = 1 / n # Werte der Zahlenfolge (a ist ein Array mit 29
    ↳ Einträgen!)

fig, ax = plt.subplots() # Zeichenfläche und Koordinatensystem
    ↳ erstellen
ax.plot(n, a, 'o') # 'o' für "keine Linien, aber dicke Punkte"
ax.set_xlabel('n') # Beschriftung horizontale Achse
ax.set_ylabel('$a_n$') # Beschriftung vertikale Achse
                        # (Mathe -Modus für Formeln mit $...$
                        ↳ aktivieren)
```

```
plt.show() # Zeichenfläche anzeigen
```



Bedeutung

Zahlenfolgen tauchen gelegentlich bei praktischen Aufgabenstellungen auf. Im Rahmen unserer Veranstaltung werden sie aber überwiegend als innermathematisches Werkzeug dienen. Insbesondere bilden sie die Basis für die Konstruktion komplexerer Werkzeuge wie dem Stetigkeitsbegriff, der Differentialrechnung und der Integralrechnung.

Beispiel (Zinseszins)

Als einfaches Beispiel für das Auftreten von Zahlenfolgen in der Praxis kann man üblicherweise auf Jahresbasis stattfindende Kapitalentwicklungen anschauen: Sei ein Startguthaben $g_0 > 0$ gegeben. Jedes Jahr wird das vorhandene Guthaben mit 1 Prozent Zinsen vergütet. Die Zinsen werden dem Guthaben zugeschlagen und im weiteren Verlauf mit verzinst. Nach einem Jahr steht somit ein größeres Guthaben g_1 zur Verfügung; nach zwei Jahren das Guthaben g_2 usw. Konkret erhalten wir die Zahlenfolge

$$(g_1, g_2, \dots) = (1.01 g_0, 1.01^2 g_0, \dots) \quad (4.30)$$

bzw.

$$g_n = 1.01^n g_0, \quad n \in \mathbb{N}. \quad (4.31)$$

In einer realen Situation wird die Folge nach ein paar Jahren abbrechen, also gar keine Zahlenfolge im Sinne der mathematischen Begriffs sein. Das Guthaben

wird nicht unendlich viele Jahre wachsen; selbst wenn, wird kein Mensch unendlich lange leben um dieses Guthaben “am Ende” in Empfang nehmen zu können. Das mathematische Werkzeug der Zahlenfolge erlaubt jedoch, den Zinseszinsprozess zu modellieren ohne sich über die konkrete Dauer des Prozesses Gedanken machen zu müssen. Wir vereinfachen (!) hier also den realen Vorgang zu einem mathematischen Modell, welches eine schwierige reale Frage (Wie lange soll die Geldanlage laufen?) umgeht, indem es die völlig realitätsferne Idee des Unendlichen einbezieht.

Monotonie, Periodizität, Beschränktheit

Führen einige Begriffe ein, die im Kontext von Zahlenfolgen gelegentlich Verwendung finden.

Merke

Eine Zahlenfolge $(a_1, a_2, \dots)$ heißt

- **monoton wachsend**, wenn $a_n \leq a_{n+1}$ für alle $n \in \mathbb{N}$ gilt,
- **streng monoton wachsend**, wenn $a_n < a_{n+1}$ für alle $n \in \mathbb{N}$ gilt,
- **monoton fallend**, wenn $a_n \geq a_{n+1}$ für alle $n \in \mathbb{N}$ gilt,
- **streng monoton fallend**, wenn $a_n > a_{n+1}$ für alle $n \in \mathbb{N}$ gilt.

Beispielsweise ist die Zahlenfolge in obigem Zinseszinsbeispiel streng monoton wachsend.

Merke

Eine Zahlenfolge $(a_1, a_2, \dots)$ heißt **periodisch**, wenn es ein $N \in \mathbb{N}$ gibt, sodass

$$a_n = a_{n+N} \quad \text{für alle } n \in \mathbb{N} \quad (4.32)$$

gilt. Jede Zahl N mit dieser Eigenschaft heißt **Periode** der Zahlenfolge. Die kleinste Zahl N mit dieser Eigenschaft heißt **Grundperiode** der Zahlenfolge.

Beispielsweise ist die durch $a_n = (-1)^n$, $n \in \mathbb{N}$, gegebene Folge periodisch mit Periode $N = 2$.

Merke

Eine Zahlenfolge $(a_1, a_2, \dots)$ heißt **beschränkt**, wenn es ein $c \in \mathbb{R}$ gibt, sodass

$$|a_n| \leq c \quad \text{für alle } n \in \mathbb{N} \quad (4.33)$$

gilt. Ansonsten heißt die Folge **unbeschränkt**.

Beispielsweise ist die durch $a_n := n$, $n \in \mathbb{N}$, gegebene Folge unbeschränkt. Hingegen ist die durch $a_n := \frac{1}{n}$, $n \in \mathbb{N}$, gegebene Folge beschränkt (mit $c = 1$ oder größer).

Konvergenz

Fast immer interessieren bei einer Zahlenfolge gar nicht die vorderen Elemente, sondern nur das Verhalten der Folge “im Unendlichen”. Werden die Folgeelemente beliebig groß (oder klein) oder nähern sie sich einem konkreten Wert an usw.? Zahlenfolgen können regelrecht als Werkzeug für das Arbeiten mit “dem Unendlichen” angesehen werden.

Merke

Eine Zahlenfolge $(a_1, a_2, \dots)$ heißt **konvergent**, wenn es eine Zahl $g \in \mathbb{R}$ gibt, sodass für jedes $\varepsilon > 0$ ein $m \in \mathbb{N}$ existiert, sodass

$$|a_n - g| < \varepsilon \quad \text{für alle } n \geq m \quad (4.34)$$

gilt. Die Zahl g heißt dann **Grenzwert** der Folge. Man schreibt auch

$$\lim_{n \rightarrow \infty} a_n = g \quad (4.35)$$

oder

$$a_n \rightarrow g \quad \text{für } n \rightarrow \infty. \quad (4.36)$$

Eine Folge ist also konvergent, wenn ihre Elemente a_n für hinreichend großes n in einem beliebig klein gewählten Intervall $(g - \varepsilon, g + \varepsilon)$ bleiben.

Beispiel

Die durch $a_n := \frac{1}{n}$ gegebene Folge konvergiert gegen 0 (Video 420, Notizen dazu).

Merke

Eine Zahlenfolge $(a_1, a_2, \dots)$ heißt **divergent**, wenn sie nicht konvergent ist. Man unterscheidet bestimmt divergente Folgen und unbestimmt divergente Folgen.

Eine divergente Folge heißt **bestimmt divergent**, wenn für jede Zahl $c \in \mathbb{R}$ ein $m \in \mathbb{N}$ existiert, sodass

$$a_n > c \quad \text{für alle } n \geq m \quad (4.37)$$

gilt, oder wenn für jede Zahl $c \in \mathbb{R}$ ein $m \in \mathbb{N}$ existiert, sodass

$$a_n < c \quad \text{für alle } n \geq m \quad (4.38)$$

gilt. Man schreibt dann auch

$$\lim_{n \rightarrow \infty} a_n = \infty \quad \text{bzw.} \quad \lim_{n \rightarrow \infty} a_n = -\infty. \quad (4.39)$$

Andernfalls heißt die Folge **unbestimmt divergent**.

Bestimmte Divergenz liegt also vor, wenn eine Folge über jede Schranke wächst oder unter jede Schranke fällt. Ersteres ist bei obigem Zinseszinsbeispiel der Fall.

Beispiele für unbestimmte Divergenz sind (nicht konstante) periodische Folgen.

Berechnen von Grenzwerten

Das Berechnen von Grenzwerten von Zahlenfolgen ist nur sehr selten nötig. Wir werden zwar viel mit Zahlenfolgen arbeiten, benötigen dabei aber nur das sprachliche Konstrukt und die damit möglichen Formulierungen zu Abläufen mit unendlich vielen Einzelschritten. Geben deshalb hier nur einige Hinweise für das konkrete Ausrechnen von Grenzwerten mit auf den Weg.

Manuelles Rechnen Für einfache Zahlenfolgen kann der Grenzwert (so er denn existiert) manuell berechnet werden. Dazu bedient man sich einer Tabelle von Grenzwerten für sehr elementare Folgen und einiger Rechenregeln zu Grenzwerten beim Bilden von komplexeren Zahlenfolgen aus diesen elementaren Folgen.

Merke (Grenzwerte elementarer Zahlenfolgen)

Hier einige nützliche Grenzwerte:

- konstante Folge:

$$\lim_{n \rightarrow \infty} c = c \quad \text{für beliebiges } c \in \mathbb{R}, \quad (4.40)$$

- Potenzen:

$$\lim_{n \rightarrow \infty} n^x = \begin{cases} 0, & \text{für } x < 0, \\ 1, & \text{für } x = 0, \\ \infty, & \text{für } x > 0, \end{cases} \quad (4.41)$$

- Euler'sche Zahl:

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = e \approx 2.71828. \quad (4.42)$$

Merke (Rechenregeln für Grenzwerte)

Mit Hilfe folgender Eigenschaften von Grenzwerten können aus Grenzwerten elementarer Folgen Grenzwerte komplexerer Folgen berechnet werden. Dabei sind $(a_1, a_2, \dots)$ und $(b_1, b_2, \dots)$ stets zwei konvergente (!) Folgen.

- Die elementweise Summe zweier Folgen konvergiert gegen die Summe der Grenzwerte:

$$\lim_{n \rightarrow \infty} (a_n + b_n) = \lim_{n \rightarrow \infty} a_n + \lim_{n \rightarrow \infty} b_n. \quad (4.43)$$

- Die Produkt zweier Folgen konvergiert gegen das Produkt der Grenzwerte:

$$\lim_{n \rightarrow \infty} (a_n b_n) = \left(\lim_{n \rightarrow \infty} a_n\right) \left(\lim_{n \rightarrow \infty} b_n\right). \quad (4.44)$$

- Das Vielfache einer Folge konvergiert gegen das Vielfache des Grenzwerts:

$$\lim_{n \rightarrow \infty} (c a_n) = c \lim_{n \rightarrow \infty} a_n \quad \text{für beliebiges } c \in \mathbb{R}. \quad (4.45)$$

- Sind alle Folgeelemente und der Grenzwert von Null verschieden, so konvergiert das Reziproke einer Folge gegen das Reziproke des Grenzwerts:

$$\lim_{n \rightarrow \infty} \frac{1}{a_n} = \frac{1}{\lim_{n \rightarrow \infty} a_n}. \quad (4.46)$$

Beispiel

Es gilt

$$\lim_{n \rightarrow \infty} \frac{3n^2}{n^2 + 2n} = 3 \quad (4.47)$$

(Video 425, Notizen dazu).

Beachte, dass die obigen Rechenregeln nur für konvergente Folgen gelten! Bei divergenten Folgen, insbesondere auch bei bestimmt divergenten Folgen, ist nicht immer klar, wie sich Summen und Produkte verhalten (vgl. Übung 3).

Numerisches Rechnen Das Verhalten einer Zahlenfolge im Unendlichen kann durch (numerisches) Berechnen später Elemente der Folge und ggf. deren grafischer Darstellung erhellt werden. Allerdings sind die daraus gewonnenen Informationen über Konvergenz/Divergenz/Grenzwert als reine Vermutung anzusehen. Prinzipiell ist denkbar, dass eine Folge beispielsweise in den ersten 10 Milliarden Elementen konstant ist und erst dann divergiert. Es kann also nie genügen, nur endlich viele Elemente zu betrachten.

Symbolisches Rechnen Das Python-Paket SymPy für symbolisches Rechnen bietet auch das Berechnen von Grenzwerten von Zahlenfolgen:

```
import sympy

n = sympy.symbols('n')

sympy.limit(3 * n ** 2 / (n ** 2 + 2), n, sympy.oo)
```

```
3
```

```
n = sympy.symbols('n')

sympy.limit((1 + 2 / n) ** n, n, sympy.oo)
```

```
exp(2)
```

Zahlenreihen

Eine spezielle Form von Zahlenfolgen sind sogenannte Zahlenreihen, oft auch nur **Reihen** genannt. Dies sind Zahlenfolgen, bei denen jedes Folgeelement eine Summe ist, wobei von Element zu Element ein Summand dazukommt. Ist $(a_n)_{n \in \mathbb{N}}$ eine Zahlenfolge, so kann man eine Zahlenreihe $(b_n)_{n \in \mathbb{N}}$ als

$$b_n := \sum_{k=1}^n a_k = a_1 + a_2 + \cdots + a_n \quad (4.48)$$

bilden.

Die Frage nach der Konvergenz einer solchen Reihe ist also die Frage nach der Endlichkeit der aus unendlich vielen Summanden bestehenden Summe. Man verwendet für den Grenzwert einer Reihe das Symbol

$$\sum_{n=1}^{\infty} a_n := \lim_{n \rightarrow \infty} b_n. \quad (4.49)$$

Die Zahlen b_n werden als **Partialsummen** bezeichnet.

Da wir Zahlenreihen nur am Rande bei der Wahrscheinlichkeitsrechnung antreffen werden, vertiefen wir das Thema hier nicht weiter.

4.4 Funktionen

Der Begriff der Funktion wurde bereits eingeführt (vgl. Section 4.2).

Sprachliche Aspekte

Funktionen sind ein fundamentales sprachliches Mittel in der Mathematik, welches vielfältige Ausdrucksmöglichkeiten bietet. Um in dieser Vielfalt nicht die Orientierung zu verlieren, ist präzises Formulieren wichtig!

Beispiel

Die Funktion

$$f : \mathbb{R}^4 \rightarrow \mathbb{R}, \quad f(x_1, y_1, x_2, y_2) := \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2} \quad (4.50)$$

ordnet den Koordinaten x_1, y_1, x_2, y_2 der Punkte (x_1, y_1) und (x_2, y_2) deren Abstand in der Ebene zu.

Inhaltlich identisch, aber formal verschieden, ist die Funktion

$$g : \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R}, \quad g((x_1, y_1), (x_2, y_2)) := \sqrt{(x_1 - x_2)^2 + (y_1 - y_2)^2}, \quad (4.51)$$

die den zwei Punkten direkt den Abstand zuordnet ohne Umweg über die Koordinaten.

Sind P_1 und P_2 zwei Punkte in der Ebene, so ist $g(P_1, P_2)$ ein korrekter mathematischer Ausdruck. Der Ausdruck $f(P_1, P_2)$ ist hingegen nicht definiert, da f nur für vier Argumente definiert wurde. Der Ausdruck $f(a)$ für $a \in \mathbb{R}^4$ ist wiederum korrekt; $g(a)$ hingegen nicht.

Funktionen (und allgemeinere Abbildungen) sind ein sprachliches Konstrukt, das auch von Computern verstanden wird. Dort sind Funktionen fast so elementar wichtig wie in der Mathematik. Das obige Beispiel kann am Computer wie folgt umgesetzt werden (beachte, dass man statt f und g aussagekräftigere Bezeichnungen wählen sollte):

```
# Definition einer Funktion (entspricht f im Beispiel)
def abstand_koord(x1, y1, x2, y2):

    return ((x1 - x2) ** 2 + (y1 - y2) ** 2) ** 0.5

# Definition einer weiteren Funktion (entspricht g im Beispiel)
def abstand_punkte(punkt1, punkt2):

    x1, y1 = punkt1
    x2, y2 = punkt2
```

```

    return ((x1 - x2) ** 2 + (y1 - y2) ** 2) ** 0.5

# Aufruf bzw. Auswertung der Funktionen
print(abstand_koord(1, 1, 3, 2))
print(abstand_punkte((1, 1), (3, 2)))

# alternativ
punkt1 = (1, 1)
punkt2 = (3, 2)
print(abstand_punkte(punkt1, punkt2))
print(abstand_koord(punkt1, punkt2)) # Zeile führt zu Fehlermeldung!

```

```

2.23606797749979
2.23606797749979
2.23606797749979

```

```

- - - - -
↪ - - - - -
↪ - - - - -
TypeError                                Traceback (most recent call
↪ last)
Cell In[1], line 21
    17 # alternativ
    18 punkt1 = (1, 1)
    19 punkt2 = (3, 2)
    20 print(abstand_punkte(punkt1, punkt2))
- - -> 21 print(abstand_koord(punkt1, punkt2)) # Zeile führt zu
↪ Fehlermeldung!

TypeError: abstand_koord() missing 2 required positional arguments:
↪ 'x2' and 'y2'

```

Beachte, dass `punkt1` und `punkt2` in der Definition der Funktion `abstand_punkte` **lokale Variablen** sind, im alternativen Aufruf weiter unten dann aber **globale Variablen**. Auch wenn diese Unterscheidung nicht zum Kernbereich dieser Veranstaltung gehört, sondern den Programmiergrundlagen zuzuordnen ist, sei zum besseren Verständnis die (kurze) Lektüre von Local versus Global Names empfohlen.

Visualisierung

Funktionen werden gern grafisch dargestellt, da viele Eigenschaften (die wir später näher untersuchen werden) leicht aus grafischen Darstellungen abgelesen werden können. Unter Nutzung des Computers sind auch komplexere, insbesondere dreidimensionale Vi-

sualierungen möglich. Visualisiert wird der Graph der Funktion, also eine Menge von Punkten in einem Koordinatensystem.

Für zweidimensionale Darstellungen eignet sich das Python-Paket Matplotlib sehr gut. Für dreidimensionale Darstellungen gibt es eine Erweiterung für Matplotlib (`mpl_toolkits.mplot3d`), jedoch bietet Plotly mehr Möglichkeiten.

Beispiel (einfacher 2D-Plot)

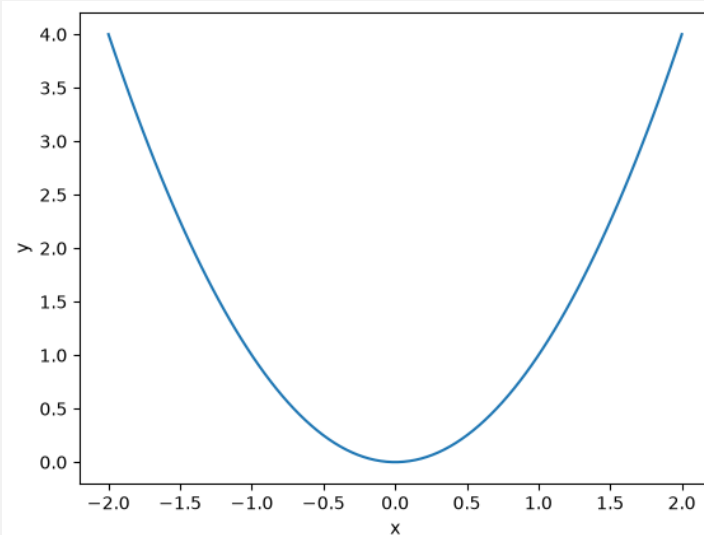
Funktionen, die einer reellen Zahl wieder eine reelle Zahl zuordnen, können in einem einfachen zweidimensionalen Plot dargestellt werden. Visualisieren die Funktion

$$f: \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) = x^2. \quad (4.52)$$

```
import numpy as np
import matplotlib.pyplot as plt

x = np.linspace(-2, 2, 100) # Intervall und Stützstellenanzahl
    ↪ für x
y = x ** 2 # Funktionswerte an den Stützstellen

fig, ax = plt.subplots() # Zeichenfläche erstellen
ax.plot(x, y) # Linie für Funktionsverlauf hinzufügen
ax.set_xlabel('x') # Beschriftung x-Achse
ax.set_ylabel('y') # Beschriftung y-Achse
plt.show() # Zeichenfläche anzeigen
```



Beachte, dass die Darstellung am Computer praktisch immer durch Auswerten der Funk-

tion an endlich vielen Stützstellen (d.h. Werten für x) erfolgt. Die entstehenden Punkte werden dann durch Geradenstücke verbunden. Ist die Stützstellenanzahl zu gering gewählt, kann die grafische Darstellung stark vom tatsächlichen Funktionsverlauf abweichen.

Beispiel (Kurve in der Ebene)

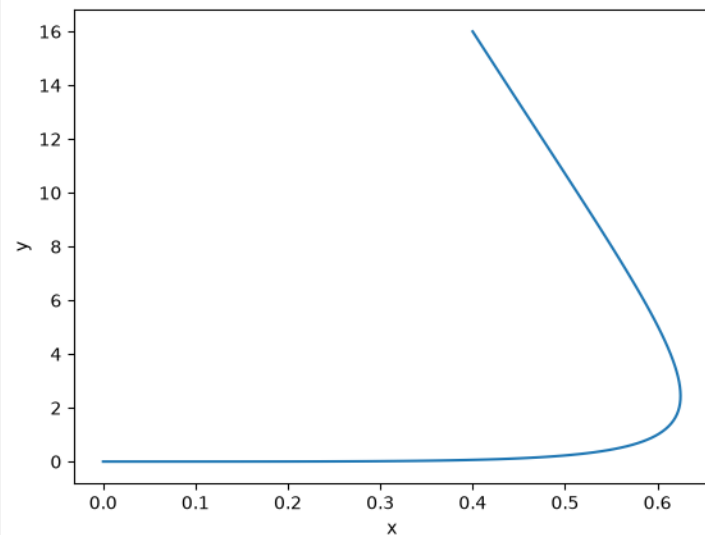
Die Position eines beweglichen Objekts in der Ebene kann durch eine Funktion der Form $t \mapsto (x(t), y(t))$ beschrieben werden. Dabei sind x und y Funktionen von $\mathbb{R}$ nach $\mathbb{R}$. Jedem Zeitpunkt t wird also eine Position $(x(t), y(t))$ zugeordnet. Funktionen dieser Form nennt man auch **Kurven in der Ebene**. Betrachten hier konkret

$$x(t) := t - 0.4t^2, \quad y(t) = t^4, \quad t \in [0, 2]. \quad (4.53)$$

Zwei Varianten für die Visualisierung von Kurven sind gebräuchlich: die Darstellung in der Ebene (ohne Zeitbezug) und die Darstellung im Raum (mit Zeitbezug).

```
t = np.linspace(0, 2, 100)
x = t - 0.4 * t ** 2
y = t ** 4

fig, ax = plt.subplots()
ax.plot(x, y)
ax.set_xlabel('x')
ax.set_ylabel('y')
plt.show()
```



```
import plotly.graph_objects as go
```

```

t = np.linspace(0, 2, 100)
x = t - 0.4 * t ** 2
y = t ** 4

fig = go.Figure()
fig.add_trace(
    go.Scatter3d(x=t, y=x, z=y, mode='lines', line_width=5)
)
fig.update_layout(
    scene_xaxis_title_text='t',
    scene_yaxis_title_text='x',
    scene_zaxis_title_text='y',
    margin=dict(l=0, r=0, b=0, t=0) # breiten Rand vermeiden
)
fig.show()

```

Beispiel (Kurven im Raum)

Die Position eines beweglichen Objekts im Raum kann durch eine Funktion der Form $t \mapsto (x(t), y(t), z(t))$ beschrieben werden. Dabei sind x , y und z Funktionen von $\mathbb{R}$ nach $\mathbb{R}$. Jedem Zeitpunkt t wird also eine Position $(x(t), y(t), z(t))$ zugeordnet. Funktionen dieser Form nennt man auch **Kurven im Raum**. Betrachten hier konkret

$$x(t) := t - 0.4t^2, \quad y(t) = t^4, \quad z(t) = t^2 - 0.5t^3, \quad t \in [0, 2]. \quad (4.54)$$

Die Darstellung kann als 3D-Plot erfolgen, wobei der Zeitbezug verloren geht (analog zum 2D-Plot bei Kurven in der Ebene).

```

t = np.linspace(0, 2, 100)
x = t - 0.4 * t ** 2
y = t ** 4
z = t ** 2 - 0.5 * t ** 3

fig = go.Figure()
fig.add_trace(
    go.Scatter3d(x=x, y=y, z=z, mode='lines', line_width=5)
)
fig.update_layout(
    scene_xaxis_title_text='x',
    scene_yaxis_title_text='y',
    scene_zaxis_title_text='z',

```

```

    margin=dict(l=0, r=0, b=0, t=0) # breiten Rand vermeiden
)
fig.show()

```

Beispiel (Funktionen von zwei Veränderlichen)

Funktionen, die jedem Punkt in der Ebene eine Zahl zuordnen, können wahlweise als farbiger 2D-Plot oder als 3D-Plot dargestellt werden. Visualisieren hier als Beispiel

$$f : [0, 3] \times [0, 2] \rightarrow \mathbb{R}, \quad f(x, y) = (x^2 - 0.5x^3)(y - 0.4y^2). \quad (4.55)$$

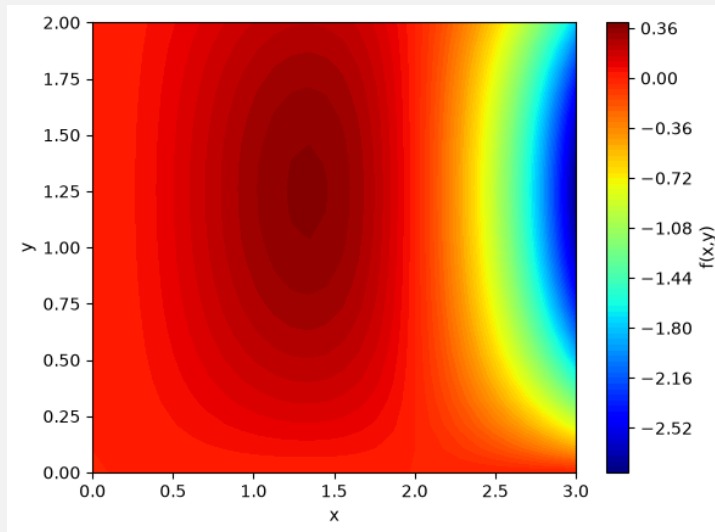
Bei der zweidimensionalen Darstellung werden die verschiedenen Funktionswerte durch Farben symbolisiert.

```

x, y = np.meshgrid(
    np.linspace(0, 3, 30), # Stützstellenanzahl nicht zu hoch
    np.linspace(0, 2, 20)
)
# x und y sind 2D -Arrays, die für jeden Gitterpunkt die x - bzw.
# y -Koordinate enthalten
f = (x ** 2 - 0.5 * x ** 3) * (y - 0.4 * y ** 2)

fig, ax = plt.subplots()
cont_plot = ax.contourf(
    x, y, f,
    levels=100, # Anzahl der Farbabstufungen
    cmap='jet' # Auswahl des Farbskala
)
plt.colorbar(cont_plot, label='f(x,y)') # Zuordnung der Farben
ax.set_xlabel('x')
ax.set_ylabel('y')
plt.show()

```



```
x, y = np.meshgrid(
    np.linspace(0, 3, 30),
    np.linspace(0, 2, 20)
)
f = (x ** 2 - 0.5 * x ** 3) * (y - 0.4 * y ** 2)

fig = go.Figure()
fig.add_trace(
    go.Surface(x=x, y=y, z=f, colorscale='Jet', showscale=False)
)
fig.update_layout(
    scene_xaxis_title_text='x',
    scene_yaxis_title_text='y',
    scene_zaxis_title_text='f(x,y)',
    margin=dict(l=0, r=0, b=0, t=0) # breiten Rand vermeiden
)
fig.show()
```

Bei der Visualisierung von Funktionen zweier Veränderlicher ist zu beachten, dass die Anzahl der Stützstellen nicht zu hoch gewählt wird. Wählt man beispielsweise in beiden Variablen je 100 Stützstellen, so entsteht ein zweidimensionales Gitter aus 10000 Punkten. Die Funktion wird also für die Visualisierung an diesen 10000 Stellen ausgewertet. Insbesondere die 3D-Darstellung von über 10000 Einzelobjekten, aus denen die Grafik im Webbrowser dann zusammengesetzt wird, kann den Webbrowser an seine Leistungsgrenzen bringen.

Die Möglichkeiten der Visualisierung sind nahezu grenzenlos. Wichtig ist, dass beim Betrachten einer Visualisierung einer Funktion stets geklärt wird, **was eigentlich dar-**

gestellt ist! Dass die grafische Darstellung nicht nur auf zahlenwertige Funktionen beschränkt ist, zeigt das nächste Beispiel.

Beispiel (Vektorfelder)

Strömungen in der Ebene können als Funktionen beschrieben werden, die jedem Punkt in der Ebene einen Vektor zuordnen. Die Richtung des Vektors gibt dabei die Richtung der Strömung an; die Länge des Vektors beschreibt die Geschwindigkeit der Strömung in dem Punkt. Die Visualisierung solcher Strömungen kann in einem 2D-Plot erfolgen, indem an einigen Punkten der Ebene die Strömungsvektoren eingezeichnet werden. Dabei wird zugunsten der besseren Übersichtlichkeit die Länge der Vektoren so skaliert, dass sie nicht überlappen. Dargestellt ist also nicht die tatsächliche Länge der Strömungsvektoren, sondern nur die relative Länge zueinander. Als Beispiel visualisieren wir die durch

$$f(x, y) = \begin{bmatrix} 1 \\ xy \end{bmatrix}, \quad x \in [-3, 3], \quad y \in [-2, 2] \quad (4.56)$$

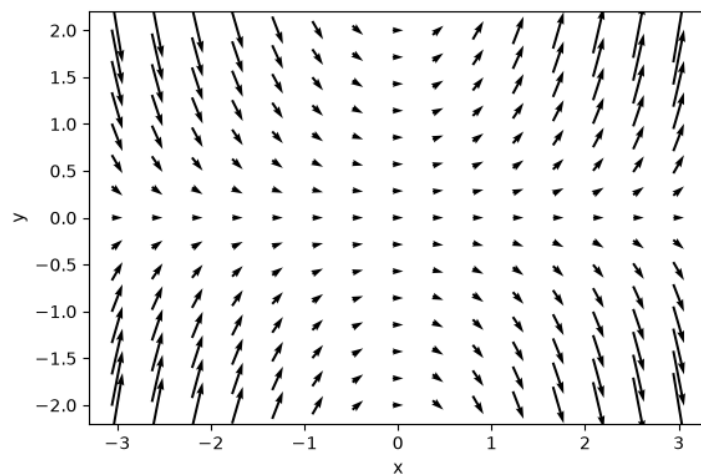
beschriebene Strömung. Der Umgang mit Vektoren wird später in der Veranstaltung genauer thematisiert.

```
# Gitterpunkte, an denen Strömungsvektoren dargestellt werden
↳ sollen
x, y = np.meshgrid(
    np.linspace(-3, 3, 15),
    np.linspace(-2, 2, 15)
)

# Strömungsvektoren
u = np.ones_like(x) # Array mit Einsen in der Form wie das Array
↳ x
v = x * y

# Plot
fig, ax = plt.subplots()
ax.quiver(
    x, y, u, v,
    pivot='middle', # Pfeilmitte an Gitterpunkt (nicht
↳ Pfeilanzfang)
    angles='xy' # Richtung der Pfeile folgt Verzerrung der Grafik
)
ax.set_aspect('equal', 'box') # gleiche Einheiten an beiden
↳ Achsen
ax.set_xlabel('x')
```

```
ax.set_ylabel('y')
plt.show()
```



Achtung!

Da bei der Visualisierung mit dem Computer die zu visualisierende Funktion stets nur an endlich vielen Stützstellen ausgewertet wird, können in ungünstigen Fällen wichtige Eigenschaften der Funktion in der Visualisierung fehlen. Auch muss insbesondere bei auf ganz $\mathbb{R}$ oder ganz $\mathbb{R}^2$ definierten Funktionen im Vorfeld geklärt werden, welcher Teil des Definitionsbereichs relevant ist.

Beispiel (ungünstige Visualisierung)

Die Funktion

$$f : \mathbb{R} \setminus \{1000\} \rightarrow \mathbb{R}, \quad f(x) = \frac{1}{(x - 1000)^{10}} \quad (4.57)$$

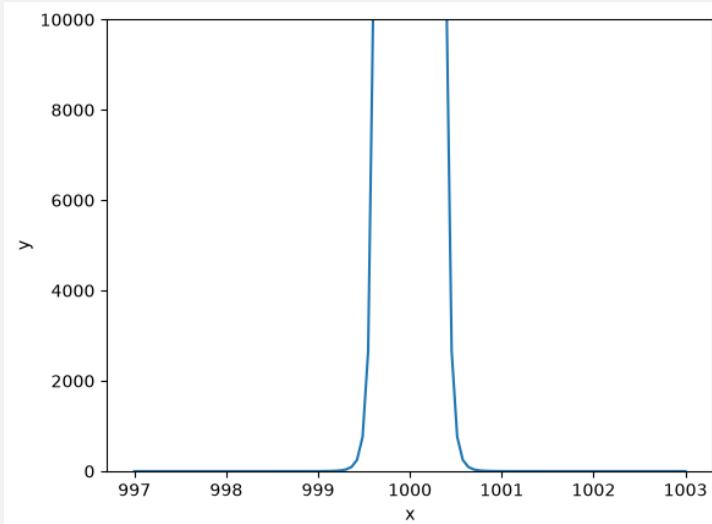
hat eine Polstelle bei $x = 1000$. Je nach Visualisierung ist diese sichtbar oder auch nicht.

Brauchbare Visualisierung:

```
x = np.linspace(997, 1003, 100)
y = 1 / (x - 1000) ** 10

fig, ax = plt.subplots()
ax.plot(x, y)
ax.set_ylim(0, 10000) # y -Achse oben abschneiden
```

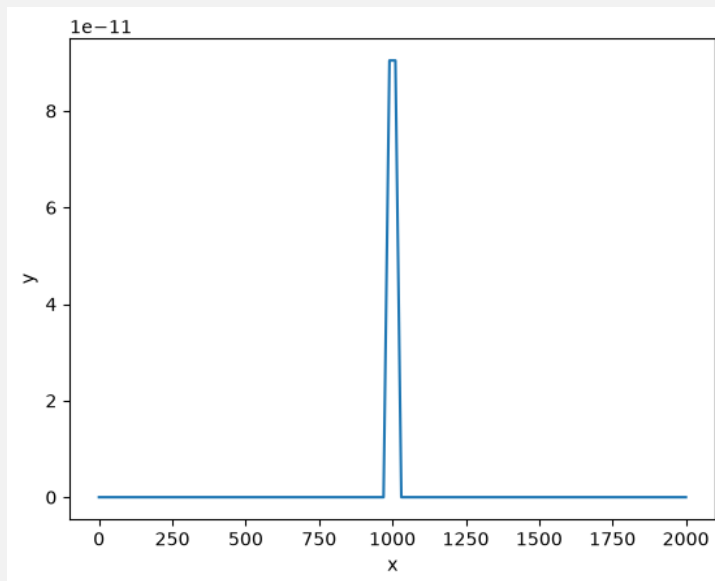
```
ax.set_xlabel('x')
ax.set_ylabel('y')
plt.show()
```



Unbrauchbare Visualisierung (Polstelle nicht als solche erkennbar; beachte, dass die dargestellten y -Werte alle praktisch Null sind):

```
x = np.linspace(0, 2000, 100)
y = 1 / (x - 1000) ** 10

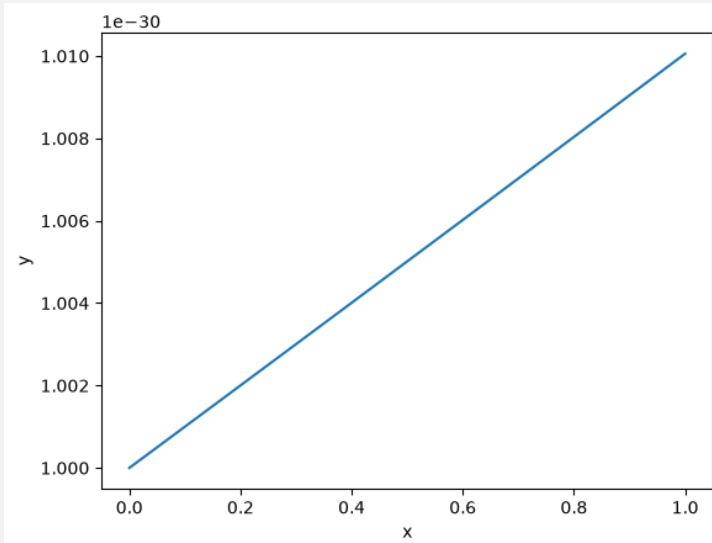
fig, ax = plt.subplots()
ax.plot(x, y)
ax.set_xlabel('x')
ax.set_ylabel('y')
plt.show()
```



Unbrauchbare Visualisierung (Polstelle nicht im dargestellten Bereich):

```
x = np.linspace(0, 1, 100)
y = 1 / (x - 1000) ** 10

fig, ax = plt.subplots()
ax.plot(x, y)
ax.set_xlabel('x')
ax.set_ylabel('y')
plt.show()
```



Einfache Eigenschaften

Betrachten einige Eigenschaften von Funktionen, deren Vorhandensein gelegentlich vorteilhaft ausgenutzt werden kann.

Nebenbemerkung

Vor dem Computerzeitalter mussten Funktionen aufwendig analytisch untersucht werden um eine Vorstellung vom Aussehen des Funktionsgraphs zu erhalten (Nullstellen, Extremstellen, Wendestellen, Monotonie, Krümmung, Polstellen, Asymptoten, Periodizität,...). Die Visualisierung mit dem Computer spart hier viel Mühe. Entsprechend werden wir im Rahmen dieser Veranstaltung keine klassischen “Kurvendiskussionen” durchführen.

Monotonie und Umkehrfunktion

Merke

Eine Funktion $f : D \rightarrow \mathbb{R}$ mit $D \subseteq \mathbb{R}$ heißt auf einer Teilmenge $A \subseteq D$

- **monoton wachsend**, wenn $f(x_1) \leq f(x_2)$ für alle $x_1, x_2 \in A$ mit $x_1 < x_2$ gilt,
- **streng monoton wachsend**, wenn $f(x_1) < f(x_2)$ für alle $x_1, x_2 \in A$ mit $x_1 < x_2$ gilt,

- **monoton fallend**, wenn $f(x_1) \geq f(x_2)$ für alle $x_1, x_2 \in A$ mit $x_1 < x_2$ gilt,
- **streng monoton fallend**, wenn $f(x_1) > f(x_2)$ für alle $x_1, x_2 \in A$ mit $x_1 < x_2$ gilt.

Die Monotonie von Funktionen spielt vor allem bei zwei Tätigkeiten eine Rolle: dem Lösen von Gleichungen und dem Umformen von Ungleichungen.

Soll eine Gleichung der Form

$$f(x) = y \quad (4.58)$$

für gegebenes $y \in \mathbb{R}$ gelöst werden und weiß man, dass die Funktion im Suchbereich für x (das kann z.B. ein Intervall oder ganz $\mathbb{R}$ sein) streng monoton ist, so kann man die **Umkehrfunktion** f^{-1} bilden und damit die Lösung direkt berechnen:

$$x = f^{-1}(y). \quad (4.59)$$

Für nicht streng monotone Funktionen ist die Sache schwieriger, da dann mehrere Lösungen existieren können, die nicht über die (dann nicht existierende) Umkehrfunktion berechnet werden können. Dieser Aspekt ist sowohl beim manuellen Rechnen als auch beim Lösen von Gleichungen mit dem Computer relevant.

Beispiel

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := x^2 \quad (4.60)$$

ist nicht monoton und es existiert auch keine Umkehrfunktion. Betrachten wir die gleiche Abbildungsvorschrift jedoch nur auf einer geeigneten Teilmenge der reellen Zahlen, so ist die Funktion streng monoton und besitzt demnach auch eine Umkehrfunktion (Video 440, Notizen dazu).

Gelegentlich möchte/muss man Ungleichungen umformen. Dabei ist zu beachten, dass manche Umformschritte, die bei Gleichungen zu einer äquivalenten Gleichung führen, bei Ungleichungen nur mit Vorsicht zu nutzen sind. Beispielsweise muss bei der Multiplikation beider Seiten der Ungleichung mit einer negativen Zahl das Relationszeichen “gedreht” werden (aus $\leq$ wird $\geq$ und umgekehrt; analog für strenge Ungleichheit). Allgemein gilt:

Merke

Das Anwenden einer streng monoton wachsenden Funktion auf beide Seiten einer Ungleichung ist eine äquivalente Umformung.

Das Anwenden einer Streng monoton fallenden Funktion auf beide Seiten einer Ungleichung ist eine äquivalente Umformung, wenn dabei das Relationszeichen “gedreht” wird.

Beispiel

Wir möchten die Ungleichung

$$\frac{-2}{x} \leq 4 \quad (4.61)$$

nach x umstellen, d.h. durch äquivalente Umformungen auf die Form $x \leq \dots$ oder $x \geq \dots$ bringen. Alle Umformschritte können wir als Anwenden geeigneter Funktionen auf beide Seiten der Ungleichung formulieren:

1. Wir wenden die Funktion $u \mapsto -\frac{1}{2}u$ auf beide Seiten an. Diese Funktion ist streng monoton fallend, also muss das Relationszeichen gedreht werden:

$$\frac{1}{x} \geq -\frac{1}{2} \cdot 4. \quad (4.62)$$

2. Bilden nun auf beiden Seiten das Reziproke, d.h. wir wenden die Funktion $u \mapsto \frac{1}{u}$ auf beide Seiten an. Dies ist ebenfalls eine streng monoton fallende Funktion. Also muss das Relationszeichen nochmal gedreht werden:

$$x \leq \frac{1}{-\frac{1}{2} \cdot 4}. \quad (4.63)$$

Erhalten also die äquivalent umgeformte Ungleichung

$$x \leq -\frac{1}{2}. \quad (4.64)$$

Für weitere Beispiele und Erläuterungen zur Monotonie von Funktionen sei auf das entsprechende OMB+-Kapitel verwiesen.

Periodizität**Merke**

Eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ heißt **periodisch**, wenn es eine Zahl $T > 0$ gibt, sodass

$$f(x + T) = f(x) \quad \text{für alle } x \in \mathbb{R} \quad (4.65)$$

gilt. Jede Zahl T mit dieser Eigenschaft heißt **Periode** von f . Die kleinste Zahl T mit dieser Eigenschaft heißt **Grundperiode** von f .

Ist T eine Periode einer Funktion, so genügt die Betrachtung der Funktion auf dem Intervall $[0, T)$. Auf dem Rest der reellen Achse ist das Verhalten der Funktion vollständig durch ihr Verhalten auf diesem Intervall bestimmt.

Wichtige periodische Funktionen sind die Sinus- und die Cosinus-Funktion (siehe unten)

sowie darauf basierende komplexere Funktionen.

Beschränktheit

Merke

Eine Funktion $f : D \rightarrow \mathbb{R}$ mit $D \subseteq \mathbb{R}^n$ heißt beschränkt, wenn es eine Zahl $c > 0$ gibt, sodass

$$|f(x)| \leq c \quad \text{für alle } x \in D \quad (4.66)$$

gilt.

Nullstellen

Merke

Für eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ heißt $x \in \mathbb{R}^m$ **Nullstelle** der Funktion, wenn $f(x) = 0$ gilt.

Nebenbemerkung

Beachte, dass wir in dieser Veranstaltung keine besondere Symbolik für die Unterscheidung von Zahlen und Punkten bzw. Vektoren verwenden. In der Ingenieurliteratur sieht man gelegentlich Unterstreichungen $\underline{x}$ bzw. $\underline{0}$ oder aufgesetzte Pfeile $\vec{x}$ bzw. $\vec{0}$ oder Fettsetzungen $\mathbf{x}$ bzw. $\mathbf{0}$ für Vektoren. Diese führen jedoch an einigen Stellen zu Komplikationen und machen komplexe Formeln unübersichtlicher. Da ohnehin beim Betrachten eines mathematischen Ausdrucks stets im Detail zu klären ist, welches Symbol welche Bedeutung hat, ergibt sich die Antwort auf die Frage “Zahl oder Punkt/Vektor?” gleich mit.

In der obigen Definition des Begriffs Nullstelle steht die 0 also für den Nullvektor (sofern $m > 1$) und auch x ist ein Punkt/Vektor, wenn $n > 1$.

Die Frage nach den Nullstellen einer Funktion ist eng verknüpft mit dem Lösen von Gleichungen

$$f(x) = y, \quad (4.67)$$

wobei $y \in \mathbb{R}^m$ gegeben und $x \in \mathbb{R}^n$ gesucht ist. Solche Gleichungen kann man als

$$g(x) = 0 \quad (4.68)$$

schreiben, wenn man $g(x) := f(x) - y$ setzt. Beim Lösen von Gleichungen mit dem Computer stehen oft nur Funktionen für die Nullstellenbestimmung zur Verfügung, sodass Gleichungen, wie eben gezeigt, entsprechend umgeformt werden müssen.

Nullstellen von Funktionen können beispielsweise mit `scipy.optimize.fsolve` aus dem Python-Paket SciPy numerisch berechnet werden. Beachte, dass stets nur eine Nullstelle geliefert wird, auch wenn es mehrere gibt. Welche Nullstelle man erhält, kann man durch Vorgabe einer groben Schätzung der Nullstelle steuern.

Beispiel

Zum Lösen der Gleichung

$$x^2 + x + 1 = \frac{3}{2} \quad (4.69)$$

können wir folgenden Python-Code verwenden:

```
import scipy.optimize as so

def f(x):
    return x ** 2 + x + 1 - 3 / 2

print(so.fsolve(f, 1))
print(so.fsolve(f, -1))
```

```
[0.3660254]
[ -1.3660254]
```

Die Anzahl der Lösungen und geeignete Näherungswerte (hier 1 und -1) sind nicht immer klar ersichtlich. Deshalb birgt das numerische Lösen von Gleichungen stets das Risiko, dass man nicht alle Lösungen findet!

Beachte, dass `fsolve` auch zum Lösen von (nichtlinearen) Gleichungssystemen verwendet werden kann. Dies entspricht dem Fall $m > 1$ und $n > 1$ in der obigen Definition des Begriffs Nullstelle.

Das symbolische Lösen von Gleichungen ist mit SymPy möglich. Für das obige Beispiel:

```
import sympy

x = sympy.symbols('x')

sympy.solve(x ** 2 + x + 1 - sympy.Rational(3, 2), x)
```

```
[ -1/2 + sqrt(3)/2, -sqrt(3)/2 - 1/2]
```

Beim symbolischen Rechnen erhält man üblicherweise alle Lösungen. Jedoch werden nicht immer Lösungen gefunden, obwohl sie existieren. Beim numerischen Rechnen hingegen erhält man praktisch immer wenigstens eine Lösung (sofern eine existiert).

Grenzwerte

Die Betrachtung von Grenzwerten bei Funktionen spielt eine zentrale Rolle in der Mathematik und ihren Anwendungen. So basiert z.B. der in Physik und anderen (Natur-)Wissenschaften ganz grundlegende Begriff der Differenzierbarkeit einer Funktion auf der Betrachtung gewisser Grenzwerte.

Grenzwerte bei Funktionen sind sehr ähnlich zu Grenzwerten bei Zahlenfolgen, jedoch gibt es Unterschiede im Detail, denn der kontinuierliche Charakter der reellen Zahlen erlaubt auch andere Grenzübergänge als nur " $\rightarrow \infty$ ".

Merke (Grenzwerte im Unendlichen)

Sei $f : D \rightarrow \mathbb{R}$ mit $D \subseteq \mathbb{R}$ eine Funktion, sodass D ein Intervall der Form $[a, \infty)$ enthält. Eine Zahl $g \in \mathbb{R}$ heißt **Grenzwert** der Funktion f bei Annäherung des Arguments an ∞ , wenn für jede (!) Zahlenfolge $(x_n)_{n \in \mathbb{N}}$ in D mit $\lim_{n \rightarrow \infty} x_n = \infty$ die zugehörige Folge $(f(x_n))_{n \in \mathbb{N}}$ der Funktionswerte gegen g konvergiert. Man schreibt dann

$$\lim_{x \rightarrow \infty} f(x) = g. \quad (4.70)$$

Analog ist der Grenzwert bei Annäherung an $-\infty$ erklärt, wenn D ein Intervall der Form $(-\infty, b]$ enthält.

Nicht jede Funktion hat Grenzwerte im Unendlichen. Beispielsweise existieren die beiden Grenzwerte nicht für (nicht konstante) periodische Funktionen. Auch können die Grenzwerte ∞ oder $-\infty$ sein; in diesem Fall sagt man ebenfalls, dass die Grenzwerte nicht (als Zahl) existieren.

Merke (Grenzwerte an einer Stelle)

Sei $f : D \rightarrow \mathbb{R}$ mit $D \subseteq \mathbb{R}$ eine Funktion. Eine Zahl $g \in \mathbb{R}$ heißt **Grenzwert** der Funktion f bei Annäherung an die Stelle $x_0 \in \mathbb{R}$, wenn für jede Zahlenfolge $(x_n)_{n \in \mathbb{N}}$ in D mit $\lim_{n \rightarrow \infty} x_n = x_0$ die zugehörige Folge $(f(x_n))_{n \in \mathbb{N}}$ der Funktionswerte gegen g konvergiert. Man schreibt dann

$$\lim_{x \rightarrow x_0} f(x) = g. \quad (4.71)$$

Betrachtet man nur Folgen $(x_n)_{n \in \mathbb{N}}$ mit $x_n \leq x_0$ für alle $n \in \mathbb{N}$, so spricht man vom **linksseitigen Grenzwert** und schreibt

$$\lim_{x \rightarrow x_0 - 0} f(x) = g. \quad (4.72)$$

Betrachtet man nur Folgen $(x_n)_{n \in \mathbb{N}}$ mit $x_n \geq x_0$ für alle $n \in \mathbb{N}$, so spricht man vom **rechtsseitigen Grenzwert** und schreibt

$$\lim_{x \rightarrow x_0 + 0} f(x) = g. \quad (4.73)$$

Man kann zeigen, dass der (beidseitige) Grenzwert genau dann existiert, wenn links- und rechtsseitiger Grenzwert existieren und übereinstimmen.

Ein relativ häufiges Phänomen sind Funktionen mit Polstellen. Für eine Funktion $f : [a, b] \setminus \{u\} \rightarrow \mathbb{R}$ heißt u **Polstelle**, wenn die einseitigen Grenzwerte bei Annäherung an diese Stelle nur im Sinne von $\pm\infty$ existieren. Da Funktionen in der Praxis fast immer Messgrößen (Temperatur, Länge, Druck, Kosten,...) beschreiben, erscheinen Polstellen eher unnatürlich und deren Auftreten entsprechend unwahrscheinlich. Sie sind jedoch ein wichtiges Modellierungswerkzeug!

Beispiel (Massepunkt)

In der Physik vereinfacht man räumlich ausgedehnte Objekte gern zu Punkten im Raum. So betrachtet man bei der Berechnung von Flugbahnen nur selten die Masseverteilung des sich bewegenden Objekts im Raum, sondern modelliert das Objekt als “Massepunkt”, also als einen einzigen Punkt im Raum mit einer gegebenen Masse. Analog verfährt man beispielsweise mit Ladungen in der Elektrotechnik (“Punktladung”).

Die Anziehungskraft, die ein als Massepunkt modellierter Himmelskörper (z.B. Erde) auf andere Objekte (z.B. Satelliten) ausübt, kann durch sein Gravitationsfeld beschrieben werden. Die Anziehungskraft ist immer zum Massepunkt hin gerichtet. Der Betrag der Kraft hängt von der Entfernung zum Massepunkt ab und kann somit als Funktion $F : (0, \infty) \rightarrow \mathbb{R}$ ausgedrückt werden:

$$F(r) = \frac{G M}{r^2}. \quad (4.74)$$

Dabei sind $G > 0$ die Gravitationskonstante und $M > 0$ die Masse des Himmelskörpers. Die Funktion F hat eine Polstelle bei $r = 0$.

Natürlich wirkt in der Nähe des Erdmittelpunktes (genauer: Massenschwerpunktes der Erde) keine beliebig hohe Gravitationskraft. Sondern die Polstelle ist ein Resultat der vereinfachten Modellierung der Erde als Massepunkt. In der Nähe des Erdmittelpunktes ist das gewählte Modell zur Beschreibung der wirkenden Kräfte ungeeignet (vgl. auch Newtonsches Kugelschalentheorem).

Der Begriff des Grenzwerts spielt auch bei Funktionen mehrerer Veränderlicher eine Rolle. Dann ist jedoch nur der (allseitige) Grenzwert bei Annäherung an eine Stelle $x_0 \in \mathbb{R}^n$ sinnvoll einführbar. Die dazu benötigte Konvergenz einer Punktfolge $(x_k)_{k \in \mathbb{N}}$ in $\mathbb{R}^n$ gegen x_0 ist gleichbedeutend mit der Konvergenz jeder einzelnen Koordinatenfolge (vgl. auch Aufgabe 3.6(c) in Übung 3). Schreiben wir jeden Punkt x_k als

$$x_k = (x_k^{(1)}, \dots, x_k^{(n)}) \in \mathbb{R}^n, \quad (4.75)$$

so bedeutet $\lim_{k \rightarrow \infty} x_k = x_0$ also, dass

$$\lim_{k \rightarrow \infty} x_k^{(l)} = x_0^{(l)} \quad (4.76)$$

für jedes $l \in \{1, \dots, n\}$ gilt.

Nebenbemerkung

Bei der Arbeit mit Folgen von Punkten (oder Vektoren) treten mehrere verschiedene Indizes und Dimensionszahlen auf. Beim Lesen und Verstehen entsprechender Formeln muss im Detail geklärt werden, welches Symbol welche Rolle übernimmt. Im hier betrachteten Fall sind

- n die Anzahl der Veränderlichen der Funktion (nicht der Folgenindex!),
- k der Folgenindex,
- l der Index der Koordinaten eines Punktes.

Insbesondere wird die Schreibweise mit zwei Indizes in der Literatur nicht einheitlich gehandhabt. In unserem Fall ist der untere Index in $x_k^{(l)}$ der Folgenindex und der obere Index ist die Koordinatennummer. Dies kann auch umgekehrt gehandhabt werden! Die Klammern um den oberen Index sollen Verwechslungen mit Potenzen vermeiden.

Merke (Grenzwerte bei mehreren Veränderlichen)

Sei $f : D \rightarrow \mathbb{R}^m$ mit $D \subseteq \mathbb{R}^n$ eine Funktion. Ein Punkt $g \in \mathbb{R}^m$ heißt **Grenzwert** der Funktion f bei Annäherung an die Stelle $x_0 \in \mathbb{R}^n$, wenn für jede Punktfolge $(x_n)_{n \in \mathbb{N}}$ in D mit $\lim_{n \rightarrow \infty} x_n = x_0$ die zugehörige Punktfolge $(f(x_n))_{n \in \mathbb{N}}$ der Funktionswerte gegen g konvergiert. Man schreibt dann

$$\lim_{x \rightarrow x_0} f(x) = g. \quad (4.77)$$

Umgangssprachlich formuliert bedeutet dies, dass völlig unabhängig davon wie man sich der Stelle x_0 annähert, das Verhalten der Funktionswerte immer gleich ist. Hier ein Gegenbeispiel:

Beispiel

Für die Funktion

$$f : \mathbb{R}^2 \setminus \{(0, 0)\} \rightarrow \mathbb{R}, \quad f(x, y) := \frac{xy}{x^2 + y^2} \quad (4.78)$$

existiert in $(0, 0)$ kein Grenzwert (Video 450, Notizen dazu).

Beachte, dass die grafische Darstellung in der Nähe von $(0,0)$ fehlerhaft ist! Ursache ist das Auswerten der Funktion an nur endlich vielen Stellen.

Stetigkeit

Merke

Eine Funktion $f : D \rightarrow \mathbb{R}$ mit $D \subseteq \mathbb{R}^n$ heißt **stetig im Punkt** $x_0 \in D$, wenn der Grenzwert $\lim_{x \rightarrow x_0} f(x)$ existiert und identisch mit dem Funktionswert $f(x_0)$ ist.

Eine Funktion $f = (f_1, \dots, f_m) : D \rightarrow \mathbb{R}^m$ mit $D \subseteq \mathbb{R}^n$ heißt stetig im Punkt $x_0 \in D$, wenn jede Komponente f_l , $l = 1, \dots, m$, stetig in x_0 ist.

Eine Funktion $f : D \rightarrow \mathbb{R}^m$ mit $D \subseteq \mathbb{R}^n$ heißt **stetig**, wenn sie in jedem Punkt ihres Definitionsbereichs D stetig ist.

Stetigkeit einer Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ bedeutet anschaulich, dass man den Graph der Funktion ohne Absetzen mit einem Stift zeichnen kann, der Graph also beispielsweise keine Sprünge aufweist. Unbedingt zu beachten ist, dass Stetigkeit nur Aussagen zum Verhalten der Funktion in ihrem Definitionsbereich macht. Der Definitionsbereich selbst kann aber beispielsweise Lücken haben.

Beispiel

Die Funktion

$$f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}, \quad f(x) := \frac{1}{x} \quad (4.79)$$

ist stetig, obwohl sie eine Polstelle bei $x = 0$ besitzt. Die Stelle $x = 0$ gehört nicht zum Definitionsbereich der Funktion und findet daher bei der Frage nach Stetigkeit keine Beachtung.

Hingegen ist die Funktion

$$g : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \begin{cases} \frac{1}{x} & \text{für } x \neq 0, \\ 0 & \text{für } x = 0 \end{cases} \quad (4.80)$$

nicht stetig, da sie in $x = 0$ nicht stetig ist (Grenzwert $\lim_{x \rightarrow 0} f(x)$ existiert nicht).

Unstetigkeitsstellen bei Funktionen einer reellen Veränderlichen sind im Wesentlichen entweder Polstellen (vgl. Funktion g im Beispiel) oder Sprungstellen (links- und rechtsseitige Grenzwerte existieren, stimmen aber nicht beide mit dem Funktionswert überein). Für Funktionen mehrerer Veränderlicher ist die Klassifikation von Unstetigkeitsstellen nicht mehr so einfach machbar.

Beispiel

Die Funktion

$$f : \mathbb{R}^2 \rightarrow \mathbb{R}, \quad f(x, y) := \begin{cases} \frac{xy}{x^2+y^2} & \text{für } (x, y) \neq (0, 0), \\ 0 & \text{für } (x, y) = (0, 0) \end{cases} \quad (4.81)$$

ist nicht stetig in $(0, 0)$, da dort der Grenzwert $\lim_{(x,y) \rightarrow (0,0)} f(x, y)$ nicht existiert (vgl. früheres Beispiel zu Grenzwerten).

Der Begriff der Stetigkeit ist innerhalb der Mathematik von enormer Bedeutung, wird aber bei der täglich Arbeit im Kontext von Anwendungsproblemen nur am Rande direkt benötigt. Die Wichtigkeit des Begriffs resultiert aus der Tatsache, dass bei stetigen Funktionen der Funktionswert an einer Stelle auch Informationen über die Funktionswerte an benachbarten Stellen liefert. Insbesondere mit Blick auf den Computereinsatz, bei dem Funktionen stets nur an endlich vielen Stellen ausgewertet werden können, ist Stetigkeit eine Art Minimalforderung, um eine Chance auf brauchbare Resultate zu haben. Beispiele für das Auftreten des Stetigkeitsbegriffs in anwendungsbezogenen Situationen:

- Viele mathematische Aussagen gelten nur, wenn die beteiligten Funktionen stetig sind. So ist Stetigkeit beispielsweise eine wichtige Voraussetzung (aber nicht die einzige!) für die Existenz von Lösungen zu Optimierungsproblemen. Bei der mathematischen Modellierung von praktischen Fragestellungen ist deshalb darauf zu achten, dass die in einem Modell auftretenden Funktionen möglichst stetig sind.
- Beim Rechnen mit Computer oder Taschenrechner oder auch beim manuellen Rechnen mit gemessenen und somit fehlerbehafteten Werten werden die Berechnungen stets mit gerundeten Werten durchgeführt. Das Ausführen von Berechnungen kann man als Auswerten einer geeigneten Funktion an einer Stelle interpretieren. Die Stetigkeit einer solchen Funktion garantiert, dass die Funktionswerte an den gerundeten Stellen nicht zu sehr von den Funktionswerten an den (unbekannten) exakten Stellen abweichen. Wie groß die Rundungsfehler sein dürfen, damit die Abweichungen in den Funktionswerten nicht zu groß werden, ist eine wichtige Frage. Dieser werden wir später nachgehen, wenn wir die geeigneten mathematischen Werkzeuge zur Beantwortung zur Verfügung haben.
- Fotografische Bilddaten werden heute praktisch immer in stark komprimierter Form gespeichert. Welche Kompressionsverfahren in Frage kommen und wie stark ein Bild ohne größere Qualitätsverluste komprimiert werden kann, hängt auch an den Stetigkeitseigenschaften des kontinuierlichen (d.h. nicht in Pixel zerlegten) Bildes. Für den heute üblichen JPEG-Standard zur Bilddatenkompression kann man grob formulieren: je “stetiger” das Bild, desto stärker kann ohne nennenswerte Qualitätsverluste komprimiert werden.

Beispiele elementarer Funktionen

Elementare Funktionen und ihre Eigenschaften wurden im Rahmen der Schulmathematik behandelt. Deshalb verweisen wir zur Wiederholung auf das entsprechenden Kapitel von OMB+ und gehen hier nur auf einzelne Beispiele und Aspekte nochmal genauer ein:

Polynome Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ der Form

$$f(x) = a_n x^n + a_{n-1} x^{n-1} + \cdots + a_2 x^2 + a_1 x + a_0 \quad (4.82)$$

heißen **Polynom**. Dabei ist die gegebene Zahl $n \in \mathbb{N}$ der **Grad** des Polynoms und die gegebenen $n + 1$ Zahlen $a_0, a_1, \dots, a_n \in \mathbb{R}$ sind die **Koeffizienten** des Polynoms.

Jedes Polynom hat höchstens n Nullstellen. Jeder Polynomnullstelle kann eine sogenannte Vielfachheit zugeordnet werden: Die Polynomnullstelle x_0 hat die **Vielfachheit** $k \in \mathbb{N}$, wenn man das Polynom in der Form

$$f(x) = (x - x_0)^k (\text{Polynom vom Grad } n - k) \quad (4.83)$$

schreiben kann und k die größte Zahl ist, die eine solche Darstellung erlaubt. Der Faktor $x - x_0$ wird auch als **Linearfaktor** bezeichnet. Sind $x_1, \dots, x_m$ die Nullstellen des Polynoms f und $k_1, \dots, k_m$ die zugehörigen Vielfachheiten, so gilt

$$f(x) = (x - x_1)^{k_1} \cdots (x - x_m)^{k_m} (\text{Polynom ohne Nullstellen}). \quad (4.84)$$

Nebenbemerkung

Man kann den Zahlenraum, der sich bei uns auf die reellen Zahlen $\mathbb{R}$ beschränkt, noch erweitern zu den sogenannten komplexen Zahlen. In diesem erweiterten Zahlenraum hat jedes Polynom n -ten Grades genau n Nullstellen, wenn man jede Nullstelle entsprechend ihrer Vielfachheit zählt. Man erhält dann eine Zerlegung des Polynoms in Linearfaktoren mit einem Polynom 0-ten Grades (konstantes Polynom) als nullstellenfreies Restpolynom. Die Theorie der Polynome über komplexen Zahlen ist ein wichtiger Baustein innerhalb der Mathematik, der auch gelegentlich an der für Anwender sichtbaren Oberfläche erscheint.

Lineare Funktionen mehrerer Veränderlicher Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$ der Form

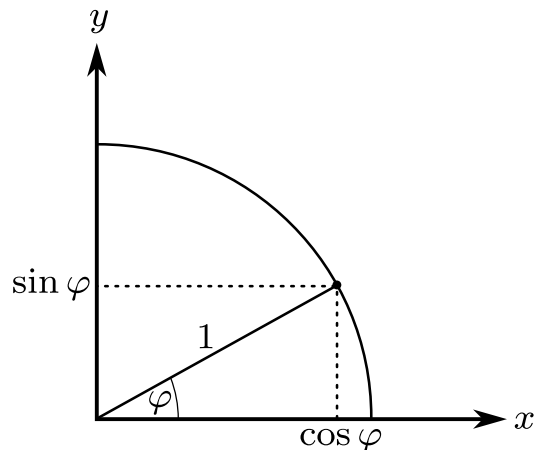
$$f(x_1, \dots, x_n) = a_1 x_1 + \cdots + a_n x_n \quad (4.85)$$

heißen **linear**. Dabei sind $a_1, \dots, a_n \in \mathbb{R}$ gegebene Zahlen.

Für $n = 2$ ist der Graph einer linearen Funktion stets eine Ebene in $\mathbb{R}^3$. Für $n > 2$ bezeichnet man die Graphen linearer Funktionen als **Hyperebenen** (in $\mathbb{R}^{n+1}$).

Lineare Funktionen werden aufgrund ihrer einfachen Struktur gern als näherungsweiser Ersatz für kompliziertere Funktionen verwendet.

Trigonometrische Funktionen Häufig benötigte Zahlen, deren direkte Erklärung zu viel Zeit benötigt, bekommen in der Mathematik üblicherweise einen Namen. So benötigt man zu einem gegebenen Winkel z.B. sehr oft die Zahlen, die wir als Cosinus und Sinus des Winkels kennen. Cosinus und Sinus eines Winkels sind gerade die x - bzw. y -Koordinate des Punktes, der entsteht indem man die Entfernung 1 vom Koordinatenursprung aus entlang eines Strahls mit Winkel φ zur (positiven) x -Achse zurücklegt.



Cosinus und Sinus geben die x - und die y -Koordinate eines Punktes auf dem Einheitskreis in Abhängigkeit des Winkels φ zwischen dem Strahl vom Ursprung durch den Punkt und der x -Achse an.

Beträgt die Entfernung des Punktes nicht 1, sondern allgemeiner $r > 0$, so ergeben sich die Koordinaten als

$$\begin{aligned} x &= r \cos \varphi, \\ y &= r \sin \varphi. \end{aligned} \tag{4.86}$$

Die Darstellung eines Punktes in der Ebene durch Abstand $r > 0$ zum Ursprung und Winkel $\varphi \in [0, 2\pi)$ zur x -Achse bezeichnet man als **Polarkoordinaten**-Darstellung des Punktes. Sie wird insbesondere in der Physik gern verwendet (Drehbewegungen und andere periodische Vorgänge), spielt aber beispielsweise auch in der Geodatenverarbeitung eine Rolle (diverse Kartenprojektionen, z.B. Azimuthalprojektionen).

Offensichtlich gilt stets

$$(\sin \varphi)^2 + (\cos \varphi)^2 = 1 \tag{4.87}$$

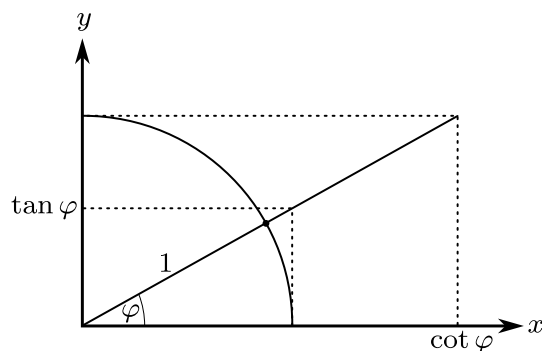
(trigonometrische Form des Satzes des Pythagoras).

Auf Basis von Cosinus und Sinus definiert man noch den Tangens

$$\tan \varphi := \frac{\sin \varphi}{\cos \varphi} \tag{4.88}$$

und den Cotangens

$$\cot \varphi := \frac{\cos \varphi}{\sin \varphi}. \tag{4.89}$$



Tangens und Cotangens sind gerade der Anstieg des Strahls durch Ursprung und Punkt bzw. das Reziproke des Anstiegs.

Sinus und Cosinus können auch für Winkel außerhalb von $[0, 2\pi)$ bestimmt werden, indem der fragliche Winkel durch Verschiebung um Vielfache von 2π in den Bereich $[0, 2\pi)$ gebracht wird. Die so entstehende **Sinusfunktion** und die **Cosinusfunktion** sind auf ganz $\mathbb{R}$ definiert und haben beide den Wertebereich $[-1, 1]$.

Sinusfunktion und Cosinusfunktion sind nicht monoton auf $[0, 2\pi)$. Auf geeigneten Teilintervallen sind sie jedoch streng monoton, sodass dort eine Umkehrfunktion existiert. Mit $\arcsin$ bezeichnet man die Umkehrfunktion der Sinusfunktion auf dem Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$. Die Umkehrfunktion der Cosinusfunktion auf dem Intervall $[0, \pi]$ wird als $\arccos$ bezeichnet. Für den Tangens wiederum existiert eine Umkehrfunktion auf $[-\frac{\pi}{2}, \frac{\pi}{2}]$, die mit $\arctan$ bezeichnet wird. Eine Umkehrfunktion des Cotangens kann analog eingeführt werden, spielt aber nur sehr selten eine Rolle.

Exponentialfunktion Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := e^x \quad (4.90)$$

mit der Euler'schen Zahl $e \approx 2.718$ wird als **Exponentialfunktion** bezeichnet. Sie taucht (meist unerwartet) in ganz vielen Situationen in praktisch allen Teilbereichen der Mathematik auf. Wir werden ihr insbesondere bei der Differentialrechnung und bei der Wahrscheinlichkeitsrechnung begegnen.

Die Exponentialfunktion ist streng monoton wachsend und besitzt somit eine Umkehrfunktion. Diese wird als (natürliche) **Logarithmusfunktion** bezeichnet und mit dem Symbol $\ln$ abgekürzt. Es gilt also stets

$$\ln(e^x) = x \quad \text{für } x \in \mathbb{R} \quad \text{und} \quad e^{\ln x} = x \quad \text{für } x > 0. \quad (4.91)$$

Nebenbemerkung (Beziehung zu Sinus und Cosinus)

Die Exponentialfunktion kann als Polynom mit “unendlichem Grad” angesehen werden: Zu $x \in \mathbb{R}$ und $n \in \mathbb{N}$ können wir das Polynom

$$P_n(x) := \sum_{k=0}^n \frac{1}{k!} x^k \quad (4.92)$$

vom Grad n bilden. Dabei ist

$$n! := 1 \cdot 2 \cdots n \quad (4.93)$$

das Produkt der Zahlen 1 bis n (sprich: “ n Fakultät”), wobei $0! := 1$ gesetzt wird, und $\sum$ ist ein Symbol zur kurzen Schreibweise für lange Summen:

$$\sum_{k=0}^n \frac{1}{k!} x^k := \frac{1}{0!} x^0 + \frac{1}{1!} x^1 + \cdots + \frac{1}{n!} x^n. \quad (4.94)$$

Man kann nun zeigen (tun wir hier nicht), dass

$$\lim_{n \rightarrow \infty} P_n(x) = e^x \quad (4.95)$$

gilt, die Polynome für festes $x \in \mathbb{R}$ bei fortwährender Erhöhung des Grades also gegen die Exponentialfunktion konvergieren.

In diesem Kontext wird auch eine enge Beziehung zu Sinus- und Cosinusfunktion sichtbar. Entfernt man alle ungeraden x -Potenzen und wechselt das Vorzeichen vor den geraden Potenzen, so erhält man die Cosinusfunktion:

$$\lim_{n \rightarrow \infty} \sum_{k=0}^n (-1)^k \frac{1}{(2k)!} x^{2k} = \cos x. \quad (4.96)$$

Analog liefert das Weglassen der geraden Potenzen die Sinusfunktion:

$$\lim_{n \rightarrow \infty} \sum_{k=0}^n (-1)^k \frac{1}{(2k+1)!} x^{2k+1} = \sin x. \quad (4.97)$$

Ohne Verzeichnungswechsel erhält man die beiden als **Cosinus hyperbolicus** und **Sinus hyperbolicus** bezeichneten Funktionen:

$$\lim_{n \rightarrow \infty} \sum_{k=0}^n \frac{1}{(2k)!} x^{2k} = \cosh x, \quad (4.98)$$

$$\lim_{n \rightarrow \infty} \sum_{k=0}^n \frac{1}{(2k+1)!} x^{2k+1} = \sinh x. \quad (4.99)$$

Für diese gilt dann

$$e^x = \cosh x + \sinh x. \quad (4.100)$$

Cosinus hyperbolicus und **Sinus hyperbolicus** tauchen gelegentlich in Anwendungen auf. So zeigt beispielsweise ein locker zwischen zwei Punkten hängendes Seil gerade die Form des Graphs des Cosinus hyperbolicus (vgl. Kettenlinie).

5 Vektoren und Matrizen

Die in diesem Kapitel vorgestellten Werkzeuge und Zusammenhänge werden auch unter dem Begriff “lineare Algebra” zusammengefasst.

5.1 Formalismus

Wir führen hier zunächst einige Begriffe und Schreibweisen ein, die ohne weitere Erklärungen nicht sonderlich sinnvoll oder gar nützlich erscheinen. In den nachfolgenden Unterkapitel füllen wir diese dann mit Leben und Sinn. So haben wir einerseits den “trockenen” Teil schnell hinter uns und konzentriert an einer Stelle zum späteren Nachschlagen; andererseits stören dann keine eher technischen als inhaltlichen Definitionen den Erklärfluss.

Vektoren

Der Begriff des Vektors wurde bereits in Section 4.2 als “Liste” von Zahlen eingeführt. Die Menge aller Vektoren der Länge n bezeichnen wir mit $\mathbb{R}^n$. Beachte, dass die Unterscheidung zwischen Punkt und Vektor nur eine Interpretationsfrage ist. Aus technischer Sicht besteht keinerlei Unterschied, da es sich in beiden Fällen um eine Liste von Zahlen handelt.

Für Vektoren führen wir zunächst zwei Rechenoperationen ein: die Addition zweier Vektoren und die Multiplikation eines Vektors mit einer Zahl.

Merke

Für zwei gleich lange Vektoren $a, b \in \mathbb{R}^n$ setzen wir

$$a + b := \begin{bmatrix} a_1 + b_1 \\ \vdots \\ a_n + b_n \end{bmatrix} \quad (5.1)$$

und bezeichnen den so entstehenden Vektor $a + b \in \mathbb{R}^n$ als **Summe** der beiden Vektoren a und b .

Für eine Zahl $\lambda \in \mathbb{R}$ und einen Vektor $a \in \mathbb{R}^n$ setzen wir

$$\lambda a := \begin{bmatrix} \lambda a_1 \\ \vdots \\ \lambda a_n \end{bmatrix} \quad (5.2)$$

und bezeichnen den so entstehenden Vektor $\lambda a \in \mathbb{R}^n$ als **skalares Vielfaches** des Vektors a mit Faktor λ .

Beide Begriffe sind also einfach komponentenweise definiert. Beachte, dass weiter unten noch das Skalarprodukt zweier Vektoren eingeführt werden wird. Dieses ist eine ganz andere Rechenoperation als das Bilden des skalaren Vielfachen!

Mit den beiden eingeführten Rechenoperationen können wir auch komplexere Ausdrücke formulieren. Eine häufig benötigte Kombination von Additionen und skalaren Vielfachen

ist die folgende:

Merke

Sind $a_1, \dots, a_k \in \mathbb{R}^n$ Vektoren und sind $\lambda_1, \dots, \lambda_k \in \mathbb{R}$ Zahlen, so bezeichnet man den Ausdruck

$$\sum_{l=1}^k \lambda_l a_l \quad (= \lambda_1 a_1 + \dots + \lambda_k a_k) \quad (5.3)$$

als **Linearkombination** der Vektoren $a_1, \dots, a_k$ mit Koeffizienten $\lambda_1, \dots, \lambda_k$.

Beachte, dass in dieser Definition jedes a_l ein Vektor ist! Im Gegensatz dazu hatten wir beim Einführen der Addition und des skalaren Vielfachen die einzelnen Komponenten des Vektors a mit $a_1, \dots, a_n$ bezeichnet. **Bei Formeln mit Vektoren ist stets ganz genau zu prüfen, welche Rolle ggf. vorhandene Indizes spielen!**

Matrizen

Die folgende Definition mag zunächst etwas sinnarm erscheinen. Es wird sich jedoch zeigen, dass Matrizen trotz ihres einfachen Aufbaus ein zentrales Werkzeug der Mathematik mit teils sehr tief liegender, vorerst in keiner Weise erkennbarer Bedeutung sind.

Merke

Eine **Matrix** ist ein rechteckiges Zahlenschema, in dem Zahlen in Zeilen und Spalten angeordnet sind. Die Menge aller Matrizen mit $m \in \mathbb{N}$ Zeilen und $n \in \mathbb{N}$ Spalten bezeichnen wir mit $\mathbb{R}^{m \times n}$ (sprich: “ $\mathbb{R}$ hoch m Kreuz n ”).

Matrizen als Ganzes werden üblicherweise mit Großbuchstaben bezeichnet. Die **Einträge** (oder auch: **Elemente**) einer Matrix werden meist mit dem zugehörigen Kleinbuchstabe und zwei unteren Indizes bezeichnet. Der erste Index gibt die Zeilennummer an, der zweite die Spaltennummer. Für eine Matrix $A \in \mathbb{R}^{m \times n}$ schreibt man also beispielsweise

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1,n-1} & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2,n-1} & a_{2n} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ a_{m-1,1} & a_{m-1,2} & \cdots & a_{m-1,n-1} & a_{m-1,n} \\ a_{m1} & a_{m2} & \cdots & a_{m,n-1} & a_{m,n} \end{bmatrix}. \quad (5.4)$$

Statt viereckiger Klammern [...] sind auch runde Klammern (...) üblich. Die beiden Indizes werden, sofern es der besseren Lesbarkeit dient, durch Komma getrennt; also $a_{i,j}$ statt a_{ij} .

Völlig analog zu Vektoren sind auch für Matrizen **Summen und skalare Vielfache** elementweise definiert. Der Begriff der Linearkombination steht entsprechend auch bei Matrizen zur Verfügung, spielt hier aber kaum eine Rolle.

Merke

Vektoren kann man als Matrix mit nur einer Spalte oder nur einer Zeile ansehen. Die Matrizen in $\mathbb{R}^{m \times 1}$ bezeichnet man daher auch als **Spaltenvektoren** und die in $\mathbb{R}^{1 \times n}$ als **Zeilenvektoren**.

Während man in der Mathematik Vektoren stets als Spaltenvektoren ansieht, wird bei der Arbeit mit dem Computer (je nach verwendeter Software) streng unterschieden zwischen Vektor (Liste von Zahlen) und Spaltenvektor (Rechteck von Zahlen mit nur einer Spalte).

```
import numpy as np

vektor = np.array([1, 4, 3])
spalte = np.array([[1], [4], [3]])
zeile = np.array([[1, 4, 3]])

print(vektor)
print(vektor.shape)
print()
print(spalte)
print(spalte.shape)
print()
print(zeile)
print(zeile.shape)
```

```
[1 4 3]
(3,)

[[1]
 [4]
 [3]]
(3, 1)

[[1 4 3]]
(1, 3)
```

Merke

Zu einer Matrix $A \in \mathbb{R}^{m \times n}$ bezeichnen wir mit $A^T \in \mathbb{R}^{n \times m}$ die Matrix, die durch Vertauschen von Zeilen und Spalten aus A entsteht:

$$A = \begin{bmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{bmatrix}, \quad A^T := \begin{bmatrix} a_{11} & \cdots & a_{1m} \\ \vdots & \ddots & \vdots \\ a_{n1} & \cdots & a_{nm} \end{bmatrix}. \quad (5.5)$$

Die Matrix A^T wird als **transponierte Matrix** oder **Transponierte** zu A bezeichnet.

Um die Transponierte zu einer Matrix zu bilden nimmt man also eine Zeile nach der anderen der ursprünglichen Matrix und schreibt sie als Spalten nebeneinander.

Das Transponieren eines Spaltenvektors liefert einen Zeilenvektor und umgekehrt. Möchte man einen (Spalten-)Vektor in Fließtext unterbringen, so schreibt man ihn gern als transponierten Zeilenvektor, also beispielsweise $[1 \ 4 \ 5]^T$ oder $[1, 4, 5]^T$.

Einige häufig auftretende Matrizen und Matrixtypen bekommen eigene Namen:

- Bei einer **Nullmatrix** sind alle Einträge Null.
- Eine **Diagonalmatrix** ist eine quadratische Matrix $A \in \mathbb{R}^{n \times n}$, bei der nur die Einträge $a_{11}, \dots, a_{nn}$ von Null verschieden sind (sie dürfen aber auch Null sein). Alle anderen Einträge sind Null. Die Einträge $a_{11}, \dots, a_{nn}$ werden auch als **Hauptdiagonale** von A bezeichnet.
- Eine **obere Dreiecksmatrix** ist eine Matrix, die unterhalb der Hauptdiagonale nur Nullen enthält.
- Eine **untere Dreiecksmatrix** ist eine Matrix, die oberhalb der Hauptdiagonale nur Nullen enthält.
- Eine **Einheitsmatrix** (bei gegebener Größe oft auch: die (!) Einheitsmatrix) ist eine Diagonalmatrix mit Einsen auf der Hauptdiagonale. Einheitsmatrizen werden oft mit I bezeichnet ohne weitere Erklärung des Symbols. Man rechnet leicht nach, dass

$$A I = A \quad \text{und} \quad I A = A \quad (5.6)$$

gilt (bei geeigneten Matrixgrößen).

- Eine **symmetrische Matrix** ist eine quadratische Matrix A , für die $A^T = A$ gilt.

Länge eines Vektors

Merke

Als **Länge** oder **Norm** eines Vektors $a \in \mathbb{R}^n$ bezeichnen wir den Wert

$$\|a\| := \sqrt{\sum_{k=1}^n a_k^2} \quad \left(= \sqrt{a_1^2 + \dots + a_n^2} \right). \quad (5.7)$$

Nebenbemerkung

Beachte, dass die Begriffe Länge und Norm bei uns völlig synonym verwendet werden. Allerdings kann man auch andere Abstandsbegriff ("Normen") einführen. In diesem Fall meint man mit "Länge" dann die oben eingeführte Norm, die auch als **euklidische Norm** bekannt ist.

Einige Eigenschaften der Norm:

- Es gilt stets

$$\|a\| \geq 0. \quad (5.8)$$

- Ein Vektor hat genau dann die Länge Null, wenn es der Nullvektor ist. In Formeln:

$$\|a\| = 0 \quad \Leftrightarrow \quad a = \begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix}. \quad (5.9)$$

- Für $\lambda \in \mathbb{R}$ und $a \in \mathbb{R}^n$ gilt

$$\|\lambda a\| = |\lambda| \|a\|. \quad (5.10)$$

- Es gilt die **Dreiecksungleichung**, d.h. für beliebige Vektoren $a, b \in \mathbb{R}^n$ gilt

$$\|a + b\| \leq \|a\| + \|b\|. \quad (5.11)$$

Die ersten beiden Eigenschaften folgen sofort aus der Definition der Norm. Der Nachweis der dritten Eigenschaft macht etwas mehr Arbeit, bringt aber keine neuen Erkenntnisse und unterbleibt deshalb an dieser Stelle.

Merke

Bildet man zu einem Vektor $a \in \mathbb{R}^n$ den Vektor

$$\tilde{a} := \frac{1}{\|a\|} a, \quad (5.12)$$

so sagt man, dass der Vektor a **normiert** wird. Der so entstehende Vektor $\tilde{a}$ zeigt in die gleiche Richtung wie a , hat aber stets die Länge 1:

$$\|\tilde{a}\| = \left\| \frac{1}{\|a\|} a \right\| = \left| \frac{1}{\|a\|} \right| \|a\| = 1. \quad (5.13)$$

Üblicherweise schreibt man auch gebrochene Faktoren nicht mit Bruchstrich unter einen Vektor, z.B.

$$\frac{1}{2} a \quad (\text{üblich}), \quad \frac{a}{2} \quad (\text{unüblich}). \quad (5.14)$$

Einzige Ausnahme bildet das Normieren:

$$\frac{1}{\|a\|} a \quad (\text{üblich}), \quad \frac{a}{\|a\|} \quad (\text{ebenfalls üblich}). \quad (5.15)$$

Skalarprodukt zweier Vektoren

Analog zur noch nicht erkennbaren Bedeutung von Matrizen, wird auch der tiefe Nutzen der folgenden, schon fast banalen Rechenoperation erst im Laufe der nächsten Unterkapitel deutlich werden.

Merke

Für zwei gleich lange Vektoren $a, b \in \mathbb{R}^n$ setzen wir

$$a \circ b := \sum_{k=1}^n a_k b_k \quad (= a_1 b_1 + a_2 b_2 + \cdots + a_n b_n) \quad (5.16)$$

und bezeichnen diesen Wert $a \circ b \in \mathbb{R}$ als **Skalarprodukt** der beiden Vektoren a und b . Gelegentlich wird auch die Bezeichnung “inneres Produkt” verwendet (vor allem im Englischen: “inner product”).

Beachte, dass wir das Symbol $\circ$ vorher schon für die Hintereinanderausführung von Abbildungen verwendet haben. Aufgrund der Vielzahl von Operationen und dem chronischen Mangel an Symbolen sind in der Mathematik viele Symbole doppelt und dreifach mit Bedeutung belegt. Die konkrete Bedeutung ergibt sich jeweils aus dem Kontext.

Später werden wir für Vektoren mit drei Komponenten noch eine andere Art von Produkt einführen, das sogenannte Kreuzprodukt oder äußere Produkt. Deshalb sollte immer

“Skalarprodukt” und nicht nur “Produkt” gesagt und geschrieben werden.

Aus der Definition des Skalarprodukts folgt sofort

$$a \circ a = \|a\|^2. \quad (5.17)$$

Weniger offensichtlich ist

$$|a \circ b| \leq \|a\| \|b\|, \quad (5.18)$$

die sogenannte **Cauchy-Schwarz-Ungleichung**, die wir hier nicht herleiten. Bei dieser gilt Gleichheit genau dann, wenn die Vektoren a und b Vielfache voneinander sind, also insbesondere bei $a = b$.

Weiter wichtige Eigenschaften, die sofort aus der Definition folgen und deren Kenntnis das Rechnen mit Skalarprodukten einfacher macht:

- Das Skalarprodukt mit dem Nullvektor ist stets Null:

$$0 \circ a = 0 \quad \text{und} \quad a \circ 0 = 0. \quad (5.19)$$

- Das Skalarprodukt ist **kommutativ**, d.h.

$$a \circ b = b \circ a. \quad (5.20)$$

- Es gilt das **Distributivgesetz**

$$(a + b) \circ c = (a \circ c) + (b \circ c) \quad \text{und} \quad a \circ (b + c) = (a \circ b) + (a \circ c). \quad (5.21)$$

- Für $\lambda \in \mathbb{R}$ gilt

$$(\lambda a) \circ b = \lambda (a \circ b) \quad \text{und} \quad a \circ (\lambda b) = \lambda (a \circ b). \quad (5.22)$$

Beachte, dass das von der Multiplikation bei Zahlen bekannte Assoziativgesetz $a(b c) = (a b) c$ hier keine Rolle spielt, da das Ergebnis des Skalarprodukts zweier Vektoren kein Vektor ist. Somit können keine Ketten von drei oder mehr Faktoren gebildet werden.

Produkt zweier Matrizen

Rein formal könnte man auch für Matrizen eine Art Skalarprodukt einführen (elementweise multiplizieren und alle Produkte addieren). Dies hat aber keinerlei Nutzen für uns und unterbleibt deshalb. Jedoch wird eine andere, etwas komplexere Form von Produkt von wesentlicher Bedeutung sein. In der folgenden Definition (und auch sonst) bezeichnen wir mit $a_{\bullet j}$ die j -te Spalte von A (Spaltenvektor) und mit $a_{i\bullet}$ die i -te Zeile von A (Zeilenvektor).

Merke

Für zwei Matrizen $A \in \mathbb{R}^{m \times p}$ und $B \in \mathbb{R}^{p \times n}$ setzen wir

$$AB := \begin{bmatrix} a_{1\bullet} \circ b_{\bullet 1} & \cdots & a_{1\bullet} \circ b_{\bullet n} \\ \vdots & \ddots & \vdots \\ a_{m\bullet} \circ b_{\bullet 1} & \cdots & a_{m\bullet} \circ b_{\bullet n} \end{bmatrix} \quad (5.23)$$

und bezeichnen die so entstehende Matrix $AB \in \mathbb{R}^{m \times n}$ als **Matrixprodukt** oder einfach als Produkt der Matrizen A und B .

Beachte, dass das Produkt nur für Matrizenpaare erklärt ist, bei denen die erste Matrix gerade so viele Spalten hat wie die zweite Zeilen besitzt. Für alle anderen Kombinationen von Matrizen ist nicht klar, was man unter ihrem Produkt verstehen soll. Die Situation ist vergleichbar mit der Division durch Null bei Zahlen: Man kann der Sache auch mit viel gutem Wille keinen Sinn beimessen; die Operation ist schlicht “nicht definiert”.

Beispiel

Die Multiplikation zweier Matrizen ist etwas fummelig, aber problemlos machbar. Man merke sich die Regel “Zeile mal Spalte” (im Sinne des Skalarprodukts). Diese Regel führt man für alle Zeilen der ersten und alle Spalten der zweiten Matrix aus. So erhält man beispielsweise

$$\begin{bmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \end{bmatrix} \begin{bmatrix} 3 & 7 & 0 \\ 5 & 3 & 0 \\ 1 & 2 & 0 \\ 2 & 5 & 0 \end{bmatrix} = \begin{bmatrix} 24 & 38 & 0 \\ 68 & 107 & 0 \end{bmatrix} \quad (5.24)$$

(Video 505, Notizen dazu)

Das Produkt zweier Matrizen kann natürlich auch mit dem Computer berechnet werden. Beachte, dass der übliche Multiplikationsoperator `*` das elementweise Produkt bildet (sofern die Matrizengrößen übereinstimmen). Für das Matrixprodukt ist `@` zu verwenden.

```
A = np.array([[1, 2, 3, 4], [5, 6, 7, 8]])
B = np.array([[3, 7, 0], [5, 3, 0], [1, 2, 0], [2, 5, 0]])
```

```
A @ B
```

```
array([[ 24,  39,   0],
       [ 68, 107,   0]])
```

Das Matrixprodukt eines Zeilenvektors mit einem Spaltenvektor liefert eine Matrix aus $\mathbb{R}^{1 \times 1}$, also praktisch nur eine Zahl. Entsprechend schreibt man das Skalarprodukt zweier

Vektoren gern als

$$a \circ b = a^T b \quad (5.25)$$

und sieht die rechte Seite (formal nicht ganz korrekt) als Zahl an. Beachte aber, dass bei umgekehrter Reihenfolge, also Spaltenvektor mal Zeilenvektor eine große quadratische Matrix entsteht!

Man kann sich überlegen, dass Matrizenmultiplikation **assoziativ** ist, d.h. bei Produkten mit drei oder mehr Faktoren ist es egal, in welcher Abfolge man die einzelnen Produkte bildet. Für Matrizen A, B, C passender Größe gilt also

$$(A B) C = A (B C), \quad (5.26)$$

sodass man auch einfach $A B C$ schreiben kann.

Allerdings ist Matrizenmultiplikation **nicht kommutativ**, d.h. man kann die Faktoren im Allgemeinen nicht vertauschen. Einerseits passen nach dem Vertausch der Faktoren die Matrixgrößen eventuell nicht mehr; andererseits kann sich das Ergebnis trotz passender Größen ändern.

Beispiel

Man rechnet leicht nach, dass

$$\begin{bmatrix} 1 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 1 \end{bmatrix} \neq \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & 0 \end{bmatrix} \quad (5.27)$$

gilt.

Wie bei Zahlen gilt auch bei Matrizen das **Distributivgesetz**, welches Addition und Multiplikation verbindet: Für Matrizen $A \in \mathbb{R}^{m \times p}$ und $B, C \in \mathbb{R}^{p \times n}$ gilt

$$A (B + C) = A B + A C. \quad (5.28)$$

In umgekehrter Reihenfolge gilt für $A, B \in \mathbb{R}^{m \times p}$ und $C \in \mathbb{R}^{p \times n}$ entsprechend

$$(A + B) C = A C + B C. \quad (5.29)$$

Beachte: Da Matrizenmultiplikation nicht kommutativ ist, folgt aus dem ersten Teil des Distributivgesetzes nicht automatisch der zweite!

Nebenbemerkung

Assoziativgesetz und Distributivgesetz bei Matrizen sind trotz ihres Namens keine "Gesetze", sondern durch logisches Schließen aus der Definition von Addition und Multiplikation herleitbare Aussagen. Das Herleiten dieser beiden Aussagen ist ohne große (geistige) Mühe möglich, aber recht länglich und unterbleibt deshalb an dieser Stelle.

Eine weitere nützliche Eigenschaft des Matrixprodukts, die sofort aus der Definition folgt, ist

$$(\lambda A) B = \lambda (A B) \quad \text{und} \quad A (\lambda B) = \lambda (A B), \quad (5.30)$$

wobei A und B Matrizen passender Größe sind und $\lambda \in \mathbb{R}$ eine beliebige Zahl ist.

Auch kann man zeigen, dass für das Zusammenspiel von Multiplizieren und Transponieren die Beziehung

$$(A B)^T = B^T A^T \quad (5.31)$$

gilt.

Hinweise zum Matrix-Vektor-Produkt

Sieht man Vektoren als Spaltenvektoren an, wie wir es stets tun, so kann man das Produkt aus Matrix und Vektor als Produkt zweier Matrizen berechnen. Das Produkt einer Matrix $A \in \mathbb{R}^{m \times n}$ mit einem Vektor $x \in \mathbb{R}^n$ liefert dann den Vektor $Ax \in \mathbb{R}^m$. Umgekehrt kann man für $x \in \mathbb{R}^m$ auch $x^T A \in \mathbb{R}^n$ bilden. Wir wollen hier einige Hinweise zur genaueren Struktur von Matrix-Vektor-Produkten sammeln, da solche Produkte sehr häufig auftauchen werden.

Schaut man sich die Arbeitsweise der Matrizenmultiplikation an, so sieht man für $A \in \mathbb{R}^{m \times n}$ und $x \in \mathbb{R}^n$, dass

$$Ax = x_1 a_{\bullet 1} + \cdots + x_n a_{\bullet n}. \quad (5.32)$$

Das Produkt liefert also gerade eine Linearkombination der Matrixspalten, wobei die Vektorkomponenten die Koeffizienten sind. Bei umgekehrter Reihenfolge (jetzt $x \in \mathbb{R}^m$!) gilt

$$x^T A = x_1 a_{1\bullet} + \cdots + x_m a_{m\bullet} \quad (5.33)$$

(Linearkombination der Zeilen).

Ebenfalls sehr einfach sieht man, dass die Spalten des Produkts zweier Matrizen $A \in \mathbb{R}^{m \times p}$ und $B \in \mathbb{R}^{p \times n}$ als Matrix-Vektor-Produkte geschrieben werden können:

$$AB = \begin{bmatrix} | & & | \\ Ab_{\bullet 1} & \cdots & Ab_{\bullet n} \\ | & & | \end{bmatrix}. \quad (5.34)$$

Dabei soll die Schreibweise mit $|$ andeuten, dass die Vektoren $Ab_{\bullet j}$ als Spalten der Matrix fungieren.

Das Berechnen von Matrix-Vektor-Produkten mit dem Computer erfolgt wie bei Produkten zweier Matrizen. Die explizite Umwandlung des Vektors in einen Spaltenvektor (Matrix mit einer Spalte) ist nicht nötig, da dies intern von NumPy erledigt wird. Wird der Vektor als linker Faktor verwendet, so wird er automatisch als Zeilenvektor interpretiert.

```
A = np.array([[1, 2], [3, 4], [5, 6]])
x = np.array([7, 8])
```

```
A @ x
```

```
array([23, 53, 83])
```

Determinante einer Matrix

Matrizen sieht man nur schlecht auf den ersten Blick wichtige Eigenschaften an. Deshalb wurden im Laufe der Zeit verschiedene Kennzahlen entwickelt, die relevante Informationen aus (beliebig großen) Matrizen in einer Zahl zusammenfassen. Führen hier zunächst nur die sogenannte Determinante einer Matrix ein; später folgt noch der Rang einer Matrix.

Wir werden Determinanten nur für Matrizen in $\mathbb{R}^{2 \times 2}$ und $\mathbb{R}^{3 \times 3}$ einführen. Für größere Matrizen sei bei Bedarf auf den Entwicklungssatz von Laplace verwiesen. Einige Hinweise zur Rechtfertigung dieses “sparsamen” Vorgehens:

- Für die praktische Arbeit (vor allem Koordinatentransformationen) spielen nur kleine Matrizen eine Rolle.
- Für größere Matrizen ist der Aufwand zum Berechnen von Determinanten so groß, dass selbst Computer schnell an ihre Leistungsgrenze kommen. Manuelles Rechnen scheidet somit aus.
- Das Berechnen von Determinanten ist ohne jede Denkarbeit möglich, also bestens für Computer geeignet.
- Die allgemeine Formel für Determinanten zu Matrizen beliebiger Größe benötigt viel Raum und viel Zeit für Erklärungen, die wir zielführender in andere Themen investieren.

Merke

Für eine Matrix $A \in \mathbb{R}^{2 \times 2}$ heißt die Zahl

$$\det A := a_{11} a_{22} - a_{12} a_{21} \quad (5.35)$$

Determinante der Matrix.

Für eine Matrix $A \in \mathbb{R}^{3 \times 3}$ heißt die Zahl

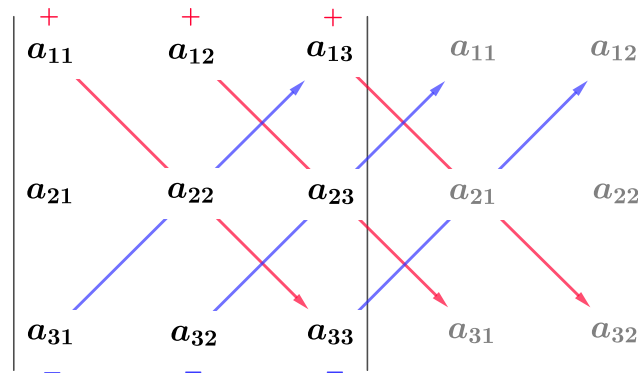
$$\begin{aligned} \det A := & a_{11} a_{22} a_{33} + a_{12} a_{23} a_{31} + a_{13} a_{21} a_{32} \\ & - a_{13} a_{22} a_{31} - a_{11} a_{23} a_{32} - a_{12} a_{21} a_{33} \end{aligned} \quad (5.36)$$

Determinante der Matrix.

Manchmal schreibt man auch $|A|$ für die Determinanten bzw.

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix}. \quad (5.37)$$

Für die Berechnung von Determinanten von Matrizen in $\mathbb{R}^{3 \times 3}$ gibt es sehr verschieden aussehende Formeln. Die hier verwendete ist als **Regel von Sarrus** bekannt.



Durch Linien verbundene Einträge werden multipliziert; die Produkte werden dann addiert (blau) bzw. subtrahiert (rot) (commons.wikimedia.org, Nutzer Kmhkmh, CC BY 4.0).

Anhand der Determinante kann man beispielsweise die Lösbarkeit von Gleichungssystemen erkennen (dazu später mehr). Ihren Hauptzweck erfüllen sie aber als Werkzeug für die platzsparende Beschreibung von Streckungs- und Spiegelungsvorgängen bei Koordinatentransformationen aller Art. Im Rahmen dieser Veranstaltung werden wir nicht tief genug in diese Thematik vorstoßen, um den Nutzen von Determinanten dort zu erkennen. Jedoch werden Determinanten in den weiterführenden Veranstaltungen im Laufe des Studiums eine Rolle spielen.

5.2 Vektorgeometrie

Vektoren und Matrizen als zunächst rein formale Werkzeuge sind sehr gut geeignet zur Beschreibung geometrischer Operationen (Verschiebungen, Drehungen, Streckungen, Projektionen,...). Beschäftigen uns hier zunächst mit Verschiebungen und einfachen Projektionen. Weitere Operationen folgen in späteren Unterkapiteln.

Verschiebungen

Möchte man eine gegebene Menge von Punkten im $\mathbb{R}^n$ in eine gegebene Richtung verschieben, so kann man diese Operation mittels Vektoraddition beschreiben. An dieser Stelle ist es üblich und sinnvoll zwischen Punkten und Vektoren zu unterscheiden. Ein Punkt $P \in \mathbb{R}^n$ gibt einen geometrischen Ort an, während ein Vektor $r \in \mathbb{R}^n$ eine Richtung angibt. Die Summe $P+r$ ist dann wieder als (verschobener) Punkt zu interpretieren. Einen formal sauberen Übergang zwischen Punkten und Vektoren schafft man sich über den Begriff des Ortsvektors:

Merke

Bezeichnen mit $O := (0, \dots, 0) \in \mathbb{R}^n$ (großes lateinisches O) den Punkt im Koordinatenursprung. Für einen Punkt $P := (p_1, \dots, p_n) \in \mathbb{R}^n$ heißt der Vektor

$$\overrightarrow{OP} := \begin{bmatrix} p_1 \\ \vdots \\ p_n \end{bmatrix} \in \mathbb{R}^n \quad (5.38)$$

Ortsvektor des Punktes.

Beachte hier, dass wir entgegen einer recht weit verbreiteten (Un-)Sitte “ $P := (p_1, \dots, p_n)$ ” statt “ $P(p_1, \dots, p_n)$ ” für die Zuordnung von Koordinaten zu einem Punkt schreiben. Letzteres ist zu leicht mit einer Funktionsauswertung zu verwechseln und verschlechtert damit die Lesbarkeit.

Zu jedem Punkt gibt es also einen Ortsvektor. Umgekehrt können wir jeden Vektor als Ortsvektor eines Punktes interpretieren, indem wir die Komponenten des Vektors als Koordinaten eines Punktes ansehen.

Mit Hilfe von Ortsvektoren können wir nun die Verschiebung von Punkten beschreiben:

Merke

Sind $P := (p_1, \dots, p_n)$ und $Q := (q_1, \dots, q_n)$ zwei Punkte im $\mathbb{R}^n$, so bezeichnen

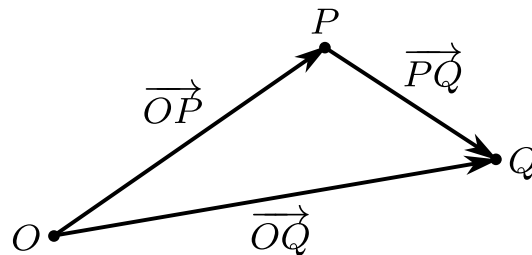
wir mit

$$\vec{PQ} := \begin{bmatrix} q_1 - p_1 \\ \vdots \\ q_n - p_n \end{bmatrix} \in \mathbb{R}^n \quad (5.39)$$

den Vektor, der P im Sinne von

$$\vec{OQ} = \vec{OP} + \vec{PQ} \quad (5.40)$$

nach Q verschiebt.



Die Verschiebung eines Punktes kann als Addition des Ortsvektors mit dem Richtungsvektor beschrieben werden.

Im Prinzip könnten wir für einen Punkt P und einen Vektor r nun die Operation

$$P + r := \vec{OP} + r \quad (5.41)$$

definieren. Üblicherweise belässt man es aber bei $\vec{OP} + r$, hat also ausschließlich Vektoren in Formeln stehen.

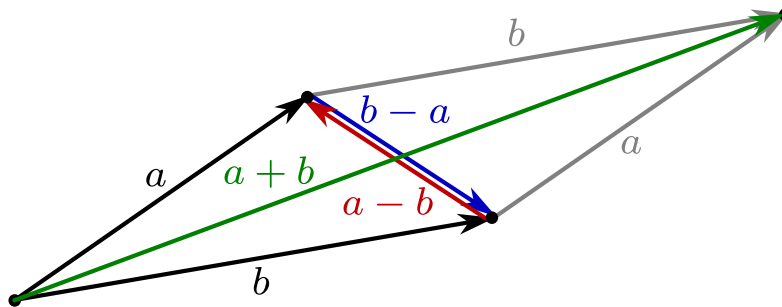
Nebenbemerkung (koordinatenfreie Geometrie)

Man kann die Arbeit mit Punkten und Vektoren in gewissem Maße auch “koordinatenfrei” betreiben. Punkte sind dann Orte in der Ebene (oder im Raum) ohne Bezug zu einem konkreten Koordinatensystem. Vektoren sind “Verschiebungen”, die man sich als lose in Ebene (oder Raum) herumliegende Pfeile vorstellt. Die “Addition” von Vektoren beschreibt man verbal als “Hintereinanderhängen” von Vektoren. Möchte man einen Punkt in eine gewisse Richtung verschieben, so wählt man sich einen Referenzpunkt O (mangels Koordinatensystem gibt es keinen Ursprung) und drückt die Verschiebung als Addition (d.h. “Hintereinanderhängen”) des Vektors, der von O zum Punkt zeigt, und des Richtungsvektors aus.

Diese Herangehensweise ist recht verlockend, da die meisten geometrischen Operationen (z.B. auch Drehungen) völlig unabhängig von Koordinatensystemen beschrieben werden können. Allerdings ist es schwierig bis unmöglich die benötigten

Begriffe wirklich sauber zu definieren. Auch ist dieser Ansatz in der Praxis ohne Relevanz, da stets ein kanonisches Koordinatensystem vorhanden ist, insbesondere bei der Arbeit mit dem Computer.

Angemerkt sei noch, dass man aus den bisherigen Überlegungen auch die geometrische Interpretation der Summe zweier Vektoren $a, b \in \mathbb{R}^n$ erhält. Die Summe entspricht der Diagonale des von den beiden Vektoren beschriebenen Parallelogramms, die am gemeinsamen Ausgangspunkt der beiden Vektoren beginnt. Die andere Diagonale ist dann gerade die Differenz.



Zwei Vektoren beschreiben ein Parallelogramm. Die beiden Diagonalen sind gerade Summe und Differenz der beiden Vektoren.

Abstand zweier Punkte

Mit dem bis hier eingeführten Formalismus der Vektorrechnung und den Betrachtungen zur Verschiebung von Punkten erhalten wir folgenden Zusammenhang:

Merke

Der **Abstand zweier Punkte** $P, Q \in \mathbb{R}^n$ ist gegeben durch die Länge des Vektors $\overrightarrow{PQ}$.

Mit Koordinaten $P = (p_1, \dots, p_n)$ und $Q = (q_1, \dots, q_n)$ gilt also

$$\text{Abstand von } P \text{ und } Q = \|\overrightarrow{PQ}\| = \sqrt{\sum_{k=1}^n (p_k - q_k)^2}. \quad (5.42)$$

Projektion eines Vektors auf einen anderen Vektor

Öfter als man denkt werden wir eine Operation benötigen, die als “Projektion eines Vektors auf einen anderen Vektor” bezeichnet wird. Insbesondere wird uns diese Operation sehr beim Wechsel zwischen verschiedenen Koordinatensystemen behilflich sein. Hier zunächst nur ein motivierendes Beispiel aus der Schulphysik.

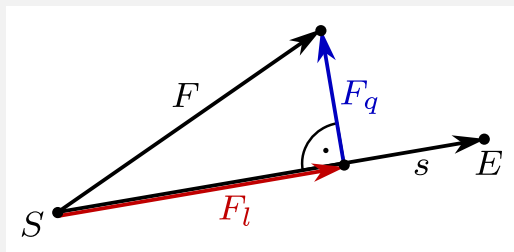
Beispiel (mechanische Arbeit)

Bewegt sich ein Objekt auf direktem Weg (also geradlinig) von einem Punkt zu einem anderen unter dem Einfluss von Kräften, so wird (mechanische) Arbeit verrichtet. Denken wir beispielsweise an einen Eisenbahnwagen auf einem geraden Streckenabschnitt, so wirken unter anderem Reibungskräfte und vielleicht eine Hangabtriebskraft (bei Steigung/Gefälle). Der Antrieb (Lok) muss zur Bewegung des Wagens Arbeit verrichten, wobei das physikalische Gesetz

$$\text{Arbeit} = \text{Weg} \cdot \text{Kraft in Wegrichtung}$$

gilt. Wirken nun auch Kräfte in andere als die Bewegungsrichtung (z.B. Wind), so hilft diese einfache Regel nicht mehr weiter. Wir müssen daher wahlweise auf ein allgemeiner formuliertes Gesetz ausweichen oder den Anteil der Kraft, der in Wegrichtung wirkt, aus der Gesamtkraft berechnen. Beide Varianten sind machbar; interessieren uns hier aber zunächst für die letztere (das allgemeinere Gesetz erhalten wir später (fast) geschenkt).

Beschreiben den zurückgelegten Weg zwischen Startpunkt S und Endpunkt E als Vektor $s := \overrightarrow{SE}$. Die wirkende Kraft können wir ebenfalls als Vektor F beschreiben. Die Richtung des Vektors entspricht der Richtung der Kraft und die Länge des Vektors entspricht dem Betrag der Kraft. Wollen den Kraftvektor in eine Summe $F_l + F_q$ zweier Vektoren zerlegen, wobei F_l längs des Weges und F_q quer zum Weg zeigt. Die Länge von F_l ist dann gerade der Betrag der Kraft in Wegrichtung, also die Größe, die wir für die obige Formel aus der Schulphysik benötigen.



Beachte, dass die Kräfte F und F_l an jedem Punkt des Weges auf das Objekt wirken, hier aber nur am Startpunkt S gezeichnet sind. Entsprechend wirkt auch die Kraft F_q an jedem Punkt des Weges, wurde hier aber an einer für die Visua-

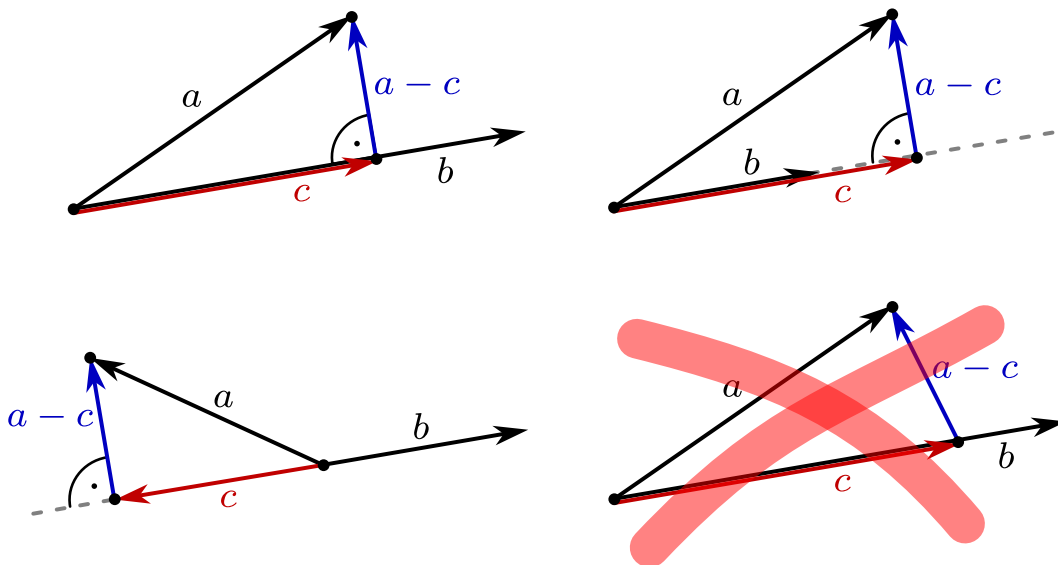
lisierung der Zerlegung von F günstigen Stelle eingetragen.

Merke

Seien $a, b \in \mathbb{R}^n$ zwei Vektoren. Ein Vektor $c \in \mathbb{R}^n$ heißt **Projektion von a auf b** , wenn die folgenden beiden Bedingungen erfüllt sind:

- c ist ein Vielfaches von b ,
- $a - c$ ist der kürzeste Vektor bezüglich aller c , die die erste Bedingung erfüllen.

Anschaulich besagt die erste Bedingung, dass c in die Richtung von b (oder genau entgegen) zeigt. Die zweite Bedingung fordert indirekt, dass $a - c$ senkrecht zu b verläuft, wie wir später sehen werden, wenn wir Winkel zwischen Vektoren betrachten.



Die Länge des Vektors b spielt keine Rolle; nur seine Richtung ist relevant für die Projektion von a auf b , wobei die Projektion c auch genau die entgegengesetzte Richtung zu b haben kann. Ist $a - c$ nicht der kürzest mögliche Vektor, so spricht man nicht von einer Projektion.

Es zeigt sich, dass die Projektion eines Vektors auf einen anderen sehr leicht mit Hilfe des Skalarprodukts der beiden gegebenen Vektoren berechnet werden kann. Wir müssen also nicht mit Winkeln hantieren, sondern nur mit einfachen Summen und Produkten von Zahlen. Dies ist der große Wert des so banal erscheinenden Skalarprodukts! Es spielt in gewisser Weise die Rolle von Winkeln zwischen Vektoren, ohne dass wir uns mit Winkeln und den dann kaum vermeidbaren Winkelfunktionen (Sinus, Cosinus,...) herumschlagen müssen. Im nächsten Abschnitt gehen wir genauer auf den Zusammenhang zwischen

Skalarprodukt und Winkeln ein.

Merke

Seien $a, b \in \mathbb{R}^n$ zwei Vektoren und sei $c \in \mathbb{R}^n$ die Projektion von a auf b . Dann gilt

$$c = \left(a \circ \frac{b}{\|b\|} \right) \frac{b}{\|b\|} \quad (5.43)$$

(Video 510, Notizen dazu).

Die Projektion von a auf b kann man also in zwei Schritten berechnen:

1. b normieren.
2. $(a \circ \tilde{b}) \tilde{b}$ berechnen, wobei $\tilde{b}$ der normierte Vektor b ist.

Beispiel (allgemeine Formel für die mechanische Arbeit)

Betrachten nochmal die aus der Schulphysik bekannte Berechnung der mechanischen Arbeit. Sei also s wieder der Wegvektor und F der Kraftvektor. Dann ist die Länge des Weges $\|s\|$ und der Betrag der in Wegerichtung wirkenden Kraft ist

$$\left\| \left(F \circ \frac{s}{\|s\|} \right) \frac{s}{\|s\|} \right\| = \left| F \circ \frac{s}{\|s\|} \right| \left\| \frac{s}{\|s\|} \right\| = \left| F \circ \frac{s}{\|s\|} \right|. \quad (5.44)$$

Die aus der Schulphysik bekannte Formel zur Berechnung der mechanischen Arbeit liefert somit

$$W = \left| F \circ \frac{s}{\|s\|} \right| \|s\| = |F \circ s|. \quad (5.45)$$

Für Kräfte, die nicht in Wegerichtung wirken, gilt also im Sinne des Skalarprodukts ebenfalls die bekannte Formel “Arbeit ist Kraft mal Weg”.

Umgekehrt erhalten wir aus der allgemeinen Formel natürlich die einfache Formel zurück: Zeigt der Kraftvektor F in Wegerichtung, so gilt

$$F = \|F\| \frac{s}{\|s\|} \quad (5.46)$$

und Einsetzen in die allgemeine Formel liefert

$$\begin{aligned} W = |F \circ s| &= \left| \left(\|F\| \frac{s}{\|s\|} \right) \circ s \right| = \left| \frac{\|F\|}{\|s\|} (s \circ s) \right| \\ &= \left| \frac{\|F\|}{\|s\|} \|s\|^2 \right| = \|F\| \|s\| = \|F\| \|s\|. \end{aligned} \quad (5.47)$$

Winkel zwischen zwei Vektoren

Wollen den in $\mathbb{R}^2$ und $\mathbb{R}^3$ intuitiv klaren Begriff des Winkels zwischen zwei Vektoren formalisieren (also ohne Bezug zu Intuition und Anschauung erklären), damit er auch in $\mathbb{R}^n$ mit $n > 3$ zur Verfügung steht.

Merke

Als **Winkel zwischen zwei Vektoren** $a, b \in \mathbb{R}^n \setminus \{0\}$ bezeichnen wir die Zahl

$$\arccos \left(\frac{a}{\|a\|} \circ \frac{b}{\|b\|} \right). \quad (5.48)$$

Dazu einige Hinweise:

- Ist einer der beiden Vektoren der Nullvektor, so ist der Begriff “Winkel” nicht sinnvoll erklärbar. Entsprechend wird in der Definition $a \neq 0$ und $b \neq 0$ vorausgesetzt (0 ist hier der Nullvektor in $\mathbb{R}^n$).
- Der Definitionsbereich von $\arccos$ ist das Intervall $[-1, 1]$. Aus der Cauchy-Schwarz-Ungleichung (vgl. Section 5.1) erhalten wir

$$\left| \frac{a}{\|a\|} \circ \frac{b}{\|b\|} \right| \leq \left\| \frac{a}{\|a\|} \right\| \left\| \frac{b}{\|b\|} \right\| = 1, \quad (5.49)$$

sodass der Ausdruck in der Definition des Winkels stets wohldefiniert ist.

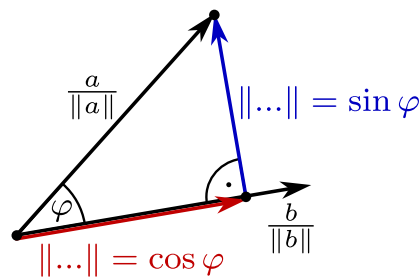
- Der Winkel zwischen zwei Vektoren liegt stets im Intervall $[0, \pi]$.
- Bezeichnen wir mit φ den Winkel zwischen den beiden Vektoren a und b , so folgt aus der Definition sofort

$$a \circ b = \|a\| \|b\| \cos \varphi. \quad (5.50)$$

Dass die obige Definition des Begriffs “Winkel” mit der üblichen Verwendung des Begriffs in Einklang steht, sieht man aus der Berechnung der Projektion von Vektoren. Bei normierten Vektoren entspricht die Länge des projizierten Vektors (laut Definition des Cosinus) gerade dem Betrag des Cosinus des Winkels zwischen den beiden Vektoren, sofern der Winkel im Bereich $[0, \frac{\pi}{2}]$ liegt. Für Winkel in $[\frac{\pi}{2}, \pi]$ ist die negative Länge des projizierten Vektors gerade der Cosinus des Winkels.

Aus der Beziehung zwischen Skalarprodukt und Winkel erhält man nützliche Informationen über das Vorzeichen des Skalarprodukts:

- Für Winkel in $[0, \frac{\pi}{2})$ (spitzer Winkel) gilt $a \circ b > 0$.
- Für senkrechte Vektoren gilt $a \circ b = 0$.



Die Länge des projizierten Vektors entspricht gerade dem Cosinus des Winkels.
Die Länge des Differenzvektors zwischen Ausgangsvektor und projiziertem Vektor entspricht gerade dem Sinus des Winkels.

- Für Winkel in $(\frac{\pi}{2}, \pi]$ (stumpfer Winkel) gilt $a \circ b < 0$.

Somit kann nun auch ganz ohne den Begriff des Winkels von Orthogonalität sprechen. Tatsächlich werden wir, wo immer möglich, die Arbeit mit Winkeln zwischen Vektoren meiden und stattdessen das viel bequemere Skalarprodukt nutzen.

Merke

Zwei Vektoren $a, b \in \mathbb{R}^n$ heißen **orthogonal** zueinander, wenn $a \circ b = 0$ gilt.

Beachte, dass nach dieser Definition der Nullvektor orthogonal zu jedem anderen Vektor ist.

5.3 Lineare Abbildungen

Begriff und Beispiele

Viele geometrische Operationen wie Drehen, Spiegeln und Strecken (nicht jedoch Verschieben!) haben eine gemeinsame Struktur, deren Kenntnis sich insbesondere für die rechnerische Umsetzung dieser Operationen vorteilhaft nutzen lässt. Auch die Umsetzung geometrischer Operationen am Computer, z.B. in CAD-Systemen, basiert ganz wesentlich auf dieser gemeinsamen Struktur.

Merke

Eine Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ heißt **linear**, wenn folgende Bedingungen erfüllt sind:

- $f(a + b) = f(a) + f(b)$ für alle $a, b \in \mathbb{R}^n$,
- $f(\lambda a) = \lambda f(a)$ für alle $a \in \mathbb{R}^n$ und alle $\lambda \in \mathbb{R}$.

Die beiden Bedingungen besagen, dass es bei linearen Abbildungen egal ist, ob man erst die Rechenoperation bei Vektoren anwendet (also Addition und skalaras Vielfaches) und dann die Abbildung, oder ob man erst die Abbildung auf die Ausgangsvektoren anwendet und anschließend die notwendigen Rechenoperationen durchführt. Man sagt auch, dass lineare Abbildungen die Struktur des $\mathbb{R}^n$ erhalten bzw. diese in den $\mathbb{R}^m$ übertragen.

Beispiel (Drehung)

Die Abbildung, die jeden Punkt $(x, y) \in \mathbb{R}^2$ um den Winkel $\alpha \in \mathbb{R}$ um den Koordinatenursprung dreht, ist offensichtlich linear, denn:

- Die Summe zweier Vektoren bildet die Diagonale des von den beiden Vektoren aufgespannten Parallelogramms. Beim Drehen um den Ursprung verändert das Parallelogramm seine Form nicht.
- Ob ein Vektor erst gestreckt wird (skalaras Vielfaches) und dann gedreht oder umgekehrt, führt offensichtlich zum selben Ergebnis.

Das gleiche gilt für die Drehung um eine durch den Ursprung verlaufende Achse im $\mathbb{R}^3$.

Beispiel (Spiegelung)

Die Abbildung, die jeden Punkt $(x, y) \in \mathbb{R}^2$ an einer durch den Koordinatenursprung verlaufenden Gerade spiegelt, ist offensichtlich linear. Ebenso ist die Punktspiegelung $(x, y) \mapsto (-x, -y)$ linear.

Gleiches gilt im $\mathbb{R}^3$ für Spiegelung an einer durch den Ursprung verlaufenden Ebene und für Punktspiegelung am Ursprung.

Beispiel (Streckung)

Die Abbildung $a \mapsto \lambda a$ für $a \in \mathbb{R}^n$ und festes $\lambda \in \mathbb{R}$ ist offensichtlich linear.

Merke

Die Hintereinanderausführung linearer Abbildungen ist linear, denn für zwei lineare Abbildungen $f : \mathbb{R}^n \rightarrow \mathbb{R}^p$ und $g : \mathbb{R}^p \rightarrow \mathbb{R}^m$ gelten

$$\begin{aligned}(g \circ f)(a + b) &= g(f(a + b)) \\ &= g(f(a) + f(b)) \\ &= g(f(a)) + g(f(b)) \\ &= (g \circ f)(a) + (g \circ f)(b)\end{aligned}\tag{5.51}$$

und

$$\begin{aligned}(g \circ f)(\lambda a) &= g(f(\lambda a)) \\ &= g(\lambda f(a)) \\ &= \lambda g(f(a)) \\ &= \lambda (g \circ f)(a).\end{aligned}\tag{5.52}$$

Matrixdarstellung

Jede lineare Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ kann als Matrix interpretiert werden. Um dies zu sehen, führen wir zunächst die Vektoren

$$e_k := \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}\tag{5.53}$$

$k = 1, \dots, n$ ein. Diese sind bis auf eine Komponenten überall Null und der Vektor e_k hat als k -te Komponente eine Eins. Mit diesen Vektoren erhalten wir nun für jedes $x \in \mathbb{R}^n$:

$$f(x) = f\left(\sum_{k=1}^n x_k e_k\right) = \sum_{k=1}^n f(x_k e_k) = \sum_{k=1}^n x_k f(e_k).\tag{5.54}$$

Beachte hier, dass die x_k die Komponenten des Vektors x sind (also Zahlen), die e_k jedoch selbst Vektoren sind. Schreiben wir die Vektoren $f(e_k) \in \mathbb{R}^m$ als Spalten in eine Matrix, so erhalten wir folgenden nützlichen Zusammenhang (vgl. dazu auch Section 5.1):

Merke

Die zu einer linearen Abbildung $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ definierte Matrix

$$A := \begin{bmatrix} | & & | \\ f(e_1) & \cdots & f(e_n) \\ | & & | \end{bmatrix} \in \mathbb{R}^{m \times n} \quad (5.55)$$

heißt **Matrixdarstellung** von f . Es gilt stets

$$f(x) = Ax. \quad (5.56)$$

Jede lineare Abbildung kann also als Matrix-Vektor-Produkt geschrieben werden. Umgekehrt ist für eine beliebige Matrix $A \in \mathbb{R}^{m \times n}$ die Abbildung $x \mapsto Ax$ stets linear. Der Wert von Matrizen liegt ganz wesentlich in der Beschreibung beliebiger linearer Abbildung durch Zahlen. Und diese Zahlen kann man sehr einfach und explizit erhalten indem man die betrachtete lineare Abbildung auf die Vektoren $e_1, \dots, e_n$ anwendet. Die Darstellung als Zahlen (Matrix) ermöglicht ein sehr bequemes Arbeiten mit linearen Abbildungen am Computer. Matrix-Vektor-Produkte sind ohne Weiteres am Computer umsetzbar.

Beispiel (Matrixdarstellung zur Drehung in der Ebene)

Ist $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Drehung in der Ebene um einen Winkel $\alpha \in \mathbb{R}$ um den Koordinatenursprung, so überlegt man sich leicht aus den Definitionen von Sinus und Cosinus, dass

$$f(e_1) = \begin{bmatrix} \cos \alpha \\ \sin \alpha \end{bmatrix}, \quad f(e_2) = \begin{bmatrix} -\sin \alpha \\ \cos \alpha \end{bmatrix} \quad (5.57)$$

gilt. Somit

$$f(x, y) = \begin{bmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}. \quad (5.58)$$

Eine Menge von Punkten, z.B. Eckpunkten eines Vielecks, lässt sich damit mühelos drehen:

```
import numpy as np
import matplotlib.pyplot as plt

points = np.array([ # ein Punkt pro Zeile
```

```

    [3, 0],
    [3, 1],
    [1, 1],
    [1, 0],
    [3, 0]
])

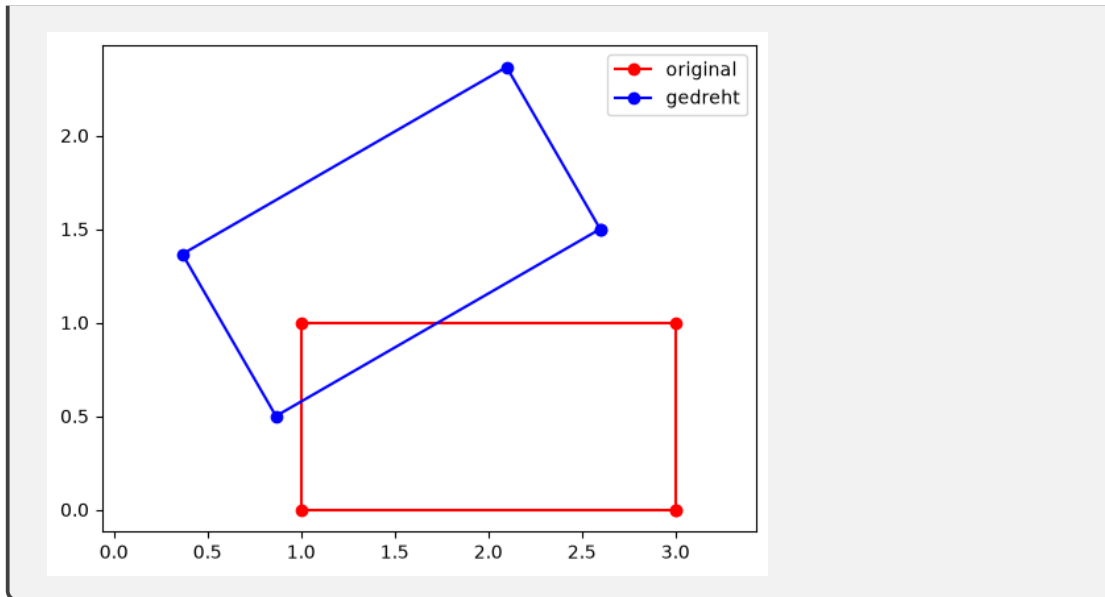
angle = 30 * np.pi / 180 # Drehwinkel

A = np.array([ # Drehmatrix
               [np.cos(angle), -np.sin(angle)],
               [np.sin(angle),  np.cos(angle)]
])

rotated = (A @ points.T).T
# points.T enthält pro Spalte einen Punkt. Das Produkt enthält
→ dann in jeder
# Spalte den entsprechenden gedrehten Punkt. Nochmaliges
→ Transponieren liefert
# ein Array mit den gedrehten Punkten in den Zeilen.

fig, ax = plt.subplots()
ax.plot(points[:, 0], points[:, 1], 'o -r', label='original')
ax.plot(rotated[:, 0], rotated[:, 1], 'o -b', label='gedreht')
ax.legend()
ax.axis('equal') # gleiche Achseneinteilung (keine Verzerrung)
plt.show()

```



Beispiel (Matrixdarstellung zur Spiegelung an der x-y-Ebene)

Ist $f : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ die Spiegelung an der x - y -Ebene, so gilt offensichtlich

$$f(e_1) = e_1, \quad f(e_2) = e_2, \quad f(e_3) = -e_3. \quad (5.59)$$

Somit

$$f(x, y, z) = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix}. \quad (5.60)$$

Im Laufe der Veranstaltung werden uns weitere, auch deutlich komplexere, lineare Abbildungen begegnen.

Komposition vs. Matrixmultiplikation

Sind $f : \mathbb{R}^n \rightarrow \mathbb{R}^p$ und $g : \mathbb{R}^p \rightarrow \mathbb{R}^m$ lineare Abbildung und sind $A \in \mathbb{R}^{n \times p}$ und $B \in \mathbb{R}^{p \times m}$ die zugehörigen Matrixdarstellungen, so gilt offensichtlich

$$(g \circ f)(x) = g(f(x)) = g(Ax) = B(Ax) = (BA)x. \quad (5.61)$$

In Worten: **Die Matrixdarstellung der Komposition (Hintereinanderausführung) zweier linearer Abbildungen ist gerade das Produkt der Matrixdarstellungen der beiden Abbildungen.**

Damit sind Sinn und Zweck der zunächst etwas seltsam anmutenden Definition des Produkts zweier Matrizen geklärt: Das Matrixprodukt wurde gerade so gewählt, dass das

Multiplizieren von Matrizen dem Hintereinanderausführen der entsprechenden linearen Abbildungen entspricht.

Diese Erkenntnis hilft einerseits beim Finden von Matrixdarstellungen für komplexere Abbildungen (erst Spiegeln, dann Drehen, dann Strecken usw.); andererseits kann man damit auch recht einfach zu diversen Standardresultaten kommen, die sonst nur mühsam herleitbar sind:

Beispiel (Formel für den doppelten Winkel)

Sei $\alpha \in \mathbb{R}$ ein beliebiger Winkel, sei $A \in \mathbb{R}^{2 \times 2}$ die Matrixdarstellung der Drehung um diesen Winkel und sei $B \in \mathbb{R}^{2 \times 2}$ die Matrixdarstellung der Drehung um den Winkel 2α , d.h.

$$A = \begin{bmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{bmatrix} \quad \text{und} \quad B = \begin{bmatrix} \cos(2\alpha) & -\sin(2\alpha) \\ \sin(2\alpha) & \cos(2\alpha) \end{bmatrix}. \quad (5.62)$$

Die Drehung um den doppelten Winkel kann als Komposition von zwei einfachen Drehungen angesehen werden. Somit gilt $B = A A$, also

$$\begin{bmatrix} \cos(2\alpha) & -\sin(2\alpha) \\ \sin(2\alpha) & \cos(2\alpha) \end{bmatrix} = \begin{bmatrix} (\cos \alpha)^2 - (\sin \alpha)^2 & -2(\cos \alpha) \sin \alpha \\ 2(\cos \alpha) \sin \alpha & (\cos \alpha)^2 - (\sin \alpha)^2 \end{bmatrix}. \quad (5.63)$$

Wir erhalten damit die bekannten Formeln

$$\cos(2\alpha) = (\cos \alpha)^2 - (\sin \alpha)^2 \quad (5.64)$$

und

$$\sin(2\alpha) = 2(\cos \alpha) \sin \alpha. \quad (5.65)$$

für Cosinus und Sinus des doppelten Winkels.

5.4 Lineare Unterräume und Untermannigfaltigkeiten

Die hier behandelten Begriffe mögen etwas abstrakt erscheinen, sind uns aber schon in der Schulmathematik in Form von Geraden und Ebenen begegnet. Auch beim Lösen von linearen Gleichungssystemen spielen diese Begriffe eine Rolle.

Unterräume

Ein Unterraum ist eine Teilmenge von $\mathbb{R}^n$, in der wir die üblichen Vektoroperationen (Addition und skalare Vielfache) ausführen können ohne dabei diese Teilmenge zu verlassen. Genauer:

Merke

Eine Teilmenge $U \subseteq \mathbb{R}^n$ heißt **linearer Unterraum**, wenn gilt:

- $a + b \in U$ für alle $a, b \in U$,
- $\lambda a \in U$ für alle $a \in U$ und alle $\lambda \in \mathbb{R}$.

Offensichtlich ist $\{0\}$ (Menge, die nur den Nullvektor enthält) ein Unterraum und auch $\mathbb{R}^n$ selbst ist ein Unterraum.

Beispiel (Gerade in der Ebene)

Die Menge $U := \{(x, y) \in \mathbb{R}^2 : y = 3x\}$ ("Gerade mit Anstieg 3") ist ein Unterraum von $\mathbb{R}^2$, denn:

- Für $(x, y) \in U$ und $(\tilde{x}, \tilde{y}) \in U$ gilt

$$y + \tilde{y} = 3x + 3\tilde{x} = 3(x + \tilde{x}), \quad (5.66)$$

also $(x + \tilde{x}, y + \tilde{y}) \in U$.

- Für $(x, y) \in U$ und $\lambda \in \mathbb{R}$ gilt

$$\lambda y = \lambda(3x) = 3(\lambda x), \quad (5.67)$$

also $(\lambda x, \lambda y) \in U$.

Der Wert von Unterräumen besteht darin, dass wir in ihnen mit Vektoren arbeiten können als würde es aushalb keine weiteren Vektoren geben. Arbeitet man beispielsweise im $\mathbb{R}^3$ nur mit Vektoren (oder Punkten), die alle in einer Ebene liegen, so unterscheidet sich das Vorgehen praktisch nicht vom Arbeiten in $\mathbb{R}^2$. Wie genau dieser Übergang zwischen Unterräumen und den entsprechenden "ganzen" Räumen vollzogen werden kann,

ist Thema der nächsten Abschnitte.

Lineare Unabhängigkeit

Wenn wir zwei Vektoren $a, b \in \mathbb{R}^n \setminus \{0\}$ betrachten, dann können diese Vielfache voneinander sein (also in die gleiche Richtung bzw. Gegenrichtung zeigen) oder nicht. In letzterem Fall sagt man, dass die beiden Vektoren linear unabhängig sind. Nimmt man einen dritten Vektor $c \in \mathbb{R}^n \setminus \{0\}$ hinzu, so liegt dieser wahlweise in der von den ersten beiden Vektoren aufgespannten Ebene oder nicht. In letzterem Fall sagt man wieder, dass die drei Vektoren linear unabhängig sind. In ersterem Fall können wir c als Linearkombination

$$c = \alpha a + \beta b \quad (5.68)$$

mit $\alpha, \beta \in \mathbb{R}$ darstellen. Anders ausgedrückt: Wir finden Zahlen $\alpha, \beta, \gamma \in \mathbb{R}$, sodass

$$\alpha a + \beta b + \gamma c = 0 \quad (5.69)$$

gilt, obwohl die drei Zahlen nicht alle Null sind (speziell ist hier $\gamma = -1$). Allgemeiner definieren wir lineare Unabhängigkeit wie folgt:

Merke

Eine Menge $\{a_1, \dots, a_r\} \subseteq \mathbb{R}^n$ von r Vektoren in $\mathbb{R}^n$ heißt **linear unabhängig**, wenn für alle Zahlen $\lambda_1, \dots, \lambda_r \in \mathbb{R}$ gilt:

$$\sum_{k=1}^r \lambda_k a_k = 0 \quad \Rightarrow \quad \lambda_1 = \dots = \lambda_r = 0. \quad (5.70)$$

Andernfalls heißt die Menge linear abhängig.

Enthält die Menge $\{a_1, \dots, a_r\}$ den Nullvektor, so ist sie offensichtlich stets linear abhängig (das entsprechende λ_k kann beliebig gewählt werden).

Ist die Menge $\{a_1, \dots, a_r\}$ linear abhängig, so gibt es ein $l \in \{1, \dots, r\}$, sodass $\lambda_l \neq 0$, aber trotzdem die Linearkombination $\sum_{k=1}^r \lambda_k a_k$ den Nullvektor liefert. Umgestellt gilt

$$a_l = \sum_{\substack{k=1 \\ k \neq l}}^r \frac{\lambda_k}{\lambda_l} a_k, \quad (5.71)$$

d.h. a_l lässt sich als Linearkombination der anderen Vektoren darstellen.

Umgekehrt gilt: Lässt sich ein Vektor als Linearkombination anderer Vektoren darstellen, so ist die entsprechende Menge von Vektoren linear abhängig (wähle $\lambda_l = -1$ wie oben für $\mathbb{R}^3$ beschrieben).

Basis und Dimension

Man kann nun der Frage nachgehen, wie viele linear unabhängige Vektoren man eigentlich finden kann. In der Ebene stellt man schnell fest, dass höchstens zwei Vektoren linear unabhängig sein können. Mengen mit drei oder mehr Vektoren sind stets linear abhängig. Entsprechend findet man im Raum nur höchstens drei linear unabhängige Vektoren. Wir können lineare Unabhängigkeit also als Grundlage für die präzise Formulierung des intuitiv klaren Begriffs “Dimension” verwenden.

Merke

Die Dimension eines linearen Unterraums U des $\mathbb{R}^n$ ist die maximale Anzahl linear unabhängiger Vektoren in U . Die Dimension von U wird mit dem Symbol $\dim U$ abgekürzt.

Der triviale Unterraum $\{0\}$ hat Dimension Null: Es gibt nur einen Vektor, den Nullvektor, und Mengen, die den Nullvektor enthalten, sind stets linear abhängig.

Die Dimension von $\mathbb{R}^n$ ist n : Die Menge $\{e_1, \dots, e_n\}$ (wobei e_k in der k -ten Komponente eins ist und sonst null) ist linear unabhängig und jeder weitere Vektor $x \in \mathbb{R}^n$ kann als Linearkombination $x = \sum_{k=1}^n x_k e_k$ geschrieben werden, führt also zu linearer Abhängigkeit.

Offensichtlich gilt für jeden Unterraum U von $\mathbb{R}^n$, dass $\dim U \in \{0, 1, \dots, n\}$.

Merke

Sei $U \subseteq \mathbb{R}^n$ ein Unterraum der Dimension $r := \dim U$. Jede Menge $\{b_1, \dots, b_r\} \subseteq U$ von r linear unabhängigen Vektoren aus U heißt **Basis** von U .

Beispielsweise bilden die Vektoren $e_1, \dots, e_n$ eine Basis in $\mathbb{R}^n$. Es gibt jedoch auch noch andere Basen. Mehr dazu später.

Koordinaten bezüglich einer Basis

Merke

Sei $U \subseteq \mathbb{R}^n$ und sei $\{b_1, \dots, b_r\} \subseteq U$ eine Basis in U . Dann existieren zu jedem $x \in U$ eindeutig bestimmte Zahlen $\lambda_1, \dots, \lambda_r \in \mathbb{R}$, sodass

$$x = \sum_{k=1}^r \lambda_k b_k \quad (5.72)$$

gilt. Diese Zahlen heißen **Koordinaten von x bezüglich der Basis $\{b_1, \dots, b_r\}$** .

Existenz und Eindeutigkeit der Koordinaten bezüglich einer Basis können leicht nachgewiesen werden. Da diese Überlegungen für uns aber keine neue Erkenntnisse bringen, verzichten wir darauf.

Offensichtlich sind die Komponenten $x_1, \dots, x_n$ eines Vektors $x \in \mathbb{R}^n$ gerade die Koordinaten bezüglich der Basis $\{e_1, \dots, e_n\}$, denn

$$x = \sum_{k=1}^n x_k e_k. \quad (5.73)$$

Beispiel (alternative Basis in der Ebene)

In $\mathbb{R}^2$ ist durch die beiden Vektoren

$$b_1 := \begin{bmatrix} 1 \\ 1 \end{bmatrix}, \quad b_2 := \begin{bmatrix} -1 \\ 1 \end{bmatrix} \quad (5.74)$$

eine Basis gegeben. Die Koordinaten des Vektors $\begin{bmatrix} 0 & 1 \end{bmatrix}^T$ bezüglich dieser Basis sind $\frac{1}{2}$ und $\frac{1}{2}$, denn

$$\begin{bmatrix} 0 \\ 1 \end{bmatrix} = \frac{1}{2} \begin{bmatrix} 1 \\ 1 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} -1 \\ 1 \end{bmatrix}. \quad (5.75)$$

Für einen beliebigen Vektor $x \in \mathbb{R}^2$ erhält man bezüglich der Basis $\{b_1, b_2\}$ die Koordinaten

$$\frac{x_1 + x_2}{2} \quad \text{und} \quad \frac{x_2 - x_1}{2}, \quad (5.76)$$

denn

$$\begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \frac{x_1 + x_2}{2} \begin{bmatrix} 1 \\ 1 \end{bmatrix} + \frac{x_2 - x_1}{2} \begin{bmatrix} -1 \\ 1 \end{bmatrix}. \quad (5.77)$$

Wie man systematisch die Koordinaten eines Vektors bezüglich einer beliebigen gegebenen Matrix berechnet, klären wir später.

Beispiel (Basis in einem nichttrivialen Unterraum)

Betrachten in $\mathbb{R}^2$ den Unterraum

$$U := \{x \in \mathbb{R}^2 : x_2 = 3x_1\}. \quad (5.78)$$

Dies ist ein eindimensionaler Unterraum und $\{b_1\}$ mit

$$b_1 := \begin{bmatrix} 1 \\ 3 \end{bmatrix} \quad (5.79)$$

ist eine Basis darin.

Wir können nun jeden Vektor $x \in U$ durch seine Koordinate (genau eine Zahl!) bezüglich der Basis $\{b_1\}$ darstellen. Diese Koordinate ist gerade x_1 , denn

$$\begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = x_1 \begin{bmatrix} 1 \\ 3 \end{bmatrix} \quad (5.80)$$

(beachte, dass $x_2 = 3x_1$ für jedes $x \in U$ gilt).

Solange wir nur Betrachtungen in U anstellen, können wir statt Vektoren mit zwei Komponenten also einfach mit einkomponentigen Vektoren (d.h. Zahlen) arbeiten. Die Arbeit im eindimensionalen Unterraum U unterscheidet sich also nicht von der Arbeit in $\mathbb{R}$.

Die Tatsache, dass man bei der Arbeit in Unterräumen mit (teils deutlich) “kürzeren” Vektoren arbeiten kann, macht man sich in den Datenwissenschaften (Data Science) und im maschinellen Lernen zu nutze. Zu analysierende Daten “leben” oft in einem sehr hochdimensionalen Grundraum (Bildraten mit 1000x1000 Pixeln z.B. in $\mathbb{R}^{1000000}$), während die eigentlich relevanten Datenpunkte sich in einem vergleichsweise niedrigdimensionalen Unterraum befinden. Durch Wahl geeigneter Basen können Ressourcen (Speicher, Rechenzeit) gespart werden.

Aufspannung

Oben hatten wir uns mit Basen in einem gegebenen Unterraum beschäftigt. Umkehrt steht man manchmal vor der Frage, was denn eigentlich der kleinste Unterraum ist, der eine gegebene Menge von Vektoren enthält.

Merke

Sei $\{b_1, \dots, b_r\} \subseteq \mathbb{R}^n$ eine Menge von r linear unabhängigen Vektoren in $\mathbb{R}^n$. Die Menge

$$\text{span}\{b_1, \dots, b_r\} := \left\{ \sum_{k=1}^r \lambda_k b_k : \lambda \in \mathbb{R}^r \right\} \quad (5.81)$$

heißt **Aufspannung** von $\{b_1, \dots, b_r\} \subseteq \mathbb{R}^n$.

Folgende Eigenschaften von $\text{span}\{b_1, \dots, b_r\}$ sieht man leicht ein:

- Die Aufspannung ist ein linearer Unterraum.
- Die Menge $\{b_1, \dots, b_r\}$ ist eine Basis in ihrer Aufspannung.
- Die Aufspannung ist der kleinste Unterraum, der $\{b_1, \dots, b_r\}$ enthält, d.h. jeder Unterraum, der $\{b_1, \dots, b_r\}$ enthält, enthält automatisch auch $\text{span}\{b_1, \dots, b_r\}$.

Nebenbemerkung

Beachte, dass manche Quellen auch Aufspannungen für Mengen von Vektoren betrachten, die nicht linear unabhängig sind. Dann ist diese Menge keine Basis mehr in ihrer Aufspannung. Die anderen beiden Eigenschaften gelten aber weiterhin.

In $\mathbb{R}^3$ ist die Aufspannung von zwei linear unabhängigen Vektoren eine Ebene. In $\mathbb{R}^2$ und $\mathbb{R}^3$ ist die Aufspannung eines Vektors (außer Nullvektor) eine Gerade. Zwei linear unabhängige Vektoren in $\mathbb{R}^2$ spannen den gesamten $\mathbb{R}^2$ auf, usw.

Orthonormalbasen

Meist ist man an besonders schönen, d.h. leicht zu handhabenden Basen, interessiert. Führen hier eine Klasse besonders schöner Basen und einige Begriffe aus deren Umfeld ein.

Merke

Ein Menge $\{a_1, \dots, a_r\} \subseteq \mathbb{R}^n$ von r linear unabhängigen Vektoren heißt **Orthogonalsystem**, wenn

$$a_k \circ a_l = 0 \quad \text{für } k \neq l. \quad (5.82)$$

gilt. Eine solche Menge heißt **Orthonormalsystem** (ONS), wenn zusätzlich

$$\|a_k\| = 1 \quad \text{für } k = 1, \dots, r \quad (5.83)$$

gilt.

Ist ein ONS eine Basis für einen Unterraum $U \subseteq \mathbb{R}^n$, so bezeichnen wir dieses ONS als **Orthonormalbasis** (ONB) für U .

Die Bedingung $a_k \circ a_l = 0$ besagt, dass die Basisvektoren senkrecht zueinander stehen. Zusammen mit der Normierung $\|a_k\| = 1$ kann man die definierende Eigenschaft eines ONS auch als

$$a_k \circ a_l = \begin{cases} 0, & \text{falls } k \neq l, \\ 1, & \text{sonst} \end{cases} \quad (5.84)$$

ausdrücken.

Offensichtlich ist die Standardbasis $e_1, \dots, e_n$ eine ONB.

Untermannigfaltigkeiten

Unterräume enthalten stets den Koordinatenursprung. Für mehr Flexibilität (beim Modellieren und beim Formulieren in der Sprache der Mathematik) führt man noch "verschobene" Unterräume ein:

Merke

Eine **lineare Untermannigfaltigkeit** ist eine Menge der Form

$$x + U := \{x + u : u \in U\} \subseteq \mathbb{R}^n, \quad (5.85)$$

wobei $x \in \mathbb{R}^n$ ein fester Vektor und $U \subseteq \mathbb{R}^n$ ein linearer Unterraum ist.

Jede Wahl eines Vektors x und eines Unterraums U liefert also eine Mannigfaltigkeit. Allerdings können verschiedenen Kombination zur gleichen Untermannigfaltigkeit führen. Insbesondere erhält man bei $x \in U$ wieder U ; jeder Unterraum ist also selbst eine Untermannigfaltigkeit.

Wählt man zwei Vektoren a, b aus einer Untermannigfaltigkeit $x + U$, so liegt deren Differenz $a - b$ stets in U . (Begründung: Es existieren $u, v \in U$ mit $a = x + u$ und $b = x + v$. Somit $a - b = x + u - (x + v) = u - v$. Da skalare Vielfache und Summen von Elementen eines Unterraums wieder im Unterraum liegen, folgt $a - b \in U$.) In diesem Sinne kann man aus einer Untermannigfaltigkeit auch ohne Kenntnis des Verschiebungsvektors x den zugrundeliegenden Unterraum zurückerhalten.

Merke

Als **Dimension** einer Untermannigfaltigkeit bezeichnen wir die Dimension des zugrundeliegenden Unterraums.

Geraden

Hatten schon gelegentlich Geraden als Unterräume des $\mathbb{R}^2$ betrachtet. Natürlich besitzen auch $\mathbb{R}^3$ und höherdimensionale Räume eindimensionale Unterräume. Diese werden allerdings nur in $\mathbb{R}^2$ und $\mathbb{R}^3$ als Geraden bezeichnet. Eindimensionale Untermannigfaltigkeiten werden in $\mathbb{R}^2$ und $\mathbb{R}^3$ ebenfalls als Geraden bezeichnet.

Eine Gerade ist durch einen Aufpunkt (bzw. dessen Ortsvektor) und einen Richtungsvektor bestimmt. Allgemein können wir jede eindimensionale Untermannigfaltigkeit in der Form

$$\{a + tb : t \in \mathbb{R}\} \quad (5.86)$$

schreiben, wobei $a, b \in \mathbb{R}^n$ fest sind. Handelt es sich um einen Unterraum, so ist $\{b\}$ eine Basis darin und t ist gerade die Koordinate des entsprechenden Vektors bezüglich dieser Basis.

Ebenen

Ebenen sind zweidimensionale Untermannigfaltigkeiten in $\mathbb{R}^3$. Gelegentlich sieht man den Begriff "Hyperebene". Dabei handelt es sich nicht (!) um eine zweidimensionale

Untermannigfaltigkeit in $\mathbb{R}^n$, sondern um eine Untermannigfaltigkeit mit Dimension $n - 1$.

Eine Ebene ist durch einen Aufpunkt (bzw. dessen Ortsvektor) und zwei linear unabhängige (!) Richtungsvektoren bestimmt. Allgemein können wir jede zweidimensionale Untermannigfaltigkeit in der Form

$$\{a + s b + t c : s, t \in \mathbb{R}\} \tag{5.87}$$

schreiben, wobei $a, b, c \in \mathbb{R}^n$ fest sind. Handelt es sich um einen Unterraum, so ist $\{b, c\}$ eine Basis darin und s, t sind gerade die Koordinaten des entsprechenden Vektors bezüglich dieser Basis.

5.5 Lineare Gleichungssysteme

Lineare Gleichungssysteme sollten im Wesentlichen aus der Schulmathematik bekannt sein. Hier konzentrieren wir uns auf die Struktur der Lösungsmengen und auf Aspekte beim Lösen mit dem Computer.

Da lineare Gleichungssysteme eine der Aufgabenstellungen sind, die besonders gut mit dem Computer behandelt werden können, bemüht man sich in allen Teilgebieten von Naturwissenschaften, Technik und Wirtschaftswissenschaften zu lösenden praktische Probleme in Form linearer Gleichungssysteme zu formulieren. Zusätzlich tauchen lineare Gleichungssysteme auch zahlreich als Teilproblem in komplexeren Modellen und zugehörigen Lösungsverfahren auf.

Matrix-Vektor-Schreibweise

Ein lineares Gleichungssystem (kurz: LGS) mit m Gleichungen und n Unbekannten hat die Form

$$\begin{aligned} a_{11}x_1 + \cdots + a_{1n}x_n &= b_1 \\ &\vdots \\ a_{m1}x_1 + \cdots + a_{mn}x_n &= b_m. \end{aligned} \quad (5.88)$$

Die a_{ij} und die b_i sind gegeben. Die x_j sind gesucht.

Setzt man

$$A := \begin{bmatrix} a_{11} & \cdots & a_{1n} \\ \vdots & \ddots & \vdots \\ a_{m1} & \cdots & a_{mn} \end{bmatrix}, \quad b := \begin{bmatrix} b_1 \\ \vdots \\ b_m \end{bmatrix}, \quad x := \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}, \quad (5.89)$$

so kann man das LGS deutlich kompakter und handhabbarer mit einem Matrix-Vektor-Produkt als

$$Ax = b \quad (5.90)$$

schreiben. Die Matrix A wird als **Systemmatrix** bezeichnet, der gegebene Vektor b als **rechte Seite**.

Struktur der Lösungsmenge

Anzahl der Lösungen LGS können keine, genau eine oder unendlich viele Lösungen besitzen. Ein einfaches Beispiel ohne Lösung ist

$$\begin{bmatrix} 1 & 0 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 1 \\ 2 \end{bmatrix}. \quad (5.91)$$

Das LGS

$$\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 1 \\ 2 \end{bmatrix}. \quad (5.92)$$

besitzt hingegen genau eine Lösung.

Die Lösungsmenge eines LGS mit unendlich vielen Lösungen besitzt eine sehr konkrete und nützliche Struktur. Zur weiteren Untersuchung unterscheiden wir zwei Typen von LGS:

Merke

Ein LGS heißt **homogen**, wenn die rechte Seite der Nullvektor ist. Andernfalls heißt das LGS **inhomogen**.

Homogene LGS Die Lösungsmenge eines homogenen LGS mit n Unbekannten ist offensichtlich ein linearer Unterraum des $\mathbb{R}^n$, denn die Summe zweier Lösungen ist wieder eine Lösung und alle Vielfachen von Lösungen sind Lösungen. In Formeln:

$$Ax = 0, A\tilde{x} = 0 \Rightarrow A(x + \tilde{x}) = Ax + A\tilde{x} = 0 + 0 = 0 \quad (5.93)$$

bzw.

$$Ax = 0, \lambda \in \mathbb{R} \Rightarrow A(\lambda x) = \lambda Ax = \lambda \cdot 0 = 0. \quad (5.94)$$

Somit kann es beispielsweise auch kein homogenes LGS gegeben, welches genau zwei Lösungen besitzt, da dann automatisch die Summe der beiden Lösungen eine dritte Lösung liefern würde.

Beachte, dass der Nullvektor eine Lösung jedes homogenen LGS ist. Falls dies die einzige Lösung ist, ist die Lösungsmenge gerade der triviale Unterraum $\{0\}$. Das andere Extrem ist ein homogenes LGS mit Nullmatrix (nur Nullen) als Systemmatrix. Dann ist die Lösungsmenge ganz $\mathbb{R}^n$.

Inhomogene LGS Bei einem inhomogenen LGS ist die Lösungsmenge stets eine Untermannigfaltigkeit oder die leere Menge. Hat ein inhomogenes LGS eine Lösung $x^* \in \mathbb{R}^n$ (und ggf. noch mehr), so können wir das LGS als

$$Ax = Ax^* \quad (5.95)$$

schreiben. Setzen wir nun $\tilde{x} := x - x^*$, so ist dies äquivalent zu dem homogenen LGS

$$A\tilde{x} = 0 \quad (5.96)$$

(alles nach links bringen und A ausklammern). Dessen Lösungsmenge ist ein Unterraum und aus der Beziehung $x = x^* + \tilde{x}$ sehen wir, dass die Lösungsmenge des ursprünglichen inhomogenen LGS eine Untermannigfaltigkeit ist. **Die Lösungsmenge eines inhomogenen LGS ist stets eine verschobene Version der Lösungsmenge des homogene LGS mit identischer Systemmatrix!**

Geometrische Interpretation Visualisiert man die Lösungsmenge eines LGS mit zwei oder drei Unbekannten, so erhält man wahlweise die leere Menge, einen Punkt, eine Gerade, eine Ebene oder den ganzen Raum. Entsprechend eng verknüpft sind Geraden und Ebenen mit linearen Gleichungssystemen.

Im $\mathbb{R}^2$ korrespondieren Geraden mit den Lösungsmengen von LGS mit zwei Unbekannten und einer Gleichung:

$$a_{11} x_1 + a_{12} x_2 = b_1. \quad (5.97)$$

Dabei dürfen a_{11} und a_{12} nicht gleichzeitig Null sein (warum?).

Im $\mathbb{R}^3$ korrespondieren Ebenen mit den Lösungsmengen von LGS mit drei Unbekannten und einer Gleichung:

$$a_{11} x_1 + a_{12} x_2 + a_{13} x_3 = b_1. \quad (5.98)$$

Dabei dürfen wieder a_{11}, a_{12}, a_{13} nicht gleichzeitig Null sein.

Nebenbemerkung (Normalenvektor)

Liegt $(x_1^*, x_2^*, x_3^*) \in \mathbb{R}^3$ in der betrachteten Ebene, d.h. es gilt

$$a_{11} x_1^* + a_{12} x_2^* + a_{13} x_3^* = b_1, \quad (5.99)$$

so folgt für jeden anderen Punkt $(x_1, x_2, x_3) \in \mathbb{R}^3$ in der Ebene

$$a_{11} (x_1 - x_1^*) + a_{12} (x_2 - x_2^*) + a_{13} (x_3 - x_3^*) = 0. \quad (5.100)$$

Dies kann mittels Skalarprodukt als

$$\begin{bmatrix} a_{11} \\ a_{12} \\ a_{13} \end{bmatrix} \circ \begin{bmatrix} x_1 - x_1^* \\ x_2 - x_2^* \\ x_3 - x_3^* \end{bmatrix} = 0 \quad (5.101)$$

ausgedrückt werden. Der Vektor $[a_{11} \ a_{12} \ a_{13}]^T$ ist also senkrecht zu jedem “in der Ebene liegenden” Vektor (genauer: senkrecht zu jedem zur Ebene parallel verlaufenden Vektor). Vektoren mit dieser Eigenschaft werden als **Normalenvektor** der Ebene bezeichnet. Man kann zeigen, dass alle Normalenvektoren einer Ebene in $\mathbb{R}^3$ Vielfache voneinander sind (anschaulich ist das klar).

Die Beschreibung einer Ebene als Lösungsmenge der Gleichung $a_{11} x_1 + a_{12} x_2 + a_{13} x_3 = b_1$ wird auch als *Normalform* der Ebenen bezeichnet. Hat der Vektor $[a_{11} \ a_{12} \ a_{13}]^T$ die Länge 1, so spricht man auch von der **Hesse’schen Normalform**. Dieser Begriff sollte aus der Schulmathematik bekannt sein.

Geraden im $\mathbb{R}^3$ entsprechen den Lösungsmengen von LGS mit drei Unbekannten und zwei Gleichungen:

$$\begin{aligned} a_{11} x_1 + a_{12} x_2 + a_{13} x_3 &= b_1 \\ a_{21} x_1 + a_{22} x_2 + a_{23} x_3 &= b_2. \end{aligned} \quad (5.102)$$

Die Lösungsmenge dieses LGS ist der Durchschnitt der Lösungsmenge der ersten Gleichung und der Lösungsmenge der zweiten Gleichung, also die Schnittgerade der beiden durch die beiden Gleichungen beschriebenen Ebenen. Sind diese Ebenen parallel oder identisch, so erhält man keine Gerade als Lösungsmenge des LGS.

Völlig analog kann man sich überlegen, dass ein Punkt in $\mathbb{R}^2$ als Schnitt zweier (nicht paralleler oder identischer) Geraden entsteht, oder, alternativ, als Lösung eines LGS mit zwei Unbekannten und zwei Gleichungen.

Manuelles Lösen

LGS löst man praktisch immer mit dem Computer (siehe unten). Dennoch sollen hier kurz zwei Ideen skizziert werden, wie man ein (kleines) LGS manuell lösen kann. Das manuelle Vorgehen ist insofern relevant als dass der Computer im Wesentlichen genauso vorgeht (aber dabei viel schneller ist). Auch begegnen uns hier einige neue Begriffe, die später in anderem Kontext benötigt werden. Praktisch relevante LGS besitzen heute tausende Unbekannte und Gleichungen, sodass manuelles Lösen nicht möglich ist.

Gauß-Algorithmus Die Grundidee des als Gauß-Algorithmus bekannten Vorgehens ist das systematische Vereinfachen des Gleichungssystems ohne dabei die Lösungsmenge zu verändern. Folgende Operationen sind dabei hilfreich und ändern die Lösungsmenge nicht:

- Vertauschen der Reihenfolge der Gleichungen,
- Multiplizieren einer Gleichung mit einem Faktor (außer Null),
- ein Vielfaches einer Gleichung zu einer anderen Gleichung hinzuaddieren,
- Gleichungen streichen, bei denen alle Koeffizienten und die rechte Seite Null sind.

Mit diesen Operationen kann man jedes LGS auf **Stufenform** bringen, d.h. die Anzahl der Unbekannten mit Nicht-Null-Koeffizient wird von Gleichung zu Gleichung kleiner. Dabei richtet man den Ablauf so ein, dass die Unbekannten in der natürlichen Reihenfolge wegfallen. So taucht x_1 dann nur noch in der ersten (oder keiner) Gleichung auf; x_2 erscheint höchstens in den ersten beiden Gleichungen usw. Welche Operation wann und wie anzuwenden ist um eine Stufenform des LGS zu erhalten, folgt einem einfachen Schema, welches man ohne nennenswerte Denkarbeit abarbeiten kann (für solche Aufgaben wurde der Computer erfunden...).

Ist das LGS in Stufenform, so kann man mit der letzten Gleichung beginnend die Lösung Gleichung für Gleichung leicht zusammensetzen. Dieser Vorgang wird als **Rückwärts-einsetzen** bezeichnet und folgt ebenfalls wieder einem klaren Schema ohne Notwendigkeit des Nachdenkens.

Da wir den Gauß-Algorithmus dem Computer überlassen, verzichten wir hier auf Details und Sonderfälle und belassen es bei einigen Beispielen:

Beispiel (Gauß-Algorithmus bei genau einer Lösung)

Das LGS

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ 2x_1 - x_2 + 9x_3 &= 6 \\ 2x_1 + 2x_2 - x_3 &= 1 \end{aligned} \quad (5.103)$$

ist äquivalent zu folgendem LGS in Stufenform:

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ -2x_2 + 8x_3 &= 6 \\ 2x_3 &= 4. \end{aligned} \quad (5.104)$$

Rückwärtseinsetzen liefert

$$x = \begin{bmatrix} -\frac{7}{2} \\ 5 \\ 2 \end{bmatrix} \quad (5.105)$$

(Video 530, Notizen dazu).

Beispiel (Gauß-Algorithmus bei Nichtlösbarkeit)

Das LGS

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ 3x_1 - \frac{3}{2}x_2 + 9x_3 &= 6 \\ 2x_1 + 2x_2 - \frac{3}{2}x_3 &= 1 \end{aligned} \quad (5.106)$$

ist äquivalent zu folgendem LGS in Stufenform:

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ x_2 - \frac{5}{2}x_3 &= 1 \\ 0 &= 9. \end{aligned} \quad (5.107)$$

Dieses LGS besitzt aufgrund der letzten Gleichung offensichtlich keine Lösung (Video 535, Notizen dazu).

Beispiel (Gauß-Algorithmus bei unendlich vielen Lösungen)

Das LGS

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ -6x_1 + 3x_2 + 9x_3 &= 6 \\ 2x_1 + 2x_2 + 3x_3 &= 1 \end{aligned} \quad (5.108)$$

ist äquivalent zu folgendem LGS in Stufenform:

$$\begin{aligned} 2x_1 + x_2 + x_3 &= 0 \\ 6x_2 + 12x_3 &= 6. \end{aligned} \quad (5.109)$$

Dieses LGS besitzt die Lösungsmenge

$$\left\{ \begin{bmatrix} -\frac{1}{2} \\ 1 \\ 0 \end{bmatrix} + t \begin{bmatrix} \frac{1}{2} \\ -2 \\ 1 \end{bmatrix} : t \in \mathbb{R} \right\} \quad (5.110)$$

(Video 540, Notizen dazu).

Die Dimension der Lösungsmenge eines LGS kann man gut an der Systemmatrix des LGS in Stufenform ablesen. Dazu folgender Begriff:

Merke

Der **Rang** einer Matrix $\mathbb{R}^{m \times n}$ ist die maximale Anzahl linear unabhängiger Zeilen von A . Der Rang von A wird mit $\text{rg} A$ abgekürzt.

Man kann zeigen, dass der Rang gleichzeitig auch der maximalen Anzahl linear unabhängiger Spalten in A entspricht. Für $A \in \mathbb{R}^{m \times n}$ gilt also $0 \leq \text{rg} A \leq \min\{m, n\}$.

Zur Berechnung des Rangs einer Matrix kann der Gauß-Algorithmus verwendet werden, denn der Rang ist gerade die Anzahl an Nicht-Null-Zeilen der Stufenform der Matrix. Beachte, dass je nach gewähltem Vorgehen unterschiedliche Stufenformen entstehen können. Alle werden aber die gleiche Anzahl an Nicht-Null-Zeilen haben.

Die Dimension der Lösungsmenge eines homogenen (!) LGS ist gerade $m - \text{rg} A$, wobei m die Anzahl der Gleichungen ist. Bei einem inhomogenen LGS mit $\text{rg} A < m$ (also Null-Zeilen in der Systemmatrix der Stufenform) besitzt das LGS entweder keine Lösung oder die Lösungsmenge hat die Dimension $m - \text{rg} A$.

Nebenbemerkung

Der Rang einer Matrix spielt auch abseits des Lösen von LGS gelegentlich eine Rolle. Grund dafür ist der folgende Zusammenhang: Der Rang von $A \in \mathbb{R}^{m \times n}$ ist identisch mit der Dimension des **Bildraums** der Abbildung $x \mapsto Ax$, also mit

der Dimension des Unterraums

$$\{Ax : x \in \mathbb{R}^n\} \quad (5.111)$$

von $\mathbb{R}^m$.

Die folgenden beiden Begriffe sind intuitiv klar, bedürfen aber einer Diskussion im Kontext des Gauß-Algorithmus:

Merke

Ein LGS mit m Gleichungen und n Unbekannten heißt

- **unterbestimmt**, wenn $m < n$,
- **überbestimmt**, wenn $m > n$.

Man kann nun etwas unbedacht Aussagen formulieren wie “unterbestimmte LGS haben unendlich viele Lösungen” oder “überbestimmte LGS sind nicht lösbar”. Dies ist aber nicht immer korrekt. Einerseits können unterbestimmte LGS trotzdem sich widersprechende Gleichungen enthalten, also keine Lösung besitzen. Andererseits können überbestimmte LGS redundante Gleichungen enthalten, sodass dennoch Lösungen existieren. Das tatsächliche Lösungsverhalten eines LGS erkennt man erst an der Stufenform!

Inverse Matrix Für LGS mit quadratischer Systemmatrix $A \in \mathbb{R}^{n \times n}$ kann man, sofern das LGS eindeutig lösbar ist, die Lösung als Matrix-Vektor-Produkt mit einer geeigneten Matrix angeben. Dazu folgender Begriff:

Merke

Sei $A \in \mathbb{R}^{n \times n}$ eine quadratische Matrix und sei $I \in \mathbb{R}^{n \times n}$ die Einheitsmatrix. Gibt es eine quadratische Matrix $B \in \mathbb{R}^{n \times n}$, sodass

$$AB = I \quad \text{und} \quad BA = I, \quad (5.112)$$

gilt, so heißt A **invertierbar**. Die Matrix B heißt dann **Inverse** zu A oder **inverse Matrix** zu A . Das übliche Symbol für die Inverse ist

$$A^{-1} := B. \quad (5.113)$$

Man kann zeigen, dass eine Matrix $A \in \mathbb{R}^{n \times n}$ genau dann invertierbar ist, wenn $\text{rg } A = n$ gilt. In diesem Fall kann man das LGS

$$Ax = b \quad (5.114)$$

auf beiden Seiten von links mit der Inversen multiplizieren:

$$A^{-1} A x = A^{-1} b. \quad (5.115)$$

Wegen $A^{-1} A = I$ und $I x = x$ erhalten wir mit

$$x = A^{-1} b \quad (5.116)$$

eine direkte Berechnungsvorschrift für die Lösung des LGS. Es bleibt somit nur die Frage zu klären, wie man zu A^{-1} kommt.

Das manuelle Aufstellen der Inversen ist sehr mühsam. Man kann das Problem auf das Lösen von n LGS zurückführen, die alle die gleiche Systemmatrix (nämlich A) besitzen; nur die rechten Seiten unterscheiden sich (das sind gerade die n Spalten von I). Hintergrund ist der Zusammenhang

$$A B = \left[\begin{array}{c|ccc|c} & & & & \\ \hline & & & & \\ A b_{\bullet 1} & & \cdots & & A b_{\bullet n} \\ & & & & \\ \hline & & & & \end{array} \right] \quad (5.117)$$

für Matrix-Vektor-Produkte (vgl. Section 5.1). Wenn nun $A B = I$ gelten soll, so müssen also die LGS

$$A b_{\bullet 1} = e_1, \quad \dots, \quad A b_{\bullet n} = e_n \quad (5.118)$$

erfüllt sein, wobei $e_1, \dots, e_n$ die Spalten von I sind (e_k hat eine Eins als k -te Komponente, sonst Nullen). Diese LGS können mit dem Gauß-Algorithmus gelöst werden.

Aufgrund des Rechenaufwands ist das Lösen von linearen Gleichungssystemen mittels inverser Matrix nur sinnvoll, wenn das LGS sehr oft für verschiedene rechte Seiten (und gleiche Systemmatrix) gelöst werden muss. Dieser Fall tritt in der Praxis gar nicht so selten auf. Aus Sicht des Anwenders (d.h. (Vermessungs-)Ingenieurs oder (Geo-)Informatikers) sind solche Fragen aber wenig relevant, da diese Themen innerhalb von auf effizientes Lösen von LGS spezialisierten Software-Bibliotheken behandelt und entschieden werden und damit dem Auge des Anwenders verborgen bleiben.

Einige nützliche Zusammenhänge beim Rechnen mit Inversen seien noch erwähnt. Diese lassen sich leicht herleiten. Für quadratische Matrizen $A, B \in \mathbb{R}^{n \times n}$ gilt:

- Sind A und B invertierbar, so ist auch $A B$ invertierbar und es gilt

$$(A B)^{-1} = B^{-1} A^{-1}. \quad (5.119)$$

- Ist A invertierbar, so ist auch A^{-1} invertierbar und es gilt

$$(A^{-1})^{-1} = A. \quad (5.120)$$

- Ist A invertierbar, so ist auch A^T invertierbar und es gilt

$$(A^T)^{-1} = (A^{-1})^T. \quad (5.121)$$

Manchmal schreibt man auch A^{-T} für die Inverse der Transponierten (bzw. für die Transponierte der Inversen).

- A ist genau dann invertierbar, wenn $\det A \neq 0$ gilt.

Symbolisches Lösen

Mit SymPy können lineare Gleichungssysteme wahlweise direkt als Liste von Gleichungen gelöst werden oder in der Matrix-Vektor-Form. Als Liste von Gleichungen:

```
import sympy

x, y, z = sympy.symbols('x y z')

sympy.linsolve([
    2 * x + y + z - 0,
    -6 * x + 3 * y + 9 * z - 6,
    2 * x + 2 * y + 3 * z - 1
], [x, y, z])
```

```
{(z/2 - 1/2, 1 - 2*z, z)}
```

Beachte, dass `sympy.linsolve` auf der rechten Seite der Gleichungen von Nullen ausgeht, wir also die rechte Seite des zu lösenden LGS zunächst auf die linke Seite bringen müssen.

Das Lösen in Matrix-Vektor-Schreibweise funktioniert nur bei invertierbarer Systemmatrix:

```
A = sympy.Matrix([
    [2, 1, 1],
    [2, -1, 9],
    [2, 2, -1]
])
b = sympy.Matrix([0, 6, 1])

A.solve(b)
```

```
Matrix([
[ -7/2],
[  5],
[  2]])
```

Da symbolisch (nicht numerisch) gelöst wird, können im LGS auch allgemeine Parameter enthalten sein:

```
c = sympy.symbols('c')
A = sympy.Matrix([
    [2, 1, c],
    [2, -1, 9],
    [2, 2, -1]
])
b = sympy.Matrix([0, 6, 1])

A.solve(b)
```

```
Matrix([
[(13*c + 15)/(6*c - 14)],
[ (-5*c - 15)/(3*c - 7)],
[      -32/(12*c - 28)]])
```

Die Inverse einer Matrix kann ebenfalls mit SymPy berechnet werden:

```
A = sympy.Matrix([
    [2, 1, 1],
    [2, -1, 9],
    [2, 2, -1]
])

A.inv()
```

```
Matrix([
[17/8, -3/8, -5/4],
[ -5/2,  1/2,   2],
[ -3/4,  1/4,  1/2]])
```

Numerisches Lösen

Mit NumPy können LGS numerisch gelöst werden. Dies ist deutlich schneller als das symbolische Lösen. Bei kleinen LGS (z.B. die obigen Beispiele) spielt dieser Geschwindigkeitsvorteil keine Rolle. Praktisch relevante LGS sind aber viel größer und symbolisch nicht mehr in akzeptabler Zeit lösbar.

Das numerische Lösen funktioniert nur, wenn die Systemmatrix invertierbar ist.

```
import numpy as np
```

```
A = np.array([
    [2, 1, 1],
    [2, -1, 9],
    [2, 2, -1]
])
b = np.array([0, 6, 1])

np.linalg.solve(A, b)
```

```
array([-3.5,  5. ,  2. ])
```

Auch die Inverse kann numerisch berechnet werden:

```
A = np.array([
    [2, 1, 1],
    [2, -1, 9],
    [2, 2, -1]
])

np.linalg.inv(A)
```

```
array([[ 2.125, -0.375, -1.25 ],
       [-2.5   ,  0.5   ,  2.   ],
       [-0.75 ,  0.25 ,  0.5  ]])
```

Nebenbemerkung

Das numerische Lösen von LGS bedarf einiger Vorsicht, da viele praktisch relevante LGS “schlecht konditioniert” sind. Das bedeutet, dass selbst sehr kleine Änderungen an der rechten Seite zu sehr großen Änderungen in der Lösung führen können. Da die rechte Seite meist (fehlerbehaftete) Messwerte enthält und da aufgrund der Funktionsweise von Computern immer kleine Rundungsfehler auftreten, kann in solch einem Fall die numerisch berechnete Lösung stark von der tatsächlichen Lösung abweichen, bis hin zur Unbrauchbarkeit der “Lösung”. Diese komplexe Thematik ist ein wesentlicher Gegenstand der Numerik als Teilgebiet von Mathematik und Informatik. Bevor man ernsthaft große LGS numerisch löst, sollte man sich unbedingt genauer mit dieser Thematik beschäftigen!

Beispiel: Koordinaten bzgl. einer Basis

Hatten noch die Frage offen gelassen, wie wir zu einer gegebenen Basis $\{b_1, \dots, b_r\}$ in einem Unterraum U von $\mathbb{R}^n$ die Koordinaten eines Vektors $x \in U$ bzgl. dieser Basis berechnen können. Wie wir nun sehen werden, führt diese Frage auf ein LGS, dessen

Lösung gerade die gesuchten Koordinaten sind.

Sei $x \in \mathbb{R}$. Gesucht sind die Zahlen $\lambda_1, \dots, \lambda_r$, sodass

$$\sum_{k=1}^r \lambda_k b_k = x \quad (5.122)$$

gilt. Schreiben wir die Basisvektoren als Spalten in eine Matrix $B \in \mathbb{R}^{n \times r}$, so soll also

$$B \lambda = x \quad (5.123)$$

gelten (vgl. dazu auch Section 5.1).

Für einen echten Unterraum (also $r < n$) ist dieses LGS überbestimmt, wird also nicht für alle $x \in \mathbb{R}^n$ eine Lösung besitzen. Man kann jedoch zeigen, dass dieses LGS für $x \in U$ stets eindeutig lösbar ist. Für $U = \mathbb{R}^n$ ist die in diesem Fall quadratische Matrix B also insbesondere invertierbar.

Für den Fall $r = n$ können wir noch weitere nützliche Beobachtungen machen: Handelt es sich bei $\{b_1, \dots, b_n\}$ um eine ONB in $\mathbb{R}^n$, so folgt aus

$$x = \sum_{k=1}^n \lambda_k b_k \quad (5.124)$$

durch Bilden des Skalarprodukts mit b_l , dass

$$x \circ b_l = \left(\sum_{k=1}^n \lambda_k b_k \right) \circ b_l = \sum_{k=1}^n \lambda_k (b_k \circ b_l) = \lambda_l, \quad (5.125)$$

also

$$\lambda = \begin{bmatrix} x \circ b_1 \\ \vdots \\ x \circ b_n \end{bmatrix} = B^T x \quad (5.126)$$

gilt. Die gesuchten Koordinaten von x bzgl. $\{b_1, \dots, b_n\}$ erhält man damit ohne nennenswerte Mühe als Matrix-Vektor-Produkt.

Da die Gleichung $x = B B^T x$ für beliebige $x \in \mathbb{R}^n$ wie gezeigt hergeleitet werden kann, muss

$$B B^T = I \quad (5.127)$$

gelten. Offensichtlich gilt auch

$$B^T B = I. \quad (5.128)$$

Die Matrix B ist also invertierbar und die Inverse ist

$$B^{-1} = B^T. \quad (5.129)$$

Matrizen, bei denen die Inverse gerade die transponierte Matrix ist, spielen in verschiedenen Teilgebieten von Mathematik (und auch Physik) eine Rolle und bekommen deshalb einen eigenen Name:

Nebenbemerkung

Eine quadratische Matrix $A \in \mathbb{R}^{n \times n}$ heißt **Orthogonalmatrix**, wenn sie invertierbar ist und $A^{-1} = A^T$ gilt.

Folgende Eigenschaften von Orthogonalmatrizen sind gelegentlich hilfreich (und leicht einzusehen):

- Es gilt

$$A^T A = I \quad \text{und} \quad A A^T = I. \quad (5.130)$$

- Die Spalten von A bilden ein Orthonormalsystem (wegen $A^T A = I$).
- Die Zeilen von A bilden ein Orthonormalsystem (wegen $A A^T = I$).
- Die Transponierte A^T ist ebenfalls eine Orthogonalmatrix (weil $A = (A^T)^T$ ihre Inverse ist).
- Das Produkt zweier Orthogonalmatrizen $A, B \in \mathbb{R}^n$ ist wieder eine Orthogonalmatrix, denn

$$(A B)^{-1} = B^{-1} A^{-1} = B^T A^T = (A B)^T. \quad (5.131)$$

Nebenbemerkung

Beachte, dass die Begriffe sich hier etwas inkonsistent eingebürgert haben. Bei Orthogonalmatrizen sind die Spalten stets normiert, auch wenn das nicht aus dem Name hervorgeht. Bei Orthonormalsystemen müssen die Vektoren jedoch nicht normiert sein. Sind die Vektoren in letzterem Fall normiert, spricht von einem Orthonormalsystem. Der Begriff "Orthonormalmatrix" ist hingegen nicht gebräuchlich.

Beispiel: Kreuzprodukt

In der Ebene steht man gelegentlich vor der Aufgabe, zu einem gegebenen Vektor $a \in \mathbb{R}^2 \setminus \{0\}$ einen zweiten Vektor $b \in \mathbb{R}^2 \setminus \{0\}$ zu finden, sodass beide Vektoren senkrecht zueinander verlaufen. Analog möchte man gelegentlich im Raum zu zwei gegebenen, linear unabhängigen Vektoren $a, b \in \mathbb{R}^3 \setminus \{0\}$ einen dritten Vektor $c \in \mathbb{R}^3 \setminus \{0\}$ finden, der sowohl senkrecht zu a als auch senkrecht zu b verläuft. Diese Aufgabenstellung können wir als LGS formulieren:

- Für $a, b \in \mathbb{R}^2 \setminus \{0\}$ soll $a \circ b = 0$ gelten (b gesucht), d.h.

$$a_1 b_1 + a_2 b_2 = 0. \quad (5.132)$$

- Für $a, b, c \in \mathbb{R}^3 \setminus \{0\}$ sollen $a \circ c = 0$ und $b \circ c = 0$ gelten (c gesucht), d.h.

$$\begin{aligned} a_1 c_1 + a_2 c_2 + a_3 c_3 &= 0 \\ b_1 c_1 + b_2 c_2 + b_3 c_3 &= 0. \end{aligned} \quad (5.133)$$

In beiden Fällen handelt es sich um homogene LGS. Der Rang der Systemmatrizen

$$\begin{bmatrix} a_1 & a_2 \end{bmatrix} \quad \text{bzw.} \quad \begin{bmatrix} a_1 & a_2 & a_3 \\ b_1 & b_2 & b_3 \end{bmatrix} \quad (5.134)$$

ist 1 bzw. 2 (a und b sind als linear unabhängig vorausgesetzt!), also um 1 niedriger als die Anzahl der Unbekannten. Somit ist die Lösungsmenge (Unterraum!) in beiden Fällen eindimensional, d.h. die Richtung des gesuchten Vektors ist bis auf das Vorzeichen festgelegt. Da wir die Länge des gesuchten Vektors frei wählen können, wählen wir diese so, dass die Darstellung des Vektors als Formel möglichst einfach wird.

Durch Einsetzen in das entsprechende LGS prüft man leicht nach:

- In $\mathbb{R}^2$ erfüllt

$$b = \begin{bmatrix} -a_2 \\ a_1 \end{bmatrix} \quad (5.135)$$

die Anforderungen.

- In $\mathbb{R}^3$ erfüllt

$$c = \begin{bmatrix} a_2 b_3 - a_3 b_2 \\ a_3 b_1 - a_1 b_3 \\ a_1 b_2 - a_2 b_1 \end{bmatrix} \quad (5.136)$$

die Anforderungen.

Merke

Für zwei Vektoren $a, b \in \mathbb{R}^3$ heißt der Vektor

$$a \times b := \begin{bmatrix} a_2 b_3 - a_3 b_2 \\ a_3 b_1 - a_1 b_3 \\ a_1 b_2 - a_2 b_1 \end{bmatrix} \quad (5.137)$$

Kreuzprodukt (oder äußeres Produkt) von a und b .

Das Kreuzprodukt verläuft stets senkrecht zu a und senkrecht zu b .

Man kann zeigen, dass die Länge $\|a \times b\|$ des Kreuzprodukts gerade dem Flächeninhalt des von a und b aufgespannten Parallelogramms entspricht.

Zu zwei linear unabhängigen Vektoren liefert uns das Kreuzprodukt gerade einen Vektor, der zusammen mit den beiden Ausgangsvektoren ein **Rechtssystem** bildet. Wählt man stattdessen gerade einen Vektor in entgegengesetzter Richtung, so erhält man ein **Linkssystem**. Da diese beiden Begriffe im Folgenden keine Rolle spielen, verweisen wir für eine genauere Erklärung auf den Wikipedia-Artikel zu Rechtssystemen (beachte insbesondere die Rolle der Determinante bei der Unterscheidung von Rechts- und Linkssystemen!).

Nebenbemerkung

Das Kreuzprodukt kann man auch für Vektoren in $\mathbb{R}^n$ für $n \neq 3$ einführen. Allerdings ist das Kreuzprodukt dort eine Abbildung, die $n - 1$ gegebene Vektoren in einen neuen Vektor überführt, der senkrecht zu allen gegebenen Vektoren verläuft. Für $n = 2$ liefert dieses verallgemeinerte Kreuzprodukt gerade den oben von uns gefundenen Vektor b .

Für das Rechnen mit Kreuzprodukten kann man sich einige nützliche Zusammenhänge überlegen, z.B. $0 \times a = 0$ und $a \times b = -(b \times a)$. Für eine Übersicht solcher Zusammenhänge sei bei Bedarf auf die entsprechende Auflistung im Wikipedia-Artikel zum Kreuzprodukt verwiesen. Bei praktischen Aufgabenstellungen aus Vermessungswesen und Informatik wird man nur äußerst selten mit Kreuzprodukten rechnen müssen. Das Ergänzen zweier Vektoren zu einer Basis in $\mathbb{R}^3$ mittels Kreuzprodukt ist hingegen öfter von Nutzen.

5.6 Koordinatentransformationen

Eine wichtige Anwendung der bisher eingeführten Werkzeuge sind Koordinatentransformationen einschließlich damit verbundener Projektionstechniken (siehe unter für ein Praxisbeispiel). Ziel ist die Wahl von für den jeweiligen Einsatzzweck besonders gut geeigneten Koordinatensystemen zur Beschreibung der Position von Punkten und anderen geometrischen Objekten. Gut geeignet ist ein Koordinatensystem immer dann, wenn die Koordinaten der betrachteten Objekte einfach zu handhaben sind, also z.B. nur wenige Koordinaten von Null verschieden sind oder relevante Objekteigenschaften direkt aus den Koordinaten abgelesen werden können.

Konzentrieren uns hier auf lineare Koordinatentransformationen. Erwähnen aber zunächst zwei wichtige nichtlineare Transformationen.

Nichtlineare Koordinatentransformationen

Für die spätere Verwendung, jedoch ohne nähere Untersuchung, seien hier Polar- und Kugelkoordinaten erwähnt. Dabei handelt es sich um sehr spezielle Koordinatensysteme, die nicht durch lineare Transformationen aus dem Standardkoordinatensystem erhalten werden können. Unter “Standardkoordinatensystem” verstehen wir in $\mathbb{R}^2$ die Angabe von Koordinaten bezüglich der Basis $\{[1 \ 0]^T, [0 \ 1]^T\}$; in $\mathbb{R}^3$ analog. Solche Koordinatensysteme und entsprechende Transformationen sind Inhalt von spezielleren Lehrveranstaltungen zum Vermessungswesen.

Polarkoordinaten Als Polarkoordinaten versteht man die Darstellung eines Punktes (x, y) in der Ebene durch zwei Werte $r \geq 0$ und $\varphi \in [0, 2\pi)$, wobei r der Abstand zum Koordinatenursprung ist und φ den Winkel zwischen positiver x -Achse und dem Strahl vom Ursprung durch (x, y) angibt. Polarkoordinaten hatten wir bereits im Abschnitt Section 4.4 eingeführt; dort findet man auch eine grafische Darstellung des Sachverhalts.

Die Transformation von Polarkoordinaten zu Standardkoordinaten ist durch

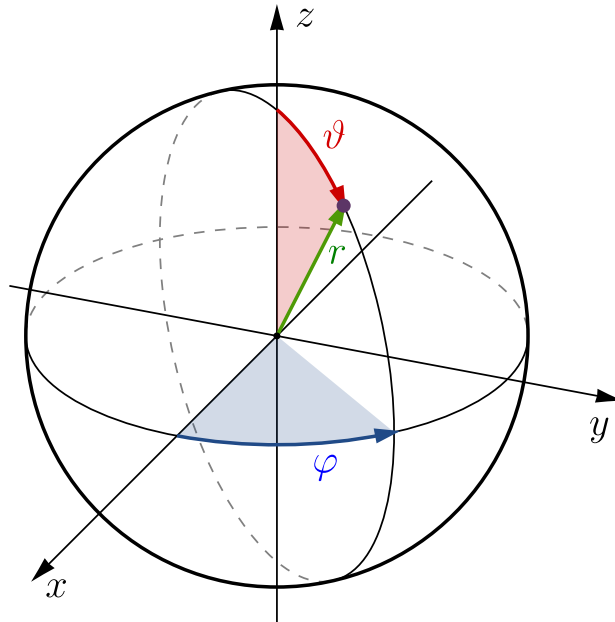
$$\begin{aligned} x &= r \cos \varphi, \\ y &= r \sin \varphi \end{aligned} \tag{5.138}$$

gegeben. In umgekehrter Richtung gilt

$$r = \sqrt{x^2 + y^2}. \tag{5.139}$$

Die Berechnung von φ lässt sich jedoch nur umständlich als Formel aufschreiben. Allerdings bieten mathematische Software-Bibliotheken (z.B. NumPy) meist eine Funktion `atan2` an. Dies ist eine spezielle Variante des Arcustangens, die zu den zwei zu übergebenden Argumenten x und y gerade den Winkel φ liefert.

Kugelkoordinaten Als Kugelkoordinaten versteht man die Darstellung eines Punktes $(x, y, z) \in \mathbb{R}^3$ im Raum durch drei Werte $r > 0$, $\varphi \in [0, 2\pi)$ und $\vartheta \in [0, \pi]$, wobei r der Abstand zum Ursprung ist, φ den Winkel zwischen positiver x -Achse und dem Strahl vom Ursprung durch $(x, y, 0)$ angibt und ϑ der Winkel zwischen positiver z -Achse und dem Strahl vom Ursprung durch (x, y, z) ist.



Lage der Winkel bei Kugelkoordinaten (Quelle: commons.wikimedia.org, Nutzer Inductiveloop, CC0, vom Autor bearbeitet).

Die Transformation von Kugelkoordinaten zu Standardkoordinaten ist durch

$$\begin{aligned} x &= r (\cos \varphi) \sin \vartheta, \\ y &= r (\sin \varphi) \sin \vartheta, \\ z &= r \cos \vartheta \end{aligned} \quad (5.140)$$

gegeben. In umgekehrter Richtung gilt

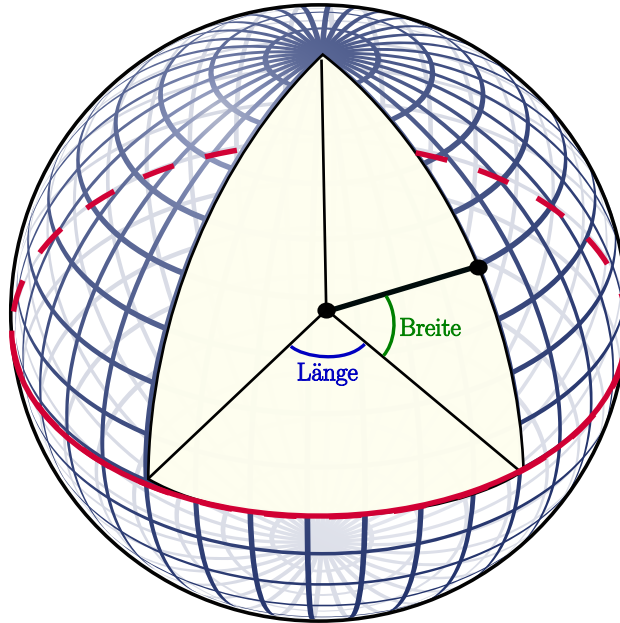
$$r = \sqrt{x^2 + y^2 + z^2}. \quad (5.141)$$

Die Berechnung der beiden Winkel lässt sich nur mühsam in Formeln ausdrücken und entfällt hier deshalb.

Beachte, dass die hier wiedergegebene und in der Mathematik übliche Form der Kugelkoordinaten von der in der Geografie üblichen Darstellung abweicht. Der Winkel φ entspricht der geografischen Länge eines Punktes, wobei in der Geografie meist das Intervall $(-\pi, \pi]$ verwendet wird. Der Winkel ϑ entspricht gerade dem Wert

$$\frac{\pi}{2} - \text{geografische Breite}. \quad (5.142)$$

Statt vom Nordpol ausgehend misst man in der Geografie also relativ zur Äquatorebene.



Lage der Winkel bei geografischen Koordinaten (Quelle: commons.wikimedia.org, Nutzer Peter Mercator, CC0, vom Autor bearbeitet).

Außerdem werden in der Geografie Breite und Länge üblicherweise im Gradmaß und nicht, wie in der Mathematik verbreitet, im Bogenmaß angegeben. In beiden Gebieten ist jedoch im Prinzip beides möglich, da sich die entsprechenden Zahlenwerte nur um einen konstanten Faktor unterscheiden:

$$\text{Winkel im Gradmaß} = \frac{180}{\pi} \cdot \text{Winkel im Bogenmaß}. \quad (5.143)$$

Koordinatensysteme

Haben den Begriff “Koordinatensystem” bisher eher intuitiv verwendet ohne genauere Definition. Präzisieren diesen Begriff nun, damit Missverständnisse vermieden werden. Beachte, dass wir im Weiteren nur geradlinige, rechtwinklige Koordinatensystem betrachten ohne dies explizit sprachlich sichtbar zu machen, also solche, die die folgenden beiden Bedingungen erfüllen:

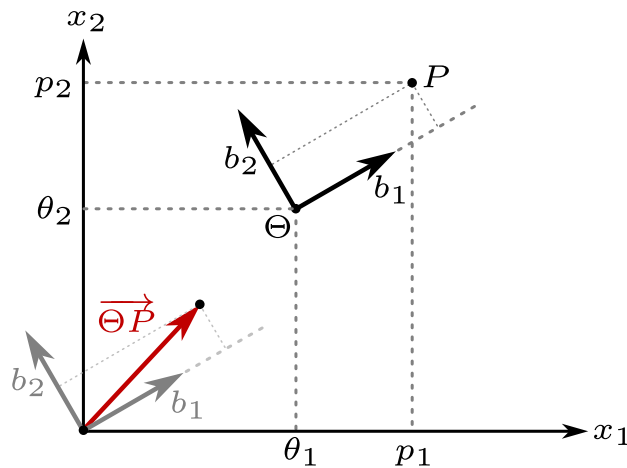
- Bei Veränderung genau einer Koordinate eines Punktes wird der Punkt entlang einer Gerade verschoben.
- Die Verschiebungsrichtungen bei Änderung jeweils genau einer Koordinate verlaufen senkrecht zueinander.

Beispielsweise liefern Polarkoordinaten und Kugelkoordinaten keine geradlinigen Koordinatensysteme, da das Verändern eines Winkels zu einer krummlinigen Bewegung des Punktes führt. Manchmal werden die hier betrachteten Koordinatensysteme als “kartesisch” bezeichnet. Diese Bezeichnung wird aber in der Literatur unterschiedlich verwendet.

Merke

Ein **Koordinatensystem** in $\mathbb{R}^n$ ist ein Paar (Θ, B) bestehend aus einem Punkt $\Theta = (\theta_1, \dots, \theta_n) \in \mathbb{R}^n$ und einer Basis $B = \{b_1, \dots, b_n\}$. Der Punkt Θ wird als **Ursprung** des Koordinatensystems bezeichnet. Die Basisvektoren geben die Richtungen der vom Ursprung ausgehenden **Koordinatenachsen** an.

Als **Koordinaten** eines Punktes $P = (p_1, \dots, p_n) \in \mathbb{R}^n$ bezüglich des Koordinatensystems (Θ, B) bezeichnen wir die Koordinaten des Vektors $\overrightarrow{\Theta P}$ bezüglich der Basis B .



Die Koordinaten des Punktes P bzgl. des Koordinatensystems (Θ, B) sind gerade die Koordinaten des Vektors $\overrightarrow{\Theta P}$ bzgl. der Basis B .

Im Folgenden bezeichnen wir mit $\{e_1, \dots, e_n\}$ stets die **Standardbasis** in $\mathbb{R}^n$ (e_k hat als k -te Komponente eine Eins, sonst Nullen). Zusammen mit dem Ursprung $(0, \dots, 0) \in \mathbb{R}^n$ beschreibt diese das **Standardkoordinatensystem** in $\mathbb{R}^n$.

Nebenbemerkung (Achsen vs. Basis)

Üblicherweise deutet man bei der grafischen Darstellungen eines Koordinatensystems die Koordinatenachsen als (lange) Pfeile an, die über den gesamten relevanten Bereich verlaufen, wie in der obigen Abbildung für die x - und die y -Achse

(Standardkoordinatensystem) geschehen. Die (kurzen) Basisvektoren selbst tauchen im Normalfall nicht auf, wurden jedoch des besseren Verständnisses halber in der Abbildung an Stelle der Achsen des Koordinatensystems (Θ, B) dargestellt.

Nebenbemerkung (griechische Buchstaben)

Das Zeichen Θ ist ein großes griechisches Theta und wurde hier aufgrund seiner Ähnlichkeit zur Null oder zu O gewählt. Das kleine griechische Theta gibt es in zwei Varianten: ϑ und θ . Für Winkel nutzen wir meist ϑ , da dieses nicht so schnell mit der Null verwechselt werden kann. Für die Komponenten eines Punktes Θ nutzen wir hingegen θ , da dieses mehr Ähnlichkeit zu Θ hat. Generell kann es nicht schaden, sich mit griechischen Buchstaben zu befassen, da diese in der Mathematik und all ihren Anwendungsgebieten (Naturwissenschaften, Technik,...) gern und oft verwendet werden. Eine gute Übersicht liefert der Wikipedia-Artikel zum griechischen Alphabet. Auch die meisten Formelsammlungen enthalten solche Übersichten.

Beachte, dass wir zwischen “Koordinaten bzgl. einer Basis” und “Koordinaten bzgl. eines Koordinatensystems” unterscheiden. Erstere werden relativ zum Punkt $(0, \dots, 0)$ angegeben, letztere relativ zum Ursprung des Koordinatensystems. Die Koordinaten bzgl. einer Basis B sind identisch mit den Koordinaten bzgl. des Koordinatensystems $(0, B)$ (hier steht 0 für den Nullpunkt $(0, \dots, 0)$).

Schreiben wir die Basisvektoren eines Koordinatensystems (Θ, B) als Spalten in eine Matrix, so können wir entsprechend den Ausführungen in Section 5.5 die Koordinaten eines Punktes P bzgl. des Koordinatensystems als Matrix-Vektor-Produkt berechnen. Um nicht zu viele verschiedene Symbole einzuführen, verwenden wir für die Basis und für die Matrix der Basisvektoren das selbe Symbol B . Je nach Kontext handelt es sich dabei also um eine Menge von Vektoren oder um eine Matrix! Für die Koordinaten $u_1, \dots, u_n$ eines Punktes $P = (p_1, \dots, p_n) \in \mathbb{R}^n$ bzgl. des Koordinatensystems (Θ, B) gilt nun

$$\begin{bmatrix} u_1 \\ \vdots \\ u_n \end{bmatrix} = B^T \begin{bmatrix} p_1 - \theta_1 \\ \vdots \\ p_n - \theta_n \end{bmatrix}. \quad (5.144)$$

Umgekehrt bekommt man aus den Koordinaten $u_1, \dots, u_n$ auch die Werte $p_1, \dots, p_n$ zurück (vorherige Formel umstellen; dabei $B^T = B^{-1}$ beachten):

$$\begin{bmatrix} p_1 \\ \vdots \\ p_n \end{bmatrix} = \begin{bmatrix} \theta_1 \\ \vdots \\ \theta_n \end{bmatrix} + B \begin{bmatrix} u_1 \\ \vdots \\ u_n \end{bmatrix}. \quad (5.145)$$

Nebenbemerkung (Koordinatensysteme in Unterräumen)

Im Prinzip müsste man auch Koordinatensysteme in Unterräumen und Untermannigfaltigkeiten des $\mathbb{R}^n$ diskutieren. Um diesen zusätzlichen Aufwand zu vermeiden, ergänzen wir eine ONB in einem Unterraum zu einer ONB in ganz $\mathbb{R}^n$ (dies ist immer möglich!) und wählen den Ursprung im Unterraum bzw. in der Untermannigfaltigkeit. Betrachten wir nun die Koordinaten eines Punktes bzgl. dieses Koordinatensystems im $\mathbb{R}^n$, so sind die Koordinaten zu den ergänzten Koordinatenachsen alle Null, sofern der Punkt im Unterraum bzw. der Untermannigfaltigkeit liegt. Die ergänzten Achsen sind also nur von formalem Interesse, werden aber nicht für die Beschreibung der relevanten Punkte/Vektoren benötigt.

Wechsel des Koordinatensystems

Als **lineare Koordinatentransformation** bezeichnen wir die Berechnung der Koordinaten eines Punktes bzgl. eines “neuen” Koordinatensystems aus den Koordinaten dieses Punktes bzgl. eines “alten” Koordinatensystems. Wie wir sehen werden, entspricht die Transformationsformel, die die alten Koordinaten in die neuen überführt, bis auf zusätzliche Verschiebungsoperationen einer linearen Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$.

Sei also (Θ, B) mit $\Theta = (\theta_1, \dots, \theta_n)$ und $B = \{b_1, \dots, b_n\}$ das “alte” Koordinatensystem und seien $u_1, \dots, u_n$ die Koordinaten eines Punktes $P = (p_1, \dots, p_n) \in \mathbb{R}^n$ bezüglich dieses Koordinatensystems. Weiter sei (Γ, C) mit $\Gamma = (\gamma_1, \dots, \gamma_n)$ und $C = \{c_1, \dots, c_n\}$ das “neue” Koordinatensystem und $v_1, \dots, v_n$ seien die Koordinaten des Punktes P bezüglich dieses Koordinatensystems. Wir wissen bereits, dass die folgenden Beziehungen gelten:

- Übergang vom Standardkoordinatensystem zu (Γ, C) :

$$\begin{bmatrix} v_1 \\ \vdots \\ v_n \end{bmatrix} = C^T \begin{bmatrix} p_1 - \gamma_1 \\ \vdots \\ p_n - \gamma_n \end{bmatrix}. \quad (5.146)$$

- Übergang von (Θ, B) zum Standardkoordinatensystem:

$$\begin{bmatrix} p_1 \\ \vdots \\ p_n \end{bmatrix} = \begin{bmatrix} \theta_1 \\ \vdots \\ \theta_n \end{bmatrix} + B \begin{bmatrix} u_1 \\ \vdots \\ u_n \end{bmatrix}. \quad (5.147)$$

Setzen wir die zweite Formel in die erste ein, so erhalten wir die gesuchte Transformati-

onsvorschrift:

$$\begin{aligned}
 \begin{bmatrix} v_1 \\ \vdots \\ v_n \end{bmatrix} &= C^T \begin{bmatrix} p_1 - \gamma_1 \\ \vdots \\ p_n - \gamma_n \end{bmatrix} = C^T \begin{bmatrix} p_1 \\ \vdots \\ p_n \end{bmatrix} - C^T \begin{bmatrix} \gamma_1 \\ \vdots \\ \gamma_n \end{bmatrix} \\
 &= C^T \begin{bmatrix} \theta_1 \\ \vdots \\ \theta_n \end{bmatrix} + C^T B \begin{bmatrix} u_1 \\ \vdots \\ u_n \end{bmatrix} - C^T \begin{bmatrix} \gamma_1 \\ \vdots \\ \gamma_n \end{bmatrix} \\
 &= C^T \begin{bmatrix} \theta_1 - \gamma_1 \\ \vdots \\ \theta_n - \gamma_n \end{bmatrix} + C^T B \begin{bmatrix} u_1 \\ \vdots \\ u_n \end{bmatrix}.
 \end{aligned} \tag{5.148}$$

Der Übergang zwischen zwei Koordinatensystemen besteht also aus dem Anwenden einer linearen Abbildung (Multiplikation mit einer Matrix) und einer Verschiebung.

Nebenbemerkung

Als Plausibilitätstest für die gefundene Transformationsformel kann man den Spezialfall betrachten, dass wahlweise das alte oder das neue Koordinatensystem das Standardkoordinatensystem ist. Dann sollte sich die Transformationsformel auf die Formel zur Berechnung der Koordinaten eines Punktes bzgl. eines Koordinatensystems reduzieren.

Nebenbemerkung

Meistens benötigt man die hier hergeleitete Transformationsformel für den Wechsel des Koordinatensystems gar nicht, wenn man, wie wir, stets das Standardkoordinatensystem als Arbeitsgrundlage nutzt, also keine koordinatenfreie Geometrie betreibt (vgl. Nebenbemerkung (koordinatenfreie Geometrie)). Man ist praktisch nur mit Übergängen zwischen dem Standardkoordinatensystem und einem anderen Koordinatensystem konfrontiert, kaum mit dem Übergang zwischen zwei verschiedenen Nicht-Standardkoordinatensystemen.

Beispiel: Orthografische Azimuthalprojektion

Problemstellung und Lösungsansatz Für eine Region (z.B. Deutschland) soll eine möglichst verzerrungsarme, zweidimensionale Kartendarstellung gefunden werden. Der Einfachheit halber legen wir uns auf orthografische Azimuthalprojektionen fest. Bei diesen ist die Verzerrung in der Kartenmitte am geringsten, sodass wir als Kartenmitte die “Mitte” der darzustellenden Region wählen. Unter “Mitte” verstehen wir hier den Punkt auf der Erdoberfläche, für den die maximale Entfernung zu allen anderen Punkten der Region am geringsten ist. Für Deutschland ist dies 10.456555° östliche Länge, 51.237923°

nördliche Breite (Quelle: Administrative units/azimuthal equidistant projections datasets).

Zu einem durch geografische Länge und Breite gegebenen Punkt auf der Erdoberfläche sind die Koordinaten bezüglich des Kartenkoordinatensystems gesucht. Der Ursprung dieses Koordinatensystems liegt im Mittelpunkt der Region und die relevante Kartenebene (die Ebene, die von den ersten beiden Koordinatenachsen aufgespannt wird) soll die als Kugeloberfläche angenommene Erdoberfläche in diesem Punkt berühren, aber nicht schneiden (Stichwort: Tangentialebene; diesen Begriff diskutieren wir nächstes Semester genauer).

Aufgabe: Man finde eine Formel, die Längen- und Breitenangaben in Kartenkoordinaten überführt.

Wir werden diese Aufgabe in drei Schritten lösen:

1. Längen- und Breitenangaben in Punkte in $\mathbb{R}^3$ überführen.
2. Eine Basis für die Beschreibung des Kartenkoordinatensystems wählen (der Ursprung ist schon geklärt).
3. Die Berechnung der Koordinaten eines Punktes bzgl. des Kartenkoordinatensystems klären.
4. Eine Gesamtformel “zusammenbauen”, die die Aufgabenstellung löst.

Punkte im Raum Seien $\lambda \in (-180, 180]$ und $\zeta \in [-90, 90]$ geografische Länge und Breite eines Punktes auf der Erdoberfläche im Gradmaß. Wir legen die als Kugeloberfläche angenommene Erdoberfläche so in den $\mathbb{R}^3$, dass

- der Erdmittelpunkt bei $(0, 0, 0)$ liegt,
- die positive x -Achse die Erdoberfläche bei 0° Länge und 0° Breite durchstößt,
- die positive y -Achse die Erdoberfläche bei 90° östlicher Länge und 0° Breite durchstößt
- die positive z -Achse den Nordpol durchstößt (90° nördliche Breite).

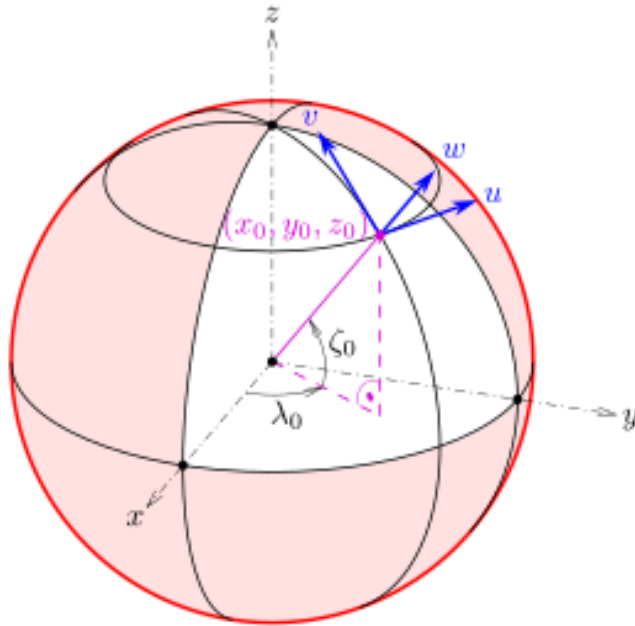
Für einen durch geografische Längen- und Breitenangaben λ und ζ beschriebenen Punkt P auf der Erdoberfläche erhalten wir die Koordinaten bzgl. des Standardkoordinatensystems in $\mathbb{R}^3$ somit als

$$\begin{aligned} x &= r \left(\cos \left(\frac{\pi}{180} \lambda \right) \right) \sin \left(\frac{\pi}{2} - \frac{\pi}{180} \zeta \right), \\ y &= r \left(\sin \left(\frac{\pi}{180} \lambda \right) \right) \sin \left(\frac{\pi}{2} - \frac{\pi}{180} \zeta \right), \\ z &= r \cos \left(\frac{\pi}{2} - \frac{\pi}{180} \zeta \right). \end{aligned} \tag{5.149}$$

Dabei ist $r = 6371$ der mittlere Erdradius in Kilometer. Die Koordinaten x, y, z sind entsprechend ebenfalls als Entfernungen vom Ursprung in Kilometer zu interpretieren.

Die Koordinaten des Mittelpunkts der betrachteten Region bezeichnen wir mit x_0, y_0, z_0 . Diese entstehen mittels obiger Formeln aus Länge λ_0 und Breite ζ_0 des Mittelpunkts der Region.

Kartenkoordinatensystem Legen nun Schritt für Schritt die Richtungen der Koordinatenachsen des Kartenkoordinatensystems fest. Beachte dabei, dass das gewünschte Koordinatensystem mit den ersten beiden Achsen die horizontale und die vertikale Achse der Karte beschreiben soll. Die dritte Achse steht senkrecht auf der Kartenebene und spielt in der zweidimensionalen Kartendarstellung keine Rolle. Wir bezeichnen die den Achsen entsprechenden Basisvektoren in dieser Reihenfolge mit $b_1, b_2, b_3 \in \mathbb{R}^3$.



Übersicht über die beteiligten Größen und Symbole. Die u -, v - und w -Achsen sind die durch die Basisvektoren b_1, b_2, b_3 definierten Koordinatenachsen (Quelle: commons.wikimedia.org, Nutzer Ag2gaeh, CC BY-SA 4.0, vom Autor bearbeitet).

Da der Vektor b_3 senkrecht zur Kartenebene verlaufen soll, entspricht dieser gerade der Richtung des Strahls vom Kugelmittelpunkt durch den Mittelpunkt (x_0, y_0, z_0) der betrachteten Region. Somit wählen wir

$$b_3 := \frac{1}{\sqrt{x_0^2 + y_0^2 + z_0^2}} \begin{bmatrix} x_0 \\ y_0 \\ z_0 \end{bmatrix}. \quad (5.150)$$

Der Vektor b_2 muss senkrecht zu b_3 verlaufen. Außerdem möchten wir die Kartendarstellung gern so einrichten, dass “Norden oben ist”. Bei einer gedachten Bewegung von

(x_0, y_0, z_0) ausgehend mit Richtung b_2 muss also die z -Achse durchstoßen werden. Elementargeometrische Berechnungen (Schnittpunkt der zu b_3 senkrechten Ebene durch (x_0, y_0, z_0) mit der z -Achse; b_2 ist dann der Vektor von (x_0, y_0, z_0) zu diesem Schnittpunkt zuzüglich Normierung) liefern

$$b_2 := \frac{\tilde{b}_2}{\|\tilde{b}_2\|} \quad \text{mit} \quad \tilde{b}_2 := \begin{bmatrix} -x_0 \\ -y_0 \\ \frac{x_0^2 + y_0^2 + z_0^2}{z_0} - z_0 \end{bmatrix}. \quad (5.151)$$

Der Vektor b_1 ergibt sich nun als Kreuzprodukt von b_2 und b_3 zuzüglich Normierung:

$$b_1 := \frac{b_2 \times b_3}{\|b_2 \times b_3\|} \quad (5.152)$$

(auf eine explizite Formel verzichten wir hier; diese Rechnung erledigt für konkrete x_0, y_0, z_0 später der Computer).

Transformationsformel Schreiben wir die Basisvektoren b_1, b_2, b_3 als Spalten in eine Matrix $B \in \mathbb{R}^{3 \times 3}$, so erhalten wir für die Koordinaten u, v, w eines Punktes $(x, y, z) \in \mathbb{R}^3$ bzgl. des Kartenkoordinatensystems die Formel

$$\begin{bmatrix} u \\ v \\ w \end{bmatrix} = B^T \begin{bmatrix} x - x_0 \\ y - y_0 \\ z - z_0 \end{bmatrix}. \quad (5.153)$$

Umsetzung am Computer Zum Test der Transformation möchten wir die Grenze von Deutschland in einer Karte darstellen. Die Datei `germany.txt` enthält geografische Länge und Breite von 3000 Punkten auf der Grenze. Diese wurden zufällig aus der durch 159230 Punkte dargestellten Grenze von Deutschland in OpenStreetMap ausgewählt (Stand: 23.08.2026). Das Einlesen der Punkte kann mit `numpy.genfromtxt` erfolgen. Anschließend erfolgt die oben hergeleitete Transformation in Kartenkoordinaten.

```
import numpy as np
import matplotlib.pyplot as plt

# Erdradius
r = 6371

# Mittelpunkt von Deutschland
lon0 = 10.456555
lat0 = 51.237923

# Punkte auf der Grenze von Deutschland
lonlat = np.genfromtxt('germany.txt')
```

```

lon = lonlat[:, 0]
lat = lonlat[:, 1]

# Transformation zu Punkten im Raum
fac = np.pi / 180 # Umrechnungsfaktor Gradmaß zu Bogenmaß
x0 = r * np.cos(fac * lon0) * np.sin(np.pi / 2 - fac * lat0)
y0 = r * np.sin(fac * lon0) * np.sin(np.pi / 2 - fac * lat0)
z0 = r * np.cos(np.pi / 2 - fac * lat0)
x = r * np.cos(fac * lon) * np.sin(np.pi / 2 - fac * lat)
y = r * np.sin(fac * lon) * np.sin(np.pi / 2 - fac * lat)
z = r * np.cos(np.pi / 2 - fac * lat)

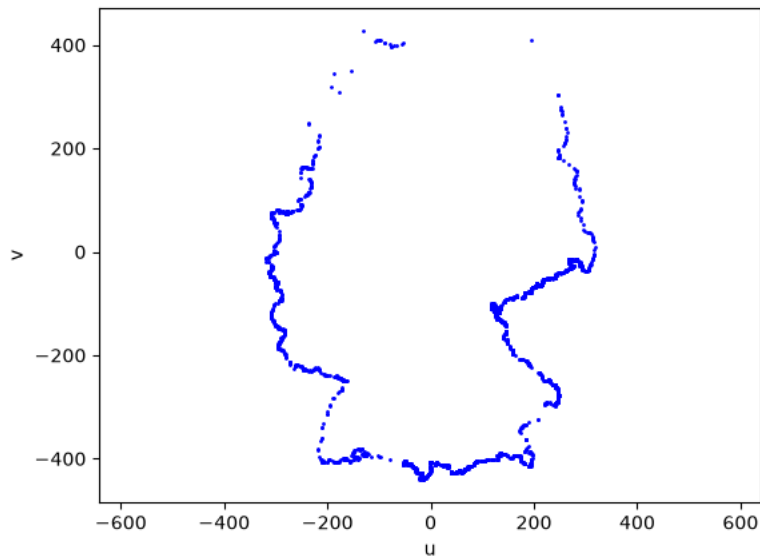
# Basis für Kartenkoordinatensystem
b3 = np.array([x0, y0, z0])
b3 = b3 / np.linalg.norm(b3)
b2 = np.array([-x0, -y0, (x0 ** 2 + y0 ** 2 + z0 ** 2) / z0 - z0])
b2 = b2 / np.linalg.norm(b2)
b1 = np.cross(b2, b3)
b1 = b1 / np.linalg.norm(b1)

# Koordinaten bzgl. des Kartenkoordinatensystems
B = np.column_stack([b1, b2, b3])
uvw = B.T @ np.vstack([x - x0, y - y0, z - z0])

# jede Spalte von uvw entspricht einem Punkt
u = uvw[0, :]
v = uvw[1, :]
w = uvw[2, :]

# Visualisierung
fig, ax = plt.subplots()
scatter_plot = ax.scatter(u, v, s=1, c='blue')
ax.axis('equal') # gleiche Achseneinteilung (keine Verzerrung)
ax.set_xlabel('u')
ax.set_ylabel('v')
plt.show()

```



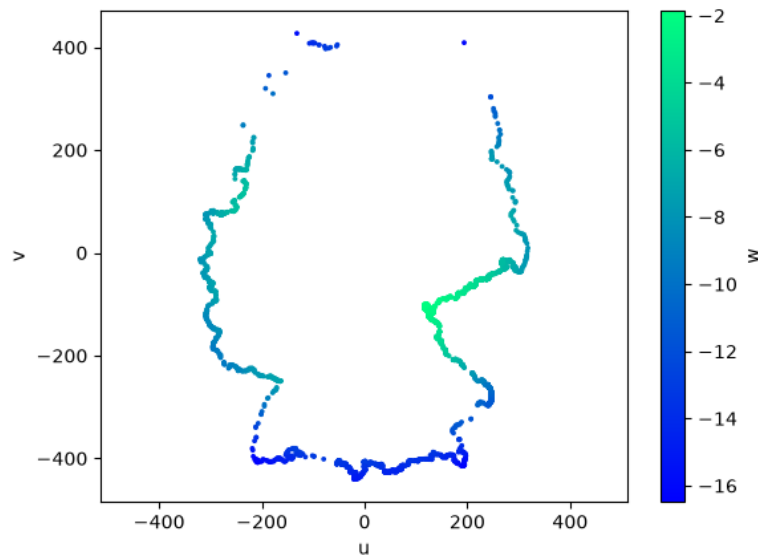
Die Ursache für das Fehlen von Punkten in den Küstenregionen liegt im einfachen Grenzverlauf dort. Dieser lässt sich mit deutlich weniger Punkten beschreiben, sodass bei der zufälligen Auswahl von Punkten entsprechend weniger Punkte in Küstenregionen ausgewählt werden.

Wir können auch die w -Koordinate visualisieren, indem wir die einzelnen Punkte entsprechend einfärben:

```
fig, ax = plt.subplots()

scatter_plot = ax.scatter(u, v, s=3, c=w, cmap='winter')
plt.colorbar(scatter_plot, label='w')

ax.axis('equal') # gleiche Achseneinteilung (keine Verzerrung)
ax.set_xlabel('u')
ax.set_ylabel('v')
plt.show()
```



Wir sehen dass die w -Koordinate zur Kartenmitte hin näher bei Null liegt als zum Kartenrand hin. Diese Beobachtung deckt sich mit der Theorie, denn die w -Koordinate gibt den Abstand eines Punktes auf der Erdoberfläche von der Kartenebene an. Insbesondere müssen alle w -Koordinaten negativ sein. An der Nordgrenze und an der Südgrenze beträgt der Abstand zwischen Kartenebene und der (als Kugeloberfläche angenommenen) Erdoberfläche bis zu 16 Kilometer.

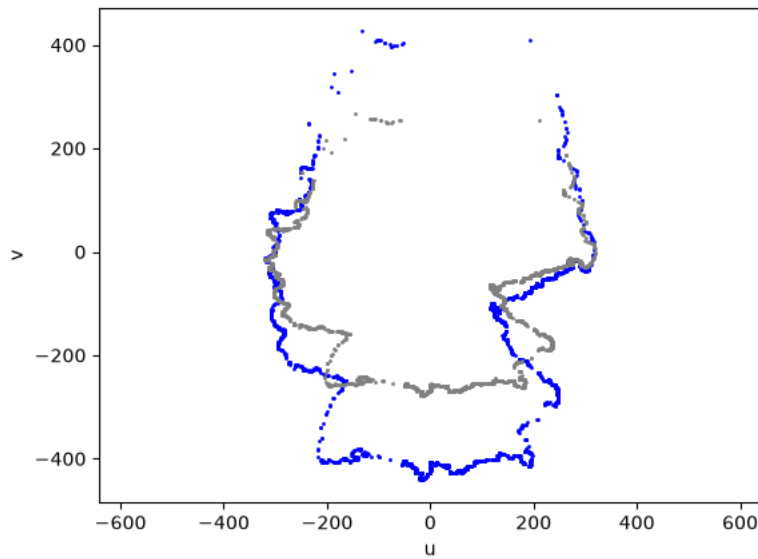
Ohne Kenntnis der Techniken zur Koordinatentransformation hätten wir nur die Längen- und Breitenwerte direkt als Koordinaten in einer zweidimensionalen Darstellung verwenden können. Dies liefert eine sogenannte Plattkarte. Zum Vergleich legen wir beide Varianten übereinander und skalieren die Plattkarte so, dass bei $v \approx 0$ die u -Koordinaten übereinstimmen (der Skalierungsfaktor wurde hier manuell durch Probieren verschiedener Werte gewählt).

```
fig, ax = plt.subplots()

scatter_plot = ax.scatter(u, v, s=1, c='blue')

fac = 70 # Skalierungsfaktor für Plattkarte
ax.scatter(fac * (lon - lon0), fac * (lat - lat0), s=1, c='gray')

ax.axis('equal') # gleiche Achseneinteilung (keine Verzerrung)
ax.set_xlabel('u')
ax.set_ylabel('v')
plt.show()
```



Offensichtlich führt die (triviale) Darstellung als Plattkarte (grau) zu einer deutlich stärkeren Verzerrung als die mittels linearer Koordinatentransformation erhaltene orthografische Azimuthalprojektion (blau).

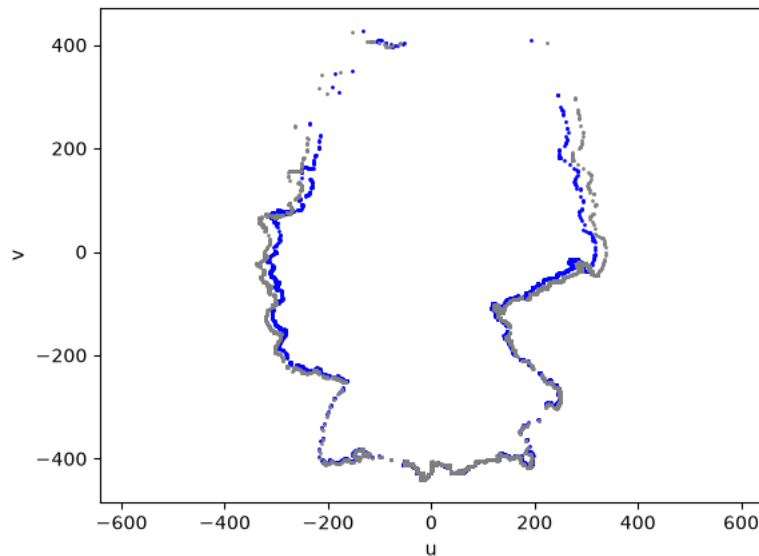
Auch bei unterschiedlicher Skalierung in horizontaler und vertikaler Richtung ist die Plattkarte nicht mit der sauber projizierten Karte in Deckung zu bringen. Wählt man die Skalierungsfaktoren beispielsweise so, dass im Süden gute Übereinstimmung entsteht, sind die Abweichungen im Norden groß:

```
fig, ax = plt.subplots()

scatter_plot = ax.scatter(u, v, s=1, c='blue')

ax.scatter(74 * (lon - lon0), 111 * (lat - lat0), s=1, c='gray')

ax.axis('equal') # gleiche Achseneinteilung (keine Verzerrung)
ax.set_xlabel('u')
ax.set_ylabel('v')
plt.show()
```



Einordnung In diesem Beispiel haben wir gesehen, wie wir mit den bisher beleuchteten Werkzeugen und Zusammenhängen durch Kugelkoordinaten oder Standardkoordinaten im $\mathbb{R}^3$ gegebene Punkte durch zwei Koordinaten in einer geeignet gewählten Ebene (Karte) darstellen können. Die Kartenebene ist in der mathematischen Sprache eine lineare Untermannigfaltigkeit des $\mathbb{R}^3$. Betrachten wir jedoch nur die ersten beiden Koordinaten eines Punktes bzgl. des passend zu dieser Untermannigfaltigkeit gewählten Koordinatensystems, so können wir die ganze Komplexität der dreidimensionalen Situation außer Acht lassen. Letztlich sind die wesentlichen Informationen durch zwei Zahlen (Koordinaten) darstellbar, die wir übersichtlich in unserer Zeichenebene (Papier, Laptop-Display,...) in einem einfachen zweidimensionalen Koordinatensystem darstellen können. In diesem Sinne macht die Mathematik das Leben hier deutlich einfacher!

Das Vernachlässigen der dritten Dimension führt hier nur zu geringem Informationsverlust, da die Erdkrümmung (wenigstens für kleine Kartenausschnitte) sehr gering ist. Den Gedanke, durch geeignete Wahl von Koordinatensystemen Informationen mit nur wenig Verlust deutlich einfacher darstellen zu können, werden wir im nächsten Unterkapitel nochmal aufnehmen. Dort wird dann auch klar, dass das hier demonstrierte Vorgehen auch in $\mathbb{R}^n$ mit sehr großem n praktisch relevant ist, die Mühen der Formulierung aller Begriffe und Zusammenhänge für allgemeines n also kein (scheinbar) unnötiger Ballast waren.

Die Suche nach möglichst verzerrungsarmen Kartendarstellungen hat insbesondere den Zweck, dass in der Karte gemessene Abstände möglichst nah an den tatsächlichen Abständen auf der Erdoberfläche liegen sollen. Je nach gewählter Kartenprojektion wird die Abweichung zwischen Abständen in der Karte und den realen Abständen unterschiedlich groß sein. Wir haben hier eine sehr einfache Kartenprojektion gewählt. Komplexere Kartenprojektionen und die ganze Fehlerthematik beim Messen von Abständen in Karten

werden in anderen Veranstaltungen Ihres Studiums vertieft werden.

5.7 Ausblick: Eigenwerte und Eigenvektoren

Dieses Unterkapitel ist für die Prüfung nicht relevant. Es dient lediglich der ersten Einordnung einiger Begriffe, die Ihnen im Laufe Ihres Studiums begegnen werden. Wir haben mit der Einführung der Ideen und Konzepte rund um Vektoren und Matrizen viel Vorarbeit geleistet; die Früchte unserer Saat werden dann in anderen Veranstaltungen geerntet. Gönnen wir uns hier wenigstens einen kleinen Vorgeschmack...

Idee

Koordinatensysteme wählt man üblicherweise so, dass die auszuführenden Arbeiten (Berechnungen, Visualisierungen,...) möglichst einfach werden. Sind zahlreiche Matrix-Vektor-Multiplikationen mit einer Matrix A durchzuführen, so kann man die Frage stellen, ob die Berechnung von Matrix-Vektor-Produkten in einem an die konkrete Matrix A angepassten Koordinatensystem eventuell deutlich einfacher wird.

In der Tat ist dies der Fall: Zu jeder symmetrischen (!) Matrix $A \in \mathbb{R}^{n \times n}$ gibt es eine ONB, sodass die Matrix-Vektor-Multiplikation bei Darstellung des Vektors bzgl. dieser ONB zu einer elementweisen Multiplikation des Vektors mit einem festen Vektor wird. Sind $u_1, \dots, u_n$ die Koordinaten eines Vektors $x \in \mathbb{R}^n$ bzgl. dieser speziellen Basis, so gilt also

$$Ax = \begin{bmatrix} \lambda_1 u_1 \\ \vdots \\ \lambda_n u_n \end{bmatrix} \quad (5.154)$$

mit festen Zahlen $\lambda_1, \dots, \lambda_n$, die aus den Einträgen der Matrix A berechnet werden können. Statt des aufwendigen Matrix-Vektor-Produkts mit n^2 Multiplikationen und $n(n-1)$ Additionen sind nun nur noch n Multiplikationen nötig. Für $n = 1000$ (für heutige Verhältnisse also eher klein) reduziert sich der Aufwand beispielsweise von 1000000 Multiplikationen und Additionen auf 1000 Multiplikationen.

Diese speziell an die Matrix angepasste Basis heißt **Eigenvektorbasis** von A ; die Basisvektoren heißen **Eigenvektoren**. Die festen Faktoren λ_k für die Matrix-Vektor-Multiplikation heißen **Eigenwerte**.

Präzisierung

Führen die obige Idee noch in Formeln aus.

Merke

Sei $A \in \mathbb{R}^{n \times n}$ eine symmetrische Matrix. Dann existieren eine ONB $\{b_1, \dots, b_n\}$ und Zahlen $\lambda_1, \dots, \lambda_n$, sodass

$$A b_k = \lambda_k b_k \quad \text{für } k = 1, \dots, n \quad (5.155)$$

gilt. Die Vektoren b_k heißen **Eigenvektoren** von A . Die Zahlen λ_k heißen **Eigenwerte** von A .

Nebenbemerkung

Die Existenz von Eigenvektoren und Eigenwerten ist keineswegs einfach nachzuweisen. Deshalb findet man oft folgende Definition, die die Existenzfrage zunächst beiseitelässt; ob die so definierten Objekte (Eigenvektoren, Eigenwerte) überhaupt existieren, muss dann im Nachgang geklärt werden:

Sei $A \in \mathbb{R}^{n \times n} \setminus \{0\}$ eine symmetrische Matrix. Ein Vektor $b \in \mathbb{R}^n$ heißt Eigenvektor von A , wenn es eine Zahl λ gibt, sodass $A b = \lambda b$ gilt. Die Zahl λ heißt dann Eigenwert von A .

Bei dieser Definition sind automatisch auch alle Vielfachen eines Eigenvektors wieder Eigenvektoren. Hat man sich (recht mühsam) von der Existenz von Eigenvektoren überzeugt, so muss man noch die Existenz einer Orthonormalbasis aus Eigenvektoren zeigen (ebenfalls mühsam).

Sind $u_1, \dots, u_n$ die Koordinaten eines Vektors $x \in \mathbb{R}^n$ bzgl. der Eigenvektorbasis einer Matrix A , so folgt insbesondere

$$A x = A \left(\sum_{k=1}^n u_k b_k \right) = \sum_{k=1}^n u_k A b_k = \sum_{k=1}^n u_k \lambda_k b_k, \quad (5.156)$$

d.h. die Koordinaten von $A x$ bzgl. der Eigenvektorbasis sind $\lambda_1 u_1, \dots, \lambda_n u_n$.

Auf die Berechnung von Eigenwerten und Eigenvektoren gehen wir hier nicht näher ein. Diese ist manuell recht mühsam und wird üblicherweise mit dem Computer erledigt, insbesondere da die Matrizen bei praktisch relevanten Aufgabenstellungen viel zu groß für manuelles Rechnen sind.

Nebenbemerkung

Für nicht symmetrische quadratische Matrizen kann man die Begriffe “Eigenwert” und “Eigenvektor” ebenfalls einführen. Allerdings existieren die Eigenwerte dann nicht zwingend als reelle Zahlen. Dehnt man den Zahlenbereich auf komplexe Zahlen aus, so funktioniert alles im Wesentlichen wieder wie bei symmetrischen Matrizen. Allerdings findet man nicht zwingend eine ONB aus Eigenvektoren, jedoch stets eine Basis.

Für nicht quadratische Matrizen $A \in \mathbb{R}^{m \times n}$ sind Eigenwerte und Eigenvektoren aufgrund der unterschiedlichen Länge der Vektoren $x \in \mathbb{R}^n$ und $Ax \in \mathbb{R}^m$ nicht sinnvoll einführbar. Es gibt allerdings ein sehr ähnliches Konstrukt, welches als **Singulärwertzerlegung** in der Literatur zu finden ist. Man bedient sich hier zweier an die Matrix angepasster Koordinatensysteme, eins in $\mathbb{R}^n$ und eins in $\mathbb{R}^m$, und kann die Matrix-Vektor-Multiplikation dann wieder als elementweise Multiplikation mit den sogenannten Singulärwerten der Matrix darstellen. Für symmetrische Matrizen entsprechen die Singulärwerte gerade den Eigenwerten.

Anwendungen

Gehen kurz auf einige Anwendungen von Eigenwerten und Eigenvektoren ein. Es zeigt sich, dass Eigenwerte und -vektoren sehr relevante Interpretationen in verschiedensten praktischen Problemen haben. Entsprechend häufig wird man ihnen in der Praxis begegnen; teils auch ohne dies zu bemerken.

Hauptkomponentenanalyse Im heutigen Datenzeitalter wird alles an Daten gesammelt, was verfügbar ist. Dies führt zu der Frage, wie man relevante Informationen aus umfangreichen Datensammlungen extrahiert. Sollen beispielsweise klimatische Prozesse untersucht werden, so wird man zunächst alle möglichen Parameter aufzeichnen ohne im Vorfeld detailliert zu klären, ob diese überhaupt relevant sind: Temperaturen in verschiedenen Höhen und an verschiedenen Orten, Luftfeuchte, Bewölkung, Sonnenstand, Windgeschwindigkeiten usw. (alles, was irgendwie messbar ist). Die Messungen könnten zum Beispiel stündlich oder täglich stattfinden. Am Ende stehen eine Vielzahl von sogenannten Datenpunkten zur Verfügung. Jeder Datenpunkt entspricht der Messung zu einem Zeitpunkt und enthält (als Koordinaten bzgl. der Standardbasis) die Messwerte (Temperatur, Luftfeuchte usw.). Bei n Messungen von m verschiedenen Parametern entstehen so n Datenpunkte $x_1, \dots, x_n \in \mathbb{R}^m$.

Die Anzahl m der Parameter ist meist sehr groß (100 und mehr; gelegentlich auch im Bereich von Millionen). Auch die Anzahl n der Messungen ist stets sehr groß (oft ebenfalls im Bereich von Millionen). Wie kann man anhand dieser Datenmasse nun sehen, welche Parameter oder Parameterkombination (!) eigentlich relevant sind? Dabei sehen wir einen Parameter bzw. eine gewichtete Kombination von Parametern als relevant an, wenn die entsprechenden Messwerte einen vergleichsweise großen Wertebereich aufweisen. Parameter(-kombinationen), für die bei jeder Messung fast der gleiche Wert gemessen wird, sind hingegen nicht relevant.

Die Antwort findet man mit der so genannten Hauptkomponentenanalyse. Aus mathematischer Sicht werden folgende Schritte ausgeführt:

1. Schreibe die Messpunkte $x_k \in \mathbb{R}^m$ als Zeilen in eine Matrix $X \in \mathbb{R}^{n \times m}$.
2. Berechne die Eigenwerte und Eigenvektoren zur symmetrischen Matrix $X^T X \in$

$\mathbb{R}^{m \times m}$.

3. Wähle die Eigenvektoren zu den größten Eigenwerten aus; b_1 sei der Eigenvektor zum größten Eigenwert, b_2 der zum zweitgrößten usw.
4. Bilde die Skalarprodukte $x_k \circ b_1, x_k \circ b_2, \dots$ für alle Datenpunkte x_k und die gewählten Eigenvektoren. Jedes Skalarprodukt generiert eine Parameterkombination aus dem Datenpunkt x_k . Diese so aus den ursprünglichen Parametern erzeugten Parameterkombinationen haben den höchstmöglichen Informationsgehalt, der aus der gesamten Datenmenge extrahiert werden kann. Insbesondere sind auch bei großer Parameteranzahl m meist nur sehr wenige Eigenvektoren nötig, um alle relevanten Informationen darzustellen.

Um einzusehen, dass dieses Vorgehen sinnvoll ist, benötigt man einiges an Statistik-Wissen. Wir verweisen an dieser Stelle auf weiterführende Lehrveranstaltungen. Siehe auch Principal component analysis (der deutsche Wikipedia-Artikel zum Thema ist leider qualitativ etwas unterdurchschnittlich, Stand: 23.08.2026).

Schwingungsanalysen In Elektrotechnik, Mechanik und anderen Teilgebieten von Physik und Technik werden unter anderem Schwingungsprozesse untersucht. Insbesondere geht man der Frage nach, bei welchen Erregerfrequenzen Resonanz auftritt. Manchmal ist Resonanz erwünscht (Akustik, elektromagnetische Wellen,...), meist jedoch nicht (z.B. bei Brücken, vgl. Video der Zerstörung der Tacoma Narrows Bridge (via Wikipedia)).

Das Verhalten vieler technischer Systeme mit schwingenden Komponenten lässt sich in Form von linearen Gleichungssystemen modellieren und/oder simulieren. Die Eigenfrequenzen des untersuchten Systems, also die Frequenzen, die bei entsprechender Fremderregung zu Resonanz führen, sind dann gerade die Eigenwerte der Systemmatrix des LGS. Die zugehörigen Eigenvektoren enthalten zusätzliche Informationen über die meist recht komplexen Schwingungsvorgänge.

PageRank-Algorithmus Der PageRank-Algorithmus bildete die Grundlage der ursprünglichen Google-Suchmaschine. Bei diesem werden die Verlinkungen zwischen verschiedenen Webseiten genutzt, um auf die Wichtigkeit jeder einzelnen Webseite und damit auf die Sortierung von Suchergebnissen zu schließen. Schreibt man alle Gewichte (je größer das Gewicht, desto wichtiger die Webseite) in einen langen Vektor $x \in \mathbb{R}^n$, wobei n die Anzahl der der Suchmaschine bekannten Webseiten ist, so kann man sich überlegen, dass

$$Ax = x \tag{5.157}$$

mit einer gewissen Matrix $A \in \mathbb{R}^{n \times n}$ gelten muss. In dieser Matrix ist hinterlegt, welche Webseite auf welche Website verlinkt (vereinfacht: eine Eins in Zeile i und Spalte j heißt, dass Webseite i auf Webseite j verlinkt; eine Null steht für "keine Verlinkung"). Je mehr (wichtige) Webseiten auf eine gewisse Webseite verlinken, desto wichtiger ist die Seite selbst.

Die obige Gleichung sagt aus, dass der gesuchte Gewichtsvektor ein Eigenvektor der Link-Matrix A zum Eigenwert 1 ist. Die ganze Frage des Page-Rankings, also der Sortierung von Suchergebnissen nach Wichtigkeit, ist aus mathematischer Sicht also “nur” ein bisschen Matrix-Vektor-Rechnung.

6 Differentialrechnung

Grundlagen der Differentialrechnung für Funktionen einer Veränderlichen und Anwendung in der Fehlerrechnung. Weitere Anwendungen folgen im zweiten Semester.

6.1 Idee und Begriffe

Betrachten in diesem Kapitel ausschließlich Funktionen $f : (a, b) \rightarrow \mathbb{R}$, die auf einem offenen Intervall (a, b) definiert sind. Dabei lassen wir auch $(-\infty, \infty)$, also ganz $\mathbb{R}$, als Definitionsbereich zu. Der Vorteil eines offenen Intervalls als Definitionsbereich liegt darin, dass $x \in (a, b)$ nie am Rand des Intervalls liegt, sondern stets auf beiden Seiten der Stelle x noch ein Stückchen Definitionsbereich vorhanden ist. Man sagt auch: alle Punkte eines offenen Intervalls sind innere Punkte.

Mit dem Begriff der Stetigkeit kennen wir bereits eine Möglichkeit, Eigenschaften einer Funktion in der Umgebung einer Stelle $x \in (a, b)$ als Eigenschaft an dieser Stelle zu formulieren. Wir erinnern uns: Die Frage, ob eine Funktion an einer Stelle x stetig ist, kann nur beantwortet werden, wenn man auch die Funktionswerte in der Nähe von x anschaut. Konkret geschieht dies über den Begriff des Grenzwerts.

Die Information, die uns der Stetigkeitsbegriff liefert (Existenz von Sprung- oder Polstellen), ist für manche Einsatzzwecke etwas dürftig. Deshalb werden wir die Idee des Formulierens von Umgebungseigenschaften als Eigenschaft an einer konkreten Stelle nun für die Frage umsetzen, ob die Funktion an der betrachteten Stelle “glatt” ist, also dort keinen “Knick” hat.

Wir führen zunächst die nötigen Begriffe ein. Erst am Ende dieses Kapitels folgt dann die Einordnung in größere Zusammenhänge, also die Antwort auf die Frage “Wozu brauche ich das?”.

Differenzierbarkeit an einer Stelle

Die Eigenschaft, die wir eben noch als “glatt” oder “ohne Knick” beschrieben haben, heißt Differenzierbarkeit.

Merke

Eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ heißt **differenzierbar an einer Stelle** $x_0 \in (a, b)$, wenn der Grenzwert

$$\lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} \quad (6.1)$$

existiert.

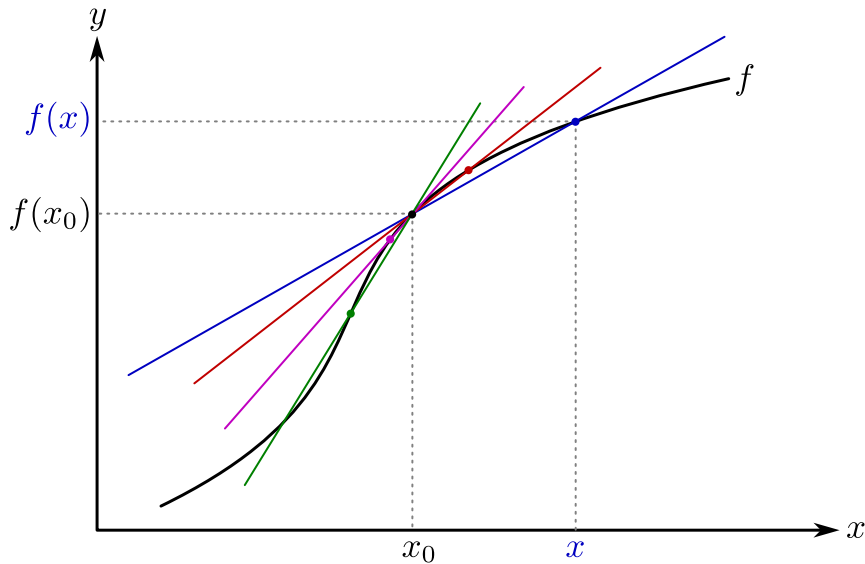
Der Ausdruck $\frac{f(x) - f(x_0)}{x - x_0}$ ist der Anstieg der Gerade durch die beiden Punkte $(x, f(x))$ und $(x_0, f(x_0))$. Der Grenzwert existiert bekanntlich genau dann, wenn die einseitigen Grenzwerte

$$\lim_{x \rightarrow x_0 - 0} \frac{f(x) - f(x_0)}{x - x_0} \quad \text{und} \quad \lim_{x \rightarrow x_0 + 0} \frac{f(x) - f(x_0)}{x - x_0} \quad (6.2)$$

existieren und übereinstimmen. Differenzierbarkeit bedeutet anschaulich also, dass die Steigung der Funktion links und rechts der betrachteten Stelle gleich ist (und natürlich

im Sinne des Grenzwerts überhaupt erst einmal existiert). Beispiele und Gegenbeispiele folgen weiter unten.

Der Ausdruck $\frac{f(x)-f(x_0)}{x-x_0}$ wird auch als **Differenzenquotient** bezeichnet, der zugehörige Grenzwert als **Differentialquotient**.



Gezeigt sind einige der Geraden, deren Anstieg beim Grenzübergang in der Definition der Differenzierbarkeit eine Rolle spielt.

Merke

Ist f in x_0 differenzierbar, so ist f dort auch stetig.

Differenzierbarkeit kann man somit als eine strengere Form von Stetigkeit ansehen. Es wird sich aber zeigen, dass das zu kurz gedacht ist. Der Begriff der Differenzierbarkeit eröffnet ganz neue Möglichkeiten und ist in der Ingenieurmathematik unverzichtbares Werkzeug sowohl für die Modellierung als auch für das Lösen praktischer Aufgabenstellungen (siehe unten).

Dass aus der Differenzierbarkeit die Stetigkeit folgt, sieht man leicht ein: Offensichtlich gilt

$$|f(x) - f(x_0)| = |x - x_0| \frac{|f(x) - f(x_0)|}{|x - x_0|}. \quad (6.3)$$

Der Faktor $|x - x_0|$ geht gegen Null, wenn $x \rightarrow x_0$. Der zweite Faktor konvergiert gegen eine feste Zahl. Folglich konvergiert das Produkt gegen Null, also $f(x) \rightarrow f(x_0)$.

Nebenbemerkung

Beachte, dass das Erweitern mit $|x - x_0|$ uns in die unschöne Situation bringt, dass wir eventuell durch Null teilen (bei $x = x_0$). Diesen Fall müsste man bei ganz strengem Vorgehen ausführlicher diskutieren. Das ist aber ein lösbares, eher technisches Problem, welches uns keine neuen Erkenntnisse bringt. Wir ersparen uns hier die Details und glauben den Mathematikern, dass hier alles seine Ordnung hat.

Ableitung

Analog zur Stetigkeit sind insbesondere solche Funktionen von Interesse, die an jeder Stelle differenzierbar sind.

Merke

Eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ heißt **differenzierbar**, wenn sie an jeder Stelle ihres Definitionsbereichs differenzierbar ist.

Ist f differenzierbar, so ist die Funktion

$$x \mapsto \lim_{\tilde{x} \rightarrow x} \frac{f(\tilde{x}) - f(x)}{\tilde{x} - x} \quad (6.4)$$

wohldefiniert auf (a, b) . Diese Funktion heißt **Ableitung** von f und wird mit f' bezeichnet. Für die Funktionswerte $f'(x)$ schreibt man manchmal auch $\frac{df(x)}{dx}$ oder $\frac{df}{dx}(x)$.

In der Ingenieurmathematik hält sich im Umfeld von Differenzierbarkeit und Ableitungen hartnäckig der Begriff **Differential**. Meist wird dieser nur sehr diffus als “unendlich kleine Zahl, die nicht Null ist” oder als “infinitesimal kleine Zahl” eingeführt, was offensichtlich Unfug ist. Der Begriff stammt aus der Zeit, als die Differentialrechnung noch nicht vollständig ausgearbeitet und verstanden war, also aus der Zeit vor 1900. Das in der Ingenieurmathematik leider nach wie vor verbreitete “Rechnen mit Differentialen” hat wenig mit mathematischer Exaktheit zu tun, sondern ist mehr als “intuitives Formelbasteln” anzusehen. Tatsächlich gibt es auch in der modernen Mathematik den Differential-Begriff. Der ist seit über 100 Jahren präzise und sinnvoll definiert, hat für uns aber keinerlei Bedeutung oder Nutzen. Deshalb meiden wir diesen Begriff und seine recht sperrige Definition hier.

Merke

Eine differenzierbare Funktion heißt **stetig differenzierbar**, wenn ihre Ableitung

stetig ist.

Beachte: Differenzierbare Funktionen sind immer stetig, aber die Ableitung kann durchaus unstetig sein (vgl. Beispiele unten).

Beispiele

Das praktische Berechnen von Ableitungen werden wir im nächsten Kapitel betrachten. Das Vorgehen wird deutlich bequemer sein als es hier zunächst den Anschein hat. Insbesondere werden wir ohne das Bilden von Grenzwerten auskommen. Hier zunächst ein paar elementare Beispiele zur Illustration der Technik.

Beispiel (konstante Funktion)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := c \quad \text{für ein } c \in \mathbb{R} \quad (6.5)$$

ist differenzierbar wegen

$$\lim_{\tilde{x} \rightarrow x} \frac{f(\tilde{x}) - f(x)}{\tilde{x} - x} = \lim_{\tilde{x} \rightarrow x} \frac{c - c}{\tilde{x} - x} = \lim_{\tilde{x} \rightarrow x} 0 = 0. \quad (6.6)$$

Die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist entsprechend gegeben durch

$$f'(x) = 0. \quad (6.7)$$

Beispiel (lineare Funktion)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := cx \quad \text{für ein } c \in \mathbb{R} \quad (6.8)$$

ist differenzierbar wegen

$$\lim_{\tilde{x} \rightarrow x} \frac{f(\tilde{x}) - f(x)}{\tilde{x} - x} = \lim_{\tilde{x} \rightarrow x} \frac{c\tilde{x} - cx}{\tilde{x} - x} = \lim_{\tilde{x} \rightarrow x} c = c. \quad (6.9)$$

Die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist entsprechend gegeben durch

$$f'(x) = c. \quad (6.10)$$

Beispiel (quadratische Funktion)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := x^2 \quad (6.11)$$

ist differenzierbar wegen

$$\begin{aligned} \lim_{\tilde{x} \rightarrow x} \frac{f(\tilde{x}) - f(x)}{\tilde{x} - x} &= \lim_{\tilde{x} \rightarrow x} \frac{\tilde{x}^2 - x^2}{\tilde{x} - x} \\ &= \lim_{\tilde{x} \rightarrow x} \frac{(\tilde{x} + x)(\tilde{x} - x)}{\tilde{x} - x} \\ &= \lim_{\tilde{x} \rightarrow x} (\tilde{x} + x) = 2x. \end{aligned} \quad (6.12)$$

Die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist entsprechend gegeben durch

$$f'(x) = 2x. \quad (6.13)$$

Beispiel (Betragsfunktion)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := |x| \quad (6.14)$$

ist für $x \neq 0$ differenzierbar mit Ableitung

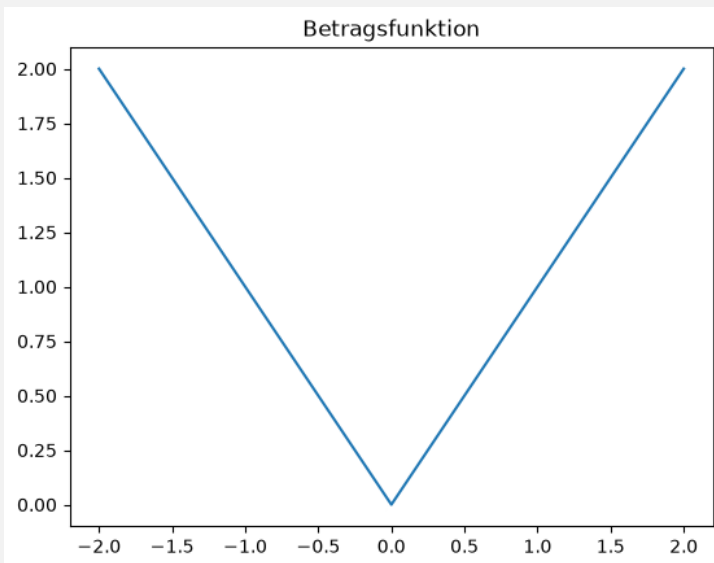
$$f'(x) = \begin{cases} -1, & \text{für } x < 0, \\ 1, & \text{für } x > 0 \end{cases} \quad (6.15)$$

(vgl. obiges Beispiel zu linearen Funktionen).

An der Stelle $x = 0$ ist die Funktion nicht differenzierbar denn links- und rechtsseitiger Grenzwert existieren zwar, stimmen aber nicht überein:

$$\begin{aligned} \lim_{\tilde{x} \rightarrow 0-0} \frac{f(\tilde{x}) - f(0)}{\tilde{x} - 0} &= \lim_{\tilde{x} \rightarrow 0-0} \frac{-\tilde{x}}{\tilde{x}} = \lim_{\tilde{x} \rightarrow 0-0} -1 = -1, \\ \lim_{\tilde{x} \rightarrow 0+0} \frac{f(\tilde{x}) - f(0)}{\tilde{x} - 0} &= \lim_{\tilde{x} \rightarrow 0+0} \frac{\tilde{x}}{\tilde{x}} = \lim_{\tilde{x} \rightarrow 0+0} 1 = 1. \end{aligned} \quad (6.16)$$

Man sagt auch: die Funktion hat einen Knick bei $x = 0$.



Beispiel (Wurzelfunktion)

Man kann zeigen (tun wir hier nicht), dass die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \sqrt{|x|} \quad (6.17)$$

für $x \neq 0$ differenzierbar ist und die Ableitung $f' : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$ durch

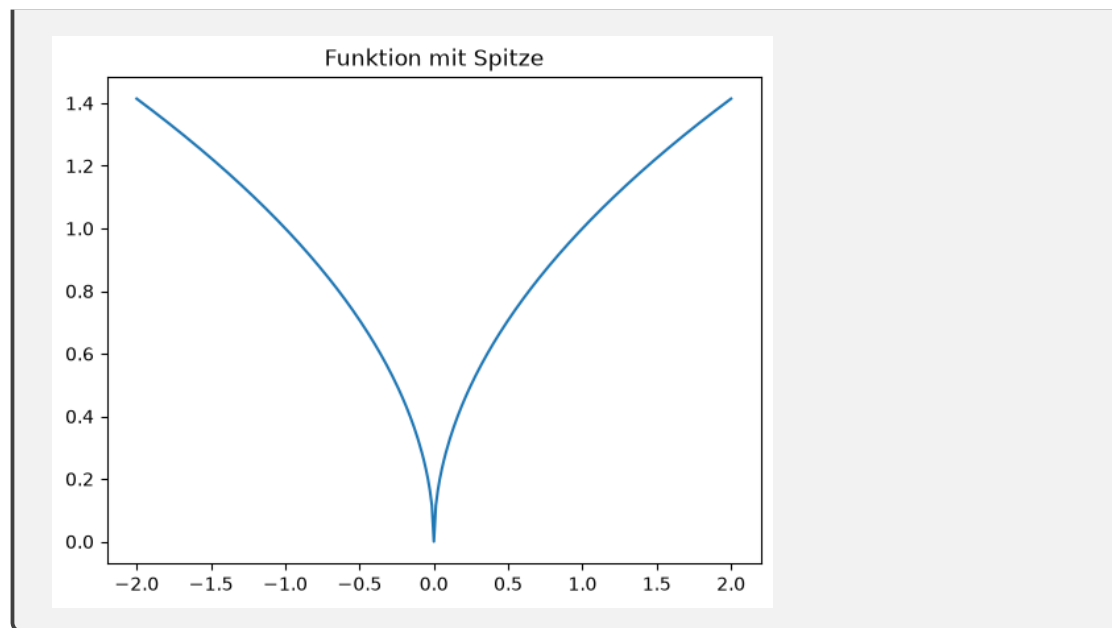
$$f'(x) = \frac{\sqrt{|x|}}{2x} \quad (6.18)$$

gegeben ist.

An der Stelle $x = 0$ ist die Funktion nicht differenzierbar denn links- und rechtsseitiger Grenzwert existieren nur im Sinne von $\pm\infty$ und stimmen nicht überein:

$$\begin{aligned} \lim_{\tilde{x} \rightarrow 0-0} \frac{f(\tilde{x}) - f(0)}{\tilde{x} - 0} &= \lim_{\tilde{x} \rightarrow 0-0} \frac{\sqrt{-\tilde{x}}}{\tilde{x}} = \lim_{\tilde{x} \rightarrow 0-0} \frac{1}{-\sqrt{-\tilde{x}}} = -\infty, \\ \lim_{\tilde{x} \rightarrow 0+0} \frac{f(\tilde{x}) - f(0)}{\tilde{x} - 0} &= \lim_{\tilde{x} \rightarrow 0+0} \frac{\sqrt{\tilde{x}}}{\tilde{x}} = \lim_{\tilde{x} \rightarrow 0+0} \frac{1}{\sqrt{\tilde{x}}} = \infty. \end{aligned} \quad (6.19)$$

Man sagt auch: die Funktion hat eine Spitze bei $x = 0$.

**Beispiel (unstetige Ableitung)**

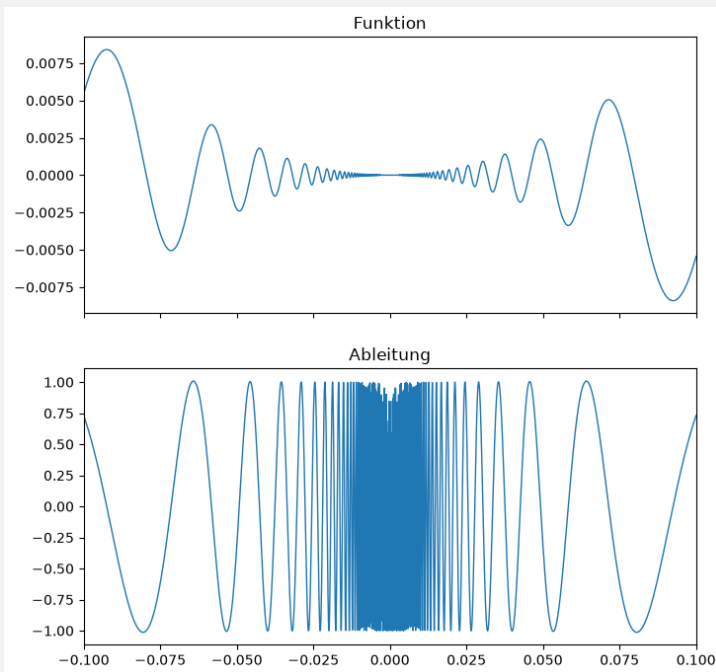
Man kann zeigen (tun wir hier aber nicht), dass die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \begin{cases} x^2 \sin \frac{1}{x}, & \text{für } x \neq 0, \\ 0, & \text{für } x = 0 \end{cases} \quad (6.20)$$

überall differenzierbar ist, die Ableitung

$$f'(x) = \begin{cases} 2x \sin \frac{1}{x} - \cos \frac{1}{x}, & \text{für } x \neq 0, \\ 0, & \text{für } x = 0 \end{cases} \quad (6.21)$$

aber in $x = 0$ nicht stetig ist.



Die Menge der differenzierbaren Funktionen ist also eine echte Obermenge der Menge der stetig differenzierbaren Funktionen.

Nebenbemerkung (nirgends differenzierbare stetige Funktionen)

Da man in der Praxis fast ausschließlich differenzierbare Funktionen antrifft oder wenigstens solche, die nur an einzelnen Stellen nicht differenzierbar sind, könnte der Eindruck entstehen, dass die Mehrheit der stetigen Funktionen fast überall differenzierbar ist. Dieser Schluss ist jedoch falsch! Betrachtet man beispielsweise die Menge aller stetigen Funktionen, die von $(0, 1)$ nach $\mathbb{R}$ abbilden, so zeigt sich, dass die große Mehrheit nirgends (!) differenzierbar ist. Unser Blick auf die Welt der Funktionen ist so eng begrenzt, dass wir große Mühe haben, wenigstens ein Beispiel dieser großen Mehrheit von stetigen, aber nirgends differenzierbaren Funktionen zu finden. Diese lassen sich nur bedingt in Formeln aufschreiben. Ein Beispiel ist die Weierstraß-Funktion. In dieser Quelle findet man auch die präzise Formulierung für “große Mehrheit”; diese benötigt allerdings Begriffe, die wir bisher nicht eingeführt haben und auch nicht einführen werden.

Die Frage, ob die Natur differenzierbare Funktionen bevorzugt oder ob nicht der Mensch beim Modellieren diese künstlich als Vereinfachung der tatsächlichen Verhältnisse hinzufügt, ist bedenkenswert!

Höhere Ableitungen

Da die Ableitung einer differenzierbaren Funktion selbst wieder eine Funktion ist, kann man auch für die Ableitung die Frage der Differenzierbarkeit stellen.

Merke

Eine Funktion f heißt **k -mal differenzierbar**, wenn $(k - 1)$ -mal die Ableitung der Ableitung gebildet werden kann, also insbesondere die (erste) Ableitung wieder differenzierbar ist, die Ableitung der Ableitung differenzierbar ist usw. Die Ableitungen werden mit $f', f'', f''', f^{(4)}, \dots, f^{(k)}$ bezeichnet. Offensichtlich gilt

$$f^{(l)} = (f^{(l-1)})' \quad \text{für } l = 1, \dots, l \quad (6.22)$$

(wobei man mit $f^{(0)}$ üblicherweise die Funktion f selbst meint).

Merke

Eine Funktion heißt **k -mal stetig differenzierbar**, wenn sie k -mal differenzierbar ist und die k -te Ableitung stetig ist.

Interpretation

Aus der Definition der Ableitung ist sofort klar, dass die Ableitung einer Funktion an einer Stelle x die **Steigung der Funktion** an dieser Stelle angibt. Bemerkenswert ist jedoch, dass der Begriff “Steigung” eigentlich nur für Geraden sinnvoll ist, also mindestens zwei Punkte benötigt, zwischen denen die Steigung, also das Verhältnis von vertikaler und horizontaler Verschiebung, betrachtet wird. Durch den Grenzprozess beim Bilden der Ableitung wird die Steigung der Funktion in der Umgebung der betrachteten Stelle x als Eigenschaft an dieser Stelle x ausgedrückt. Dies ist wieder ein Beispiel für starke Abstraktion, die durch Gewöhnung mit der Zeit in die ganz normale, scheinbar intuitiv klare Begriffswelt übergeht.

Die zweite Ableitung einer Funktion gibt die Änderungsrate des Anstiegs an einer Stelle an. Auch hier ist “Änderungsrate” eigentlich nur sinnvoll erklärbar, wenn man eine Umgebung einer Stelle anschaut und nicht nur die Stelle selbst. Durch den Grenzprozess wird die relevante Information von der Umgebung einer Stelle auf die Stelle selbst übertragen. Die Änderungsrate der Steigung kann man als **Krümmung des Funktionsgraphs** interpretieren. Wächst der Anstieg, so krümmt sich der Graph nach links (wenn man ihn gedanklich von links nach rechts durchläuft); sinkt der Anstieg, so krümmt sich der Graph nach rechts. Funktionen, deren zweite Ableitung überall positiv ist, heißen **konvex**; ist die zweite Ableitung überall negativ, so heißt die Funktion **konkav**.

Die dritte Ableitung liefert Informationen darüber, wie sich die Krümmung der Funktion

ändert. Ist die dritte Ableitung in einem Bereich positiv, so nimmt in diesem Bereich die Krümmung zu, d.h. die Kurve, die der Funktionsgraph beschreibt, wird enger; bei negativer dritter Ableitung wird sie weiter.

Ab der vierten Ableitung wird die geometrische Interpretation schwieriger. Auch ist eine geometrische Interpretation hier kaum relevant.

Einordnung

Sofern in der Schulmathematik bereits über Differenzierbarkeit und Ableitungen gesprochen wurde, könnte der Eindruck entstanden sein, dass die Hauptzwecke dieser Begriffe das Untersuchen von Funktionen ("Kurvendiskussion") und das Lösen von Optimierungsproblemen sind. Auch in der Grundlagenmathematik im Studium könnte dieser Eindruck entstehen. Um diesen Trugschluss zu vermeiden, seien hier ein paar Hinweise zu Bedeutung und Entstehung des Differenzierbarkeitsbegriffs gegeben.

Die Väter der Differentialrechnungen waren Newton und Leibniz (beide 17. Jahrhundert). Man wollte nach dem "dunklen", durch Spiritualität und Unterdrückung von Wissenschaft geprägten Mittelalter die Welt verstehen und in Formeln beschreiben, hatte aber bis auf Zahlen, Funktionen und geometrische Figuren keine Werkzeuge dafür. Insbesondere Gesetzmäßigkeiten bei Bewegungsabläufen (freier Fall, Planetenbewegungen,...) ließen sich damit nicht zufriedenstellend formulieren. Erst die für diesen Zweck entwickelte Differentialrechnung erlaubte das Formulieren der bis heute (im Wesentlichen unverändert) im Einsatz befindlichen klassischen Mechanik. In diesen Bereich fällt insbesondere das allgemein bekannte Newtonsche Gesetz:

$$\text{Kraft} = \text{Masse} \cdot \text{Beschleunigung}.$$

Bewegt sich ein Körper mit Masse m im Raum entlang einer durch

$$t \mapsto (x(t), y(t), z(t)), \quad t \in [a, b], \quad (6.23)$$

beschriebenen Kurve unter Einwirkung einer zeitlich und örtlich variierenden Kraft

$$F : \mathbb{R} \times \mathbb{R}^3 \rightarrow \mathbb{R} \quad (6.24)$$

(jeder Zeit und jedem Ort wird ein Kraftvektor zugeordnet), so kann das Newtonsche Gesetz mittels Ableitungen der Koordinatenfunktionen als

$$F(t, (x(t), y(t), z(t))) = m \begin{bmatrix} x''(t) \\ y''(t) \\ z''(t) \end{bmatrix} \quad \text{für alle } t \in (a, b) \quad (6.25)$$

formuliert werden. Newton hatte erkannt, dass dieser Zusammenhang stets gilt. Kennt man die wirkenden Kräfte und die Masse m , so kann man aus dieser Gleichung die Bahn berechnen, auf der sich der Körper bewegen wird.

Die gezeigte Form des Newtonschen Gesetzes ist ein **System von gewöhnlichen Differentialgleichungen**. Differentialgleichungen sind Gleichungen, die eine Funktion in Beziehung zu ihren Ableitungen setzt. Das Attribut “gewöhnlich” wird genutzt, wenn es sich um Funktionen einer Veränderlichen handelt. Da beim Newtonschen Gesetz vektorwertige Funktionen links und rechts stehen, handelt es sich eigentlich um drei Differentialgleichungen, also ein Differentialgleichungssystem.

Später werden wir die Differentialrechnung auf Funktionen mehrerer Veränderlicher ausdehnen. Die damit formulierbaren Differentialgleichungen werden als **partielle Differentialgleichungen** bezeichnet.

In unserer Veranstaltung spielen Differentialgleichungen keine Rolle. Dennoch bilden diese den Haupteinsatzzweck für die Differentialrechnung. Mit Differentialgleichungen lassen sich nahezu alle Phänomene aus Naturwissenschaften, Technik und auch Wirtschaftswissenschaften beschreiben: Planetenbewegungen, Wachstumsprozesse, Wärmeleitprozesse, Strömungen, die Entwicklung von Aktienkursen usw.

Eher nebenbei ist die Differentialrechnung sehr nützlich für das Lösen von Optimierungsproblemen, sowohl beim manuellen Lösen als auch beim Lösen mit dem Computer. Beispielsweise basieren die heute üblichen Verfahren zum Training von KI ganz wesentlich auf Resultaten der Differentialrechnung (wir kommen im zweiten Semester darauf zurück).

6.2 Vorgehen

Vor der flächendeckenden Einführung des Computers mussten Ableitungen per Hand berechnet werden. Entsprechend umfangreich und weit entwickelt sind die Verfahren dazu. Wir werden hier nur einen Überblick über das prinzipielle Vorgehen geben und, da heute deutlich relevanter, das Differenzieren mit dem Computer betrachten.

Manuelles Differenzieren

Der Weg über die Definition der Ableitung einer Funktion als Grenzwert ist recht mühsam für den praktischen Einsatz. Analog zur Berechnung von Grenzwerten bei Zahlenfolgen und Funktionen entnimmt man deshalb die Ableitungen einiger Grundfunktionen aus Tabellen und kombiniert diese dann anhand verschiedener Ableitungsregeln zu Ableitungen komplexerer Funktionen. Zunächst die Ableitungen einiger immer wieder auftretender Funktionen (gültig für alle x , für die f und f' sinnvolle Werte liefern):

$$f(x) = x^c, \quad c \in \mathbb{R} \quad \Rightarrow \quad f'(x) = c x^{c-1}, \quad (6.26)$$

$$f(x) = \cos x \quad \Rightarrow \quad f'(x) = -\sin x, \quad (6.27)$$

$$f(x) = \sin x \quad \Rightarrow \quad f'(x) = \cos x, \quad (6.28)$$

$$f(x) = e^x \quad \Rightarrow \quad f'(x) = e^x, \quad (6.29)$$

$$f(x) = \ln x \quad \Rightarrow \quad f'(x) = \frac{1}{x}. \quad (6.30)$$

Beachte, dass die Exponentialfunktion beim Ableiten unverändert bleibt! Dinge, die trotz des Anwendens irgendwelcher Operationen unverändert bleiben, spielen stets eine besondere Rolle. So taucht die Exponentialfunktion beispielsweise sehr oft in den Lösungen von Differentialgleichungen auf, spielt aber auch in der Wahrscheinlichkeitsrechnung eine wichtige Rolle.

Mittels folgender Ableitungsregeln, deren Gültigkeit wir hier nicht nachweisen, können Ableitungen komplexerer Funktionen auf Ableitungen einfacher Funktionen zurückgeführt werden.

- Ist f an einer Stelle x differenzierbar und ist $\lambda \in \mathbb{R}$ eine Zahl, so ist auch λf bei x differenzierbar und es gilt

$$(\lambda f)'(x) = \lambda f'(x). \quad (6.31)$$

- Sind f und g an einer Stelle x differenzierbar, so ist auch die Summe $f + g$ dort differenzierbar und es gilt

$$(f + g)'(x) = f'(x) + g'(x) \quad (6.32)$$

(Summenregel).

- Sind f und g an einer Stelle x differenzierbar, so ist auch das Produkt $f g$ dort differenzierbar und es gilt

$$(f g)'(x) = f'(x) g(x) + f(x) g'(x) \quad (6.33)$$

(**Produktregel**).

- Sind f und g an einer Stelle x differenzierbar und gilt $g(x) \neq 0$, so ist auch der Quotient $\frac{f}{g}$ bei x differenzierbar und es gilt

$$\left(\frac{f}{g}\right)'(x) = \frac{f'(x) g(x) - f(x) g'(x)}{g(x)^2} \quad (6.34)$$

(**Quotientenregel**).

- Ist f an einer Stelle x differenzierbar und ist g an der Stelle $f(x)$ differenzierbar, so ist auch die Hintereinanderausführung $g \circ f$ bei x differenzierbar und es gilt

$$(g \circ f)'(x) = f'(x) g'(f(x)) \quad (6.35)$$

(**Kettenregel**).

Beispiel (Ableitung eines Polynoms)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := 2x^3 + 5x + 1 \quad (6.36)$$

ist überall differenzierbar und die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist durch

$$f'(x) = 6x^2 + 5 \quad (6.37)$$

gegeben (Video 610, Notizen dazu).

Beispiel (Ableitung eines Produkts)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := x^2 \cos x \quad (6.38)$$

ist überall differenzierbar und die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist durch

$$f'(x) = 2x \cos x - x^2 \sin x \quad (6.39)$$

gegeben (Video 615, Notizen dazu).

Beispiel (Ableiten mittels Kettenregel)

Die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \sin e^{3x} \quad (6.40)$$

ist überall differenzierbar und die Ableitung $f' : \mathbb{R} \rightarrow \mathbb{R}$ ist durch

$$f'(x) = 3 e^{3x} \cos e^{3x} \quad (6.41)$$

gegeben (Video 620, Notizen dazu).

Ableitungen mit SymPy

Das Bilden von Ableitungen mit SymPy ist denkbar einfach:

```
import sympy

x, c = sympy.symbols('x c')

sympy.diff((x + c) ** 2 * sympy.sin(x), x)
```

```
(c + x)**2*cos(x) + (2*c + 2*x)*sin(x)
```

Numerische Differentiation

Das numerische Differenzieren erfolgt direkt über die Definition der Ableitung:

$$f'(x) = \lim_{\tilde{x} \rightarrow x} \frac{f(\tilde{x}) - f(x)}{\tilde{x} - x}. \quad (6.42)$$

Da man Grenzübergänge mit dem Computer nicht exakt umsetzen kann, wählt man $\tilde{x}$ nah bei x und betrachtet den zugehörigen Differenzenquotient als Näherung für die Ableitung. Für $\tilde{x} := x + h$ mit einem fest (also unabhängig von x) gewähltem, kleinen $h > 0$ erhält man also

$$f'(x) \approx \frac{f(x + h) - f(x)}{h}. \quad (6.43)$$

Diese Form der numerischen Differentiation bezeichnet man als **Vorwärtsdifferenz**. Da $\tilde{x}$ auch kleiner als x sein kann, ist

$$f'(x) \approx \frac{f(x - h) - f(x)}{-h} = \frac{f(x) - f(x - h)}{h}. \quad (6.44)$$

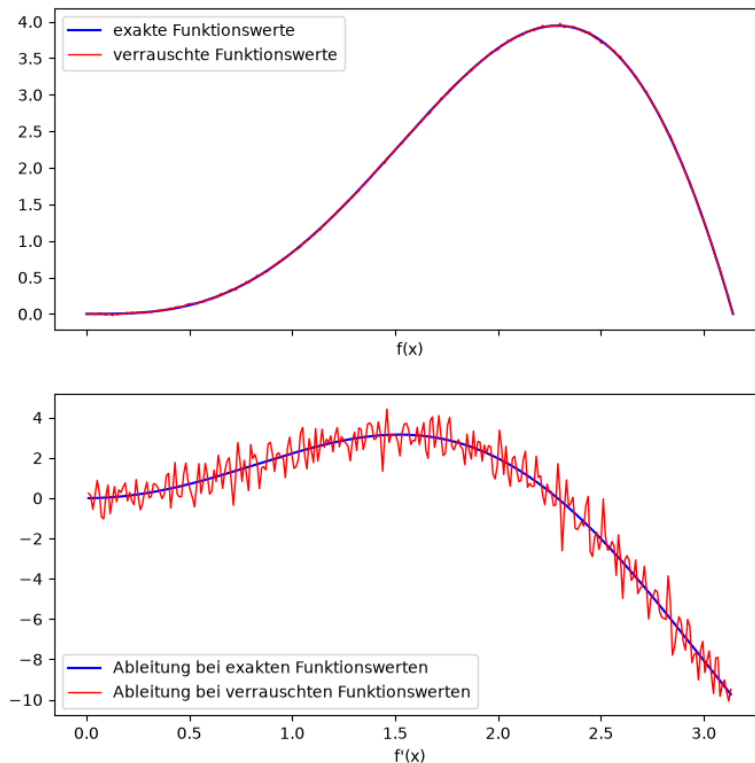
eine völlig gleichberechtigte, alternative Näherungsformel für die Ableitung (**Rückwärtsdifferenz**). Gibt es keinen Grund eine von beiden Varianten zu bevorzugen, so nutzt man gern deren

Mittelwert:

$$\begin{aligned} f'(x) &\approx \frac{1}{2} \frac{f(x+h) - f(x)}{h} + \frac{1}{2} \frac{f(x) - f(x-h)}{h} \\ &= \frac{f(x+h) - f(x-h)}{2h} \end{aligned} \quad (6.45)$$

(**zentrale Differenz**).

Auf den ersten Blick wird die Näherung für $f'(x)$ besser, je kleiner h gewählt wird. Jedoch zeigt sich in der Praxis, dass zu kleines h zu starken Oszillationen in der berechneten Näherung für f' führen kann. Ursache ist, dass kleine Fehler in den (oft aus Messungen stammenden) Werten für f durch das numerische Differenzieren verstärkt werden. Die Details dazu sind Thema der numerischen Mathematik und führen hier zu weit. Merke jedoch: **Numerisches Differenzieren ist nur mit großer Vorsicht zu benutzen!**



Im Bild sehen wir: Obwohl die verrauschten Funktionswerte kaum von den exakten Funktionswerten unterscheidbar sind, weichen die numerisch berechneten Ableitungen deutlich voneinander ab. Zum selbst ausprobieren: Erhöht man die Anzahl der Stützstellen, so wird die Abweichung der Ableitungen immer größer.

6.3 Taylor-Polynome

Als eine erste konkrete Anwendung der Differentialrechnung schauen wir uns hier die Annäherung von Funktionen durch Polynome an. Diese Technik hilft uns im nächsten Kapitel beim Verständnis der Fehlerrechnung. Weitere Anwendungen folgen im zweiten Semester.

Bedeutung

Es gibt verschiedene Verfahren zum Ersetzen (nahezu) beliebiger Funktionen durch Polynome mit sehr ähnlichem Verlauf des Funktionsgraphs. Das einfachste ist die Konstruktion sogenannter Taylor-Polynome.

Während man früher durch den Übergang zu Polynomen die manuelle Rechenarbeit vereinfachen wollte, liegt die Bedeutung der Approximation (d.h. Annäherung) von Funktionen durch Polynome heute vor allem in der Verarbeitung von Funktionen mit dem Computer. Während allgemeine Funktionen (z.B. Sinus, Exponentialfunktion, Logarithmusfunktion,...) beim numerischen Rechnen mit dem Computer nur als Sammlung endlich vieler Funktionswerte an gegebenen Stützstellen darstellbar sind, also nie exakt und vollständig vorliegen, können Polynome auch mit dem Computer exakt und vollständig beschrieben werden. Wir erinnern uns: Ein Polynom

$$a_n x^n + a_{n-1} x^{n-1} + \cdots + a_2 x^2 + a_1 x + a_0 \quad (6.46)$$

vom Grad n ist durch die $n+1$ Koeffizienten $a_0, a_1, \dots, a_n$ beschrieben. Das sind endlich viele Zahlen und diese sind mit dem Computer (bis auf Rundungsfehler) darstellbar und verarbeitbar.

Wir werden Taylor-Polynome nur in sehr einfacher Form im nächsten Kapitel wiedersehen. Wann immer Sie jedoch mit dem Computer oder dem Taschenrechner beispielsweise den Sinus eines Winkels berechnen, werden im Hintergrund Taylor-Polynome (oder sehr ähnliche Techniken) verwendet.

Konstruktion

Sei $f : (a, b) \rightarrow \mathbb{R}$ eine Funktion, wobei auch ein- und beidseitig unendliche Intervalle als Definitionsbereich zulässig seien, und sei n eine gegebene natürliche Zahl. Gesucht ist ein Polynom $T_n : \mathbb{R} \rightarrow \mathbb{R}$ vom Grad n , welches einen ähnlichen Verlauf zeigt wie die Funktion f . Ein Polynom vom Grad n besitzt $n+1$ Freiheitsgrade (d.h. frei wählbare Parameter). Wir können also davon ausgehen, dass wir $n+1$ Bedingungen formulieren können, welche den Verlauf des Polynoms T_n festlegen. Zu einer gegebenen/gewählten Stelle $x_0 \in (a, b)$ fordern wir

$$T_n(x_0) = f(x_0), \quad T'_n(x_0) = f'(x_0), \quad \dots, T_n^{(n)}(x_0) = f^{(n)}(x_0). \quad (6.47)$$

Klar ist, dass diese Forderungen nur formulierbar sind, wenn f wenigstens n -mal differenzierbar ist.

Die Ableitungen von T_n kann man explizit angeben. Man erhält so ein lineares Gleichungssystem für die $n + 1$ Unbekannten $a_0, \dots, a_n$ mit $n + 1$ Gleichungen. Die Systemmatrix ist sogar schon in Stufenform, sodass das Lösen dieses LGS im Prinzip keine Probleme macht. Allerdings ist das Rückwärtseinsetzen recht fummelig und unterbleibt deshalb an dieser Stelle. Man erhält folgendes Ergebnis:

Merke

Sei $f : (a, b) \rightarrow \mathbb{R}$ mindestens n -mal differenzierbar und sei $x_0 \in (a, b)$. Dann heißt das Polynom

$$\begin{aligned} T_n(x) &:= \sum_{k=0}^n \frac{1}{k!} f^{(k)}(x_0) (x - x_0)^k \\ &= f(x_0) + f'(x_0) (x - x_0) + f''(x_0) (x - x_0)^2 + \dots + \frac{1}{n!} f^{(n)}(x_0) (x - x_0)^n \end{aligned} \quad (6.48)$$

Taylor-Polynom vom Grad n zur Funktion f mit der **Entwicklungsstelle** x_0 .

Das Taylor-Polynom vom Grad 1 wird auch als **Tangente** an die Funktion an der Stelle x_0 bezeichnet. Dies ist eine Gerade, die an der Stelle x_0 den gleichen Funktionswert und den gleichen Anstieg wie die Funktion f hat. Das Taylor-Polynom vom Grad 2 wird gelegentlich als **Schmiegeparabel** bezeichnet. Neben Funktionswert und Anstieg stimmt hier auch die Krümmung mit der von f überein.

Beispiel

Zur Vereinfachung von Berechnungen soll die Sinusfunktion durch ein geeignetes Taylor-Polynom ersetzt werden. Wir entscheiden uns für den Grad $n = 4$, da Polynome vierten Grades manuell noch einigermaßen handhabbar sind. Als Entwicklungsstelle wählen wir $x_0 := \frac{\pi}{2}$, da wir uns für den Sinus von Winkeln in $[0, \pi]$ interessieren (aus der Konstruktion der Taylor-Polynome kann man erwarten, dass die Näherung um so schlechter wird, je weiter wir von der Entwicklungsstelle entfernt sind).

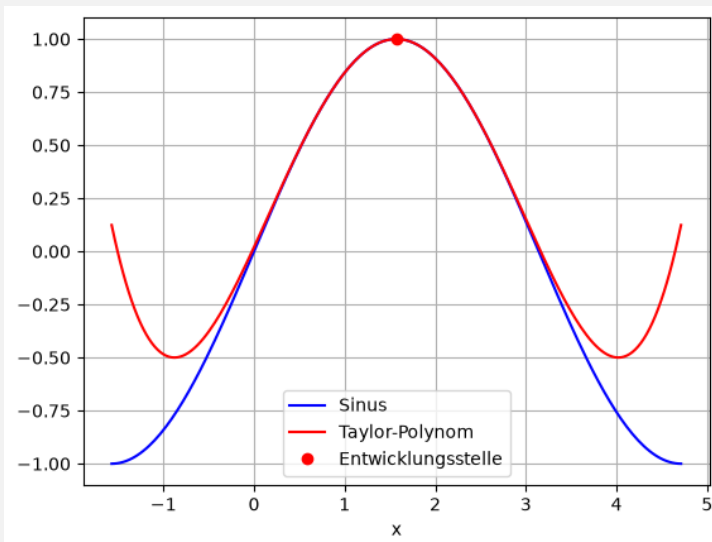
Die relevanten Ableitungen sind

$$\begin{aligned}
 f(x_0) &= \sin \frac{\pi}{2} = 1, \\
 f'(x_0) &= \cos \frac{\pi}{2} = 0, \\
 f''(x_0) &= -\sin \frac{\pi}{2} = -1, \\
 f'''(x_0) &= -\cos \frac{\pi}{2} = 0, \\
 f^{(4)}(x_0) &= \sin \frac{\pi}{2} = 1.
 \end{aligned} \tag{6.49}$$

Somit

$$T_4(x) = 1 - \frac{1}{2} \left(x - \frac{\pi}{2}\right)^2 + \frac{1}{24} \left(x - \frac{\pi}{2}\right)^4. \tag{6.50}$$

Diese Näherungsformel für die Sinusfunktion erlaubt das Berechnen von Sinuswerten unter Verwendung der vier Grundrechenarten.



Approximationsfehler

An der Entwicklungsstelle x_0 stimmen Funktion und Taylor-Polynom exakt überein. Mit wachsendem Abstand zur Entwicklungsstelle wird die Abweichung zwischen Funktion und Taylor-Polynom jedoch zunehmen. Frage: Können wir den Fehler, der durch das Ersetzen einer Funktion durch ein Taylor-Polynom entsteht, als konkrete Zahl angeben? Die Antwort ist ganz klar: Ja. Für das “Wie” gibt es verschiedenen Varianten. Wir geben hier nur eine an und verzichten auf die Herleitung.

Merke

Sei $f : (a, b) \rightarrow \mathbb{R}$ mindestens $(n + 1)$ -mal stetig differenzierbar und sei T_n das Taylor-Polynom n -ten Grades von f zur Entwicklungsstelle $x_0 \in (a, b)$. Dann existiert zu jedem $x \in (a, b)$ eine Zahl $\xi \in \mathbb{R}$, die zwischen x und x_0 liegt und

$$f(x) - T_n(x) = \frac{1}{(n+1)!} f^{(n+1)}(\xi) (x - x_0)^{n+1} \quad (6.51)$$

erfüllt.

Der Ausdruck auf der rechten Seite wird als **Lagrange-Restglied** oder Restglied in Lagrange-Form bezeichnet.

Insbesondere erhält man mit Hilfe des Lagrange-Restglieds die Fehlerabschätzung

$$|f(x) - T_n(x)| \leq \frac{1}{(n+1)!} \left(\max_{\xi \in (a,b)} |f^{(n+1)}(\xi)| \right) |x - x_0|^{n+1}. \quad (6.52)$$

Beachte hier, dass $f^{(n+1)}$ unbeschränkt sein kann, sodass das Maximum nicht existiert (bzw. ∞ ist).

Beispiel

Betrachten nochmal das Taylor-Polynom vierten Grades für die Sinusfunktion bei $x_0 = \frac{\pi}{2}$. Für $x \in [0, \pi]$ erhalten wir die Fehlerabschätzung

$$\begin{aligned} |\sin x - T_4(x)| &\leq \frac{1}{5!} \left(\max_{\xi \in [0, \pi]} |\cos \xi| \right) \left| x - \frac{\pi}{2} \right|^5 \\ &= \frac{1}{120} \left| x - \frac{\pi}{2} \right|^5 \leq \frac{1}{120} \left(\frac{\pi}{2} \right)^5 \leq 0.08, \end{aligned} \quad (6.53)$$

wobei die vorletzte Abschätzung aus $|x - \frac{\pi}{2}| \leq \frac{\pi}{2}$ folgt.

Somit stimmt bis zur ersten Nachkommastelle der mittels T_4 berechnete Näherungswert für $\sin x$ mit dem exakten Wert im Wesentlichen überein.

Ausblick

Wie eingangs erwähnt, gibt es diverse Möglichkeiten Polynome als Näherung für eine Funktion zu konstruieren. Taylor-Polynome konzentrieren sich auf eine Stelle und fordern dort die Gleichheit der Ableitungen bis zu einer gewählten Ordnung. Das andere Extrem sind **Interpolationspolynome**. Bei diesen fordert man die Gleichheit der Funktionswerte an mehreren Stellen. Möchte man ein Polynom vom Grad n konstruieren, fordert man also die Gleichheit der Funktionswerte an $n + 1$ Stellen.

Zwischen diesen beiden Extremen gibt es nahezu beliebige Mischformen. Eine in der Praxis sehr beliebte Variante sind sogenannte **Splines**. Grundidee ist, nicht ein Polynom für den gesamten Definitionsbereich zu bilden, sondern auf Teilintervallen zu arbeiten. Auf jedem Teilintervall wird ein eigenes Polynom (mit vorgegebenem Grad) konstruiert, sodass die einzelnen Polynome an den Intervallenden zusammenpassen, also gleiche Funktionswerte und gleiche Ableitungen wie die zu approximierende Funktion besitzen.

Für vektorwertige Funktionen und für Funktionen mehrerer Veränderlicher gibt es sehr ähnliche, auf Polynomen basierende Konstruktionen. Insbesondere sind die in der Computergrafik sehr beliebten Bézier-Kurven eine Form von Näherungspolynomen für Funktionen $f : [a, b] \rightarrow \mathbb{R}^2$.

6.4 Fehlerrechnung

Man steht sehr oft vor dem Problem, dass man eine interessierende Größe nicht direkt messen kann, sondern nur aus Berechnungen auf diese schließen kann. Ausgangspunkt für die nötigen Berechnungen sind gemessene Werte, die praktisch immer mit einem nur ungefähr bekannten Messfehler behaftet sind. Frage: Wie wirken sich solche Messfehler auf aus Messwerten berechnete Größen aus?

Beispiel

Um den Umfang eines Kreises (z.B. Radumfang bei einem Fahrzeug) zu ermitteln, bietet es sich an, den Durchmesser d des Rades zu messen und daraus den Umfang $u = \pi d$ zu berechnen. Liegt der Messfehler beim Durchmesser in etwa bei 1 mm, so wird der Fehler im berechneten Umfang aufgrund des Faktors π bei etwa 3 mm liegen.

Ziel

In obigem Beispiel ist der Zusammenhang von Messwert und berechnetem Wert sehr einfach und die **Fehlerfortpflanzung** beim Rechnen gut überschaubar. Wie aber bei komplexeren Zusammenhängen vorgehen? Wie wirkt sich beispielsweise die etwas ungenau gemessene Kantenlänge eines Würfels auf das daraus berechnete Volumen des Körpers aus?

Ziel dieses Kapitels ist die Entwicklung eines systematischen Vorgehens, welches uns praktisch verwertbare Aussagen über den zu erwartenden Fehler in einer aus Messwerten berechneten Größe liefert. Zunächst betrachten wir nur Größen, die von genau einem Messwert abhängen. Im zweiten Semester erweitern wir das Vorgehen auf von mehreren Messwerten abhängige Größen.

Es wird sich zeigen, dass die Differentialrechnung und insbesondere die Taylor-Polynome hier gute Dienste leisten.

Formalisierung

Die Aufgabenstellung ist also folgende: Gegeben ist eine Funktion f , die eine messtechnisch zugängliche Zahl $x_{\text{exakt}} \in \mathbb{R}$ in den gesuchten Wert $f(x_{\text{exakt}})$ überführt. Die praktische Messung liefert statt x_{exakt} allerdings nur x_{mess} . Entsprechend steht in der Praxis nur $f(x_{\text{mess}})$ zur Verfügung statt $f(x_{\text{exakt}})$.

Der exakte Wert x_{exakt} ist üblicherweise nahe bei x_{mess} (sonst ist der Messprozess unbrauchbar), muss aber nicht mit x_{mess} übereinstimmen. Wir nehmen an, dass wir eine **Fehlerschranke** $\Delta_{\max} x \in \mathbb{R}$ kennen, d.h. wir wissen, dass

$$|x_{\text{exakt}} - x_{\text{mess}}| \leq \Delta_{\max} x \quad (6.54)$$

gilt. Daraus soll nun eine Aussage über die Größe des Fehlers

$$|f(x_{\text{exakt}}) - f(x_{\text{mess}})| \quad (6.55)$$

abgeleitet werden. Da wir $f(x_{\text{exakt}})$ nicht kennen, müssen wir vom schlechtesten Fall ausgehen, d.h. wir suchen nach der maximalen Abweichung $\Delta_{\max} f$ bzgl. aller im Rahmen des maximalen Messfehlers $\Delta_{\max} x$ denkbaren Werte $f(x_{\text{exakt}})$. In jedem Fall wird also

$$|f(x_{\text{exakt}}) - f(x_{\text{mess}})| \leq \Delta_{\max} f \quad (6.56)$$

gelten.

Merke

Der Wert $\Delta_{\max} f$ wird als **absoluter Fehler** bezeichnet. Setzt man diesen ins Verhältnis zum berechneten Wert, so erhält man den **relativen Fehler** $\frac{\Delta_{\max} f}{|f(x_{\text{mess}})|}$.

Analog spricht man bei $\Delta_{\max} x$ vom absoluten Messfehler und bei $\frac{\Delta_{\max} x}{|x_{\text{mess}}|}$ vom relativen Messfehler.

Absolute Fehler schreibt man im Ingenieurwesen gern in der Form

$$f(x_{\text{exakt}}) = f(x_{\text{mess}}) \pm \Delta_{\max} f \quad \text{bzw.} \quad x_{\text{exakt}} = x_{\text{mess}} \pm \Delta_{\max} x, \quad (6.57)$$

meint damit aber stets, dass der exakte Wert höchstens (nicht genau!) um $\Delta_{\max} f$ bzw. $\Delta_{\max} x$ vom gemessenen Wert abweicht.

Relative Fehler gibt man gern in Prozent an, also beispielsweise 0.1 % statt 0.001.

Vorgehen

Theoretisch könnte man den zu erwartenden Fehler als

$$\Delta_{\max} f = \max_{x \in X} |f(x) - f(x_{\text{mess}})| \quad (6.58)$$

berechnen, wobei

$$X := [x_{\text{mess}} - \Delta_{\max} x, x_{\text{exakt}} + \Delta_{\max} x] \quad (6.59)$$

das Intervall der infrage kommenden exakten Werte ist. Dieses Maximum ist jedoch (auch mit dem Computer!) nur sehr mühsam berechenbar und damit für den praktischen Einsatz untauglich. Ersetzen wir f jedoch durch ein Taylor-Polynom, so wird die Sache handhabbar. Schon das Taylor-Polynom ersten Grades (Tangente) reicht in den meisten Fällen völlig aus. Als Entwicklungsstelle wählen wir natürlich x_{mess} , da dieser Wert einerseits ohnehin vorliegt und andererseits alle Betrachtungen in der nahen Umgebung dieser Stelle stattfinden, also dort, wo die Abweichung zwischen Funktion und dem so konstruierten Taylor-Polynom gering ist. Mit dem Restglied

$$R(\xi, x, x_{\text{mess}}) := \frac{1}{2} f''(\xi) (x - x_{\text{mess}})^2 \quad (6.60)$$

erhalten wir

$$\begin{aligned}\Delta_{\max} f &= \max_{x \in X} |f(x_{\text{mess}}) + f'(x_{\text{mess}})(x - x_{\text{mess}}) + R(\xi, x, x_{\text{mess}}) - f(x_{\text{mess}})| \\ &= \max_{x \in X} |f'(x_{\text{mess}})(x - x_{\text{mess}}) + R(\xi, x, x_{\text{mess}})|.\end{aligned}\quad (6.61)$$

Da ξ stets zwischen x und x_{mess} liegt, können wir diesen Wert nach oben abschätzen:

$$\begin{aligned}\Delta_{\max} f &\leq \max_{x \in X} |f'(x_{\text{mess}})(x - x_{\text{mess}})| + \max_{x \in X} \max_{\xi \in X} |R(\xi, x, x_{\text{mess}})| \\ &= |f'(x_{\text{mess}})| \max_{x \in X} |x - x_{\text{mess}}| + \frac{1}{2} \left(\max_{\xi \in X} |f''(\xi)| \right) \max_{x \in X} |x - x_{\text{mess}}|^2 \\ &= |f'(x_{\text{mess}})| \Delta_{\max} x + \frac{1}{2} \left(\max_{\xi \in X} |f''(\xi)| \right) (\Delta_{\max} x)^2.\end{aligned}\quad (6.62)$$

Üblicherweise wird $\Delta_{\max} x$ klein sein, insbesondere viel kleiner als 1 (Schreibweise: $\Delta_{\max} x \ll 1$). Folglich wird $(\Delta_{\max} x)^2$ deutlich kleiner sein als $\Delta_{\max} x$. Für $\Delta_{\max} x = 0.01$ erhalten wir beispielsweise $(\Delta_{\max} x)^2 = 0.0001$. Zum Zwecke der Beurteilung des zu erwartenden Fehlers genügt also die Betrachtung des ersten Summanden. Nur in wenigen Ausnahmefällen, z.B. wenn der Messfehler sehr groß ist oder wenn die Funktion f stark gekrümmt ist (große zweite Ableitung!), ist der quadratische Fehleranteil relevant. In solch einem Fall können auch Taylor-Polynome zweiten Grades verwendet werden, damit man sich die aufwendige (und ggf. gar nicht machbare) Maximierung der zweiten Ableitung spart (man vernachlässigt dann den Term dritter Ordnung).

Merke

Für den Fehler $\Delta_{\max} f$ in den aus Messwerten berechneten Funktionswerten kann man für Praxiszwecke von der Abschätzung

$$\Delta_{\max} f \lesssim |f'(x_{\text{mess}})| \Delta_{\max} x \quad (6.63)$$

ausgehen. Dabei steht $\lesssim$ für “im Wesentlichen kleiner als” und deutet an, dass es sich nicht um eine exakte Ungleichung handelt.

Die Ableitung der Funktion f am gemessenen Wert liefert also den Verstärkungsfaktor bei der Fehlerfortpflanzung.

Beachte jedoch: Bei streng mathematischer Betrachtung gilt diese Abschätzung nicht!

Nebenbemerkung

Man könnte nun einwenden, dass durch die Wahl der Maßeinheiten das Erfüllen der Bedingung $\Delta_{\max} x \ll 1$ beeinflusst werden kann. Misst man beispielsweise eine Strecke x in Meter und der maximale Messfehler beträgt $\Delta_{\max} = 0.01$ Meter, so ist die Bedingung offensichtlich erfüllt. Misst man jedoch in Millimeter, so führt der

gleiche Messfehler zu $\Delta_{\max} = 10$ Millimeter, was deutlich über 1 liegt. In ersterem Fall wäre das Vernachlässigen des quadratischen Anteils in der Fehlerabschätzung okay; in zweiterem eher nicht, da dann $(\Delta_{\max}x)^2 = 100$ viel größer ist als $\Delta_{\max}x$, also durchaus relevant ist. Allerdings ändert ein Wechsel der Einheiten auch die Funktion f . Werden die Argumente von f in Metern angegeben, so entspricht die Arbeit mit Millimetern der Funktion $x \mapsto f(0.001x)$, sodass die zweite Ableitung zu $0.001^2 f''(0.001x)$ wird (Kettenregel!). Somit ist das Maximum über die zweite Ableitung nun sehr klein, sodass der üblicherweise vernachlässigte zweite Summand in der Fehlerabschätzung dennoch klein und damit vernachlässigbar ist.

Die Forderung $\Delta_{\max}x \ll 1$ ist also nicht zwingend für das Vernachlässigen des zweiten Summands, sondern stellt lediglich eine einfache Regel für die praktische Arbeit dar. Ist diese nicht erfüllt, kann das Vernachlässigen des zweiten Summands dennoch zulässig sein. Dieser Fall sollte dann entsprechend geprüft werden.

Beispiele

Beispiel

Der Flächeninhalt eines Kreises (in Quadratmeter) soll aus dem gemessenen Radius berechnet werden. Die Messung ergab einen Radius von 3 ± 0.01 Meter. Der zu erwartende absolute Fehler beim Flächeninhalt beträgt somit 0.188 Quadratmeter. Dies entspricht einem relativen Fehler von 0.67 Prozent (Video 640, Notizen dazu)?

Hinweis: Der exakt berechnete maximale Fehler liegt bei rund 0.189 Quadratmeter, also ebenfalls rund 0.67 Prozent. Die Größe des vernachlässigten zweiten Summanden in der Fehlerabschätzung beträgt $\frac{\pi}{10000} \approx 0.000314$.

Beispiel

Verwendet man im vorhergehenden Beispiel alle Maße in Millimeter, so erhält man aus der Näherungsformel den selben relativen Fehler wie zuvor, obwohl die Bedingung $\Delta_{\max}x \ll 1$ offensichtlich nicht erfüllt ist. Dies liegt in der Tatsache begründet, dass der zweite, vernachlässigte Summand in der Fehlerabschätzung mit einem Wert von $200\pi \approx 628.3$ im Vergleich zum ersten Summand mit einem Wert von $60000\pi \approx 188495.6$ trotzdem klein ist und dessen Vernachlässigung somit keinen relevanten Einfluss auf die berechnete Fehlerschranke hat (Video 645, Notizen dazu).

Der im zweiten Beispiel beobachtete Effekt, dass das Vernachlässigen des zweiten Sum-

manden in der Fehlerabschätzung auch bei $\Delta_{\max} x > 1$ hinnehmbar ist, ist nicht (!) allgemein gültig, sondern der konkreten Struktur der betrachteten Funktion geschuldet: Für wachsendes Argument bleibt die Krümmung (zweite Ableitung) konstant, während der Anstieg (erste Ableitung) wächst. Dadurch bleibt der zweite Summand in der Fehlerabschätzung vergleichsweise (!) klein, obwohl der große maximale Messfehler dort quadratisch eingeht.

7 Integralrechnung

Grundlagen der Berechnung von Flächeninhalten krummlinig begrenzter Flächen.

7.1 Flächeninhalt

Wir werden uns hier intensiv mit der Frage beschäftigen, was wir eigentlich unter dem Inhalt einer krummlinig begrenzten (also nahezu beliebig geformten) Fläche verstehen. Eine Fläche ist dabei eine Teilmenge des $\mathbb{R}^2$. Nicht betrachtet werden gekrümmte Flächen, also z.B. Teilmengen von Kugeloberflächen.

Später wird sich zeigen, dass die Berechnung von Flächeninhalten eng mit Funktionen (zur Beschreibung des Randes) und Ableitungen von Funktionen verknüpft ist. Dieser zunächst nicht sichtbare Zusammenhang wird uns ein mächtiges Werkzeug zur praktischen Berechnung von Flächeninhalten nahezu beliebiger Flächen an die Hand geben.

Problemstellung

Schon in der Schulmathematik wurden Flächeninhalte berechnet. Das lief nach dem Schema “suche den Typ der Fläche in der Formelsammlung und nutze die dort angegebene Formel für den Flächeninhalt”. Aber was verstehen ganz allgemein unter “Flächeninhalt”? Wie erklären wir jemandem, was zu tun ist, wenn beispielsweise der Flächeninhalt einer eiförmigen Fläche zu berechnen ist?

Aus praktischer Sicht können wir den Flächeninhalt einer beliebigen Fläche beschreiben als “Menge an Farbe, die zum Anstreichen der Fläche benötigt wird”. Oder etwas präziser (da die Schichtdicke der Farbe nicht klar ist): Man versehe die Fläche mit einem 1 Zentimeter hohen Rand und fülle die Fläche bis zur Oberkante des Randes mit Wasser. Die benötigte Wassermenge in Kubikzentimeter ist dann gerade (per Definition!) der Flächeninhalt in Quadratzentimeter. Praktisch ist das Problem also durchaus lösbar.

Offen ist, wie man die benötigte Wassermenge im Vorfeld berechnen kann. Für sehr große oder sehr kleine Flächen ist die beschriebene Form der experimentellen Inhaltsbestimmung kaum umsetzbar. Auch ist der Aufwand recht groß (gekrümmten Rand bauen; alles muss dicht sein usw.).

Das Frage der sinnvollen, allgemeinen Definition des Inhalts einer Fläche wurde erst um 1900 zufriedenstellend beantwortet. Für mathematische Verhältnisse ist dieses Thema also noch regelrecht “neu”.

Idee

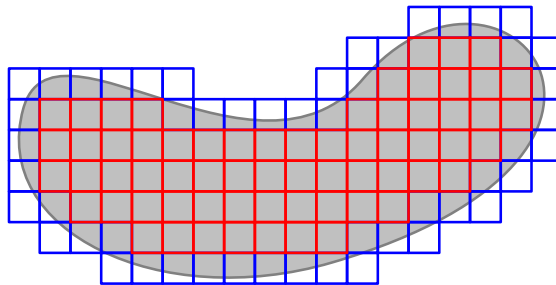
Die Grundidee zur mathematisch exakten Definition des Begriffs “Flächeninhalt” ist folgende.

1. Zerlege die Fläche in Teilflächen mit berechenbarem Flächeninhalt.
2. Addiere die Inhalte aller Teilflächen.

Für manche Flächen ist dies leicht machbar. Für andere jedoch werden sehr viele Teilflächen benötigt oder eine Zerlegung ist nicht möglich, da der krummlinige Rand beim

Zerlegen nicht verschwindet und stets nicht berechenbare Teilflächen übrig bleiben.

Deshalb modifizieren wir die Idee etwas: Zum einen legen wir uns auf eine möglichst einfache Form von Teilflächen fest, um die Berechnung der einzelnen Flächeninhalte zu ermöglichen. Die einfachste Fläche mit Blick auf Inhaltsberechnung ist das Quadrat. Zum anderen fordern wir nicht mehr, dass die Fläche vollständig in Teilflächen zerlegt werden muss, sondern wir überdecken die Fläche mit Quadraten gleicher Größe und zählen wie viele Quadrate vollständig innerhalb der Fläche liegen und wie viele Quadrate wenigstens einen Teil der Fläche enthalten. Die Summen der jeweiligen Quadratinhalte liefern den Flächeninhalt einer etwas kleineren und einer etwas größeren Fläche. Je kleiner die Quadrate sind, desto näher werden sich diese beiden Werte kommen. Betrachten wir den Grenzübergang “Kantenlänge gegen Null”, so sollten beide Werte übereinstimmen. Diese Zahl betrachten wir nun als Flächeninhalt der unregelmäßigen Fläche.



Die roten Quadrate liegen vollständig im Inneren der Fläche. Die blauen Quadrate überlappen mit der Fläche. Der Flächeninhalt der grauen Fläche muss zwischen der Summe der Flächen der roten Quadrate und der Summe der Flächen der roten und blauen Quadrate liegen. Je kleiner die Quadrate, desto kleiner der Unterschied zwischen beiden Schranken.

Umsetzung

Setzen die Idee zur Definition des Flächeninhalts beliebiger Flächen nun mathematisch exakt um. Zu $F \subseteq \mathbb{R}^2$ ist also eine Zahl $A(F) \geq 0$ gesucht, die unserer intuitiven Vorstellung von “Flächeninhalt” entspricht. Insbesondere soll $A(F)$ für einfache Flächen F natürlich die bisher mittels konkreter Formeln berechneten Flächeninhalte liefern.

Zur Überdeckung von F verwenden wir ein Raster aus achsenparallelen Quadraten $[x, x+l] \times [y, y+l]$ mit Kantenlänge $l > 0$, welches der Einfachheit halber so ausgerichtet ist, dass der Punkt $(0,0)$ gerade ein Eckpunkt eines der Quadrate ist. Die x - und die y -Koordinaten aller Eckpunkte aller Quadrate sind also ganzzahlige Vielfache von l .

Mit $i(l) \in \mathbb{N}$ bezeichnen wir die Anzahl von Quadraten, die vollständig von F überdeckt werden (innere Überdeckung). Mit $a(l) \in \mathbb{N}$ bezeichnen wir die Anzahl von Quadraten, die mindestens einen Punkt von F enthalten (äußere Überdeckung). Offensichtlich gilt also stets $i(l) \leq a(l)$. Damit beide Zahlen existieren (also nicht “unendlich” sind), müssen wir unsere Betrachtungen auf **beschränkte Flächen** eingrenzen, d.h. auf solche Flächen

F , für die es eine Zahl $c \in \mathbb{R}$ gibt mit

$$\sqrt{x^2 + y^2} \leq c \quad \text{für alle } (x, y) \in F. \quad (7.1)$$

Merke

Sei $F \subseteq \mathbb{R}^2$ beschränkt und seien $i(l)$ und $a(l)$ für $l > 0$ wie oben beschrieben. Die Menge F heißt **Jordan-messbar**, wenn die beiden Grenzwerte

$$\lim_{l \rightarrow \infty} i(l) l^2 \quad \text{und} \quad \lim_{l \rightarrow \infty} a(l) l^2 \quad (7.2)$$

existieren und gleich sind. In diesem Fall werden die beiden (gleichen) Grenzwerte als **Flächeninhalt** von F bezeichnet.

Die Werte $i(l) l^2$ und $a(l) l^2$ sind gerade die Flächeninhalte der inneren und der äußeren Überdeckung zur Kantenlänge l .

Beispiele und Einordnung

Beispiel (Flächeninhalt eines Rechtecks)

Wenn unsere Definition von “Flächeninhalt” sinnvoll ist, sollten wir damit die üblichen Formeln für Flächeninhalte von einfachen Flächen herleiten können. Tun dies für das Rechteck $[0, c] \times [0, d]$.

Für festes $l > 0$ sei $m(l) \in \mathbb{N}$ die Anzahl der Quadratspalten der inneren Überdeckung. Entsprechend bezeichne $n(l) \in \mathbb{N}$ die Anzahl der Quadratzeilen. Es gilt also $i(l) = m(l) n(l)$. Die Zahlen $m(l)$ und $n(l)$ können wir mittels der Kantenlängen c und d des Rechtecks ausdrücken:

$$m(l) = \left\lfloor \frac{c}{l} \right\rfloor \quad \text{und} \quad n(l) = \left\lfloor \frac{d}{l} \right\rfloor. \quad (7.3)$$

Dabei bezeichnet $\lfloor \dots \rfloor$ das Abrunden des eingeschlossenen Wertes. Für den Grenzwert des Inhalts der inneren Überdeckung erhalten wir

$$\begin{aligned} \lim_{l \rightarrow 0} i(l) l^2 &= \lim_{l \rightarrow 0} \left\lfloor \frac{c}{l} \right\rfloor \left\lfloor \frac{d}{l} \right\rfloor l^2 \\ &\geq \lim_{l \rightarrow 0} \left(\frac{c}{l} - 1 \right) \left(\frac{d}{l} - 1 \right) l^2 \\ &= \lim_{l \rightarrow 0} (c - l)(d - l) = c d. \end{aligned} \quad (7.4)$$

Analog folgt

$$\begin{aligned}\lim_{l \rightarrow 0} a(l) l^2 &= \lim_{l \rightarrow 0} \left\lceil \frac{c}{l} \right\rceil \left\lceil \frac{d}{l} \right\rceil l^2 \\ &\leq \lim_{l \rightarrow 0} \left(\frac{c}{l} + 1 \right) \left(\frac{d}{l} + 1 \right) l^2 \\ &= \lim_{l \rightarrow 0} (c + l) (d + l) = c d,\end{aligned}\tag{7.5}$$

wobei $\lceil \dots \rceil$ für das Aufrunden des eingeschlossenen Wertes steht.

Aus

$$c d \leq \lim_{l \rightarrow 0} i(l) l^2 \leq \lim_{l \rightarrow 0} a(l) l^2 \leq c d \tag{7.6}$$

folgt nun die Gleichheit beider Grenzwerte. Der Flächeninhalt ist entsprechend unserer Definition des Begriffs also wie erwartet $c d$.

Bereits bei einem Rechteck ist die konkrete Berechnung des Flächeninhalts recht mühsam. Bei komplexeren Fläche ist die Berechnung direkt über die Definition kaum machbar. Deshalb werden wir uns im nächsten Unterkapitel ein alternatives Vorgehen überlegen.

Nebenbemerkung

Es gibt durchaus beschränkte Mengen, die nicht Jordan-messbar sind, denen wir also auf dem beschriebenen Weg keinen Flächeninhalt zuordnen können. Solche Mengen sind aber praktisch nicht relevant.

Beispiel (nicht Jordan-messbare Menge)

Die Menge $F := [0, 1] \times ([0, 1] \cap \mathbb{Q})$ ist nicht Jordan-messbar. Kein noch so kleines Quadrat wird von F vollständig überdeckt, da jedes Quadrat auch Punkte mit nichtrationalen y -Koordinaten enthält. Also $i(l) = 0$ für alle $l > 0$. Andererseits besteht jede äußere Überdeckung aus mindestens einem Quadrat. Also gilt stets $a(l) \geq 1$. Somit können die Grenzwerte der Inhalte der inneren und der äußeren Überdeckungen nicht gleich sein.

Nebenbemerkung

Es gibt verschiedene Varianten von “Messbarkeit” in der Mathematik. Neben der hier eingeführten Jordan-Messbarkeit spielt die sogenannte Lebesgue-Messbarkeit eine wesentliche Rolle. Diese ist eine Verallgemeinerung der Jordan-Messbarkeit (jede Jordan-messbare Menge ist auch Lebesgue-messbar) und deckt mehr Mengen ab. Jedoch gibt es trotzdem Mengen, die auch nicht Lebesgue-messbar sind.

Die Konstruktion des Flächeninhalts auf Basis der Jordan-Messbarkeit kann völlig analog auch in drei oder mehr Dimensionen durchgeführt werden. In $\mathbb{R}^3$ erhält man mittels Überdeckungen aus Würfeln auf diese Weise eine Definition des Begriffs “Volumen”, die auch sehr unregelmäßige Körper mit einschließt.

7.2 Normalbereiche

Für die praktische Berechnung von Flächeninhalten unregelmäßig geformter Flächen müssen wir zwei Probleme lösen:

1. Die Fläche muss mittels Formeln beschreibbar sein.
2. Die Berechnung des Flächeninhalts muss auch bei komplexeren Flächen mit vertretbarem Aufwand durchführbar sein.

In diesem Unterkapitel widmen wir uns vor allem dem ersten Problem. Die Lösung des zweiten Problems bereiten wir vor, gehen sie aber erst im nächsten Unterkapitel vollständig an.

Begriff

Wir beschränken alle nachfolgenden Betrachtungen auf folgenden Typ von Flächen:

Merke

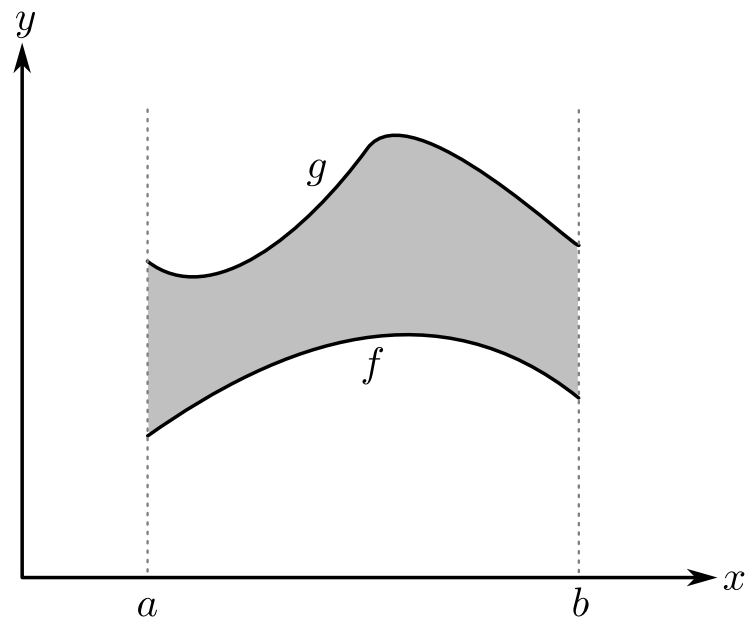
Lässt sich eine Menge $B \subseteq \mathbb{R}^2$ in der Form

$$B = \{(x, y) \in \mathbb{R}^2 : a \leq x \leq b, f(x) \leq y \leq g(x)\} \quad (7.7)$$

schreiben, wobei $f : [a, b] \rightarrow \mathbb{R}$ und $g : [a, b] \rightarrow \mathbb{R}$ zwei stetige Funktionen über einem Intervall $[a, b]$ sind, die $f(x) \leq g(x)$ für alle $x \in [a, b]$ erfüllen, so nennen wir B einen **Normalbereich bzgl. der x -Achse**.

Man kann auch **Normalbereiche bzgl. der y -Achse** einführen. Diese können jedoch durch drehen um 90° oder durch Spiegelung an der Gerade $y = x$ stets in Normalbereiche bzgl. der x -Achse überführt werden, sodass deren weitere Betrachtung hier unterbleibt.

Beispiele



Ein Normalbereich wird durch ein Intervall für die x -Koordinaten und durch zwei Funktionen beschrieben. Die beiden Funktionsgraphen bilden die untere und die obere Grenze des Normalbereichs.

Beispiel

Das Viereck $B \subseteq \mathbb{R}^2$, welches durch die Geraden

$$x = 1, \quad y = \frac{1}{2}x + 1, \quad x = 3, \quad y = x + 2 \quad (7.8)$$

begrenzt wird, ist ein Normalbereich:

$$B = \left\{ (x, y) : 1 \leq x \leq 3, \frac{1}{2}x + 1 \leq y \leq x + 2 \right\} \quad (7.9)$$

(Video 660, Notizen dazu).

Beispiel

Das krummlinig begrenzte Dreieck $B \subseteq \mathbb{R}^2$, welches durch die Kurven

$$y = x^2, \quad y = 6 - x, \quad x = 4 \quad (7.10)$$

eingeschlossen wird, ist ein Normalbereich:

$$B = \{(x, y) : 2 \leq x \leq 3, 6 - x \leq y \leq x^2\} \quad (7.11)$$

(Video 665, Notizen dazu).

Beispiel

Die Kreisscheibe $K \subseteq \mathbb{R}^2$ mit Radius 2 um den Punkt $(3, 1)$ ist ein Normalbereich:

$$K = \left\{ (x, y) : -2 \leq x \leq 2, 1 - \sqrt{4 - (x - 3)^2} \leq y \leq 1 + \sqrt{4 - (x - 3)^2} \right\} \quad (7.12)$$

(Video 670, Notizen dazu).

Flächeninhalt von Normalbereichen

Berechnen nun den Flächeninhalt eines Normalbereichs

$$B := \{(x, y) \in \mathbb{R}^2 : a \leq x \leq b, f(x) \leq y \leq g(x)\} \quad (7.13)$$

auf Basis unserer Definition des Flächeninhalts. Der Einfachheit halber nehmen wir zunächst an, dass f und g nur positive Werte annehmen. Dann können wir den Flächeninhalt von B als Differenz der Flächeninhalte zwischen g und der x -Achse sowie f und der x -Achse ausdrücken. Unsere Aufgabe beschränkt sich damit auf das Berechnen des Flächeninhalts zwischen der x -Achse und einem Funktionsgraph. Beschreiben das Vorgehen hier für die Funktion f .

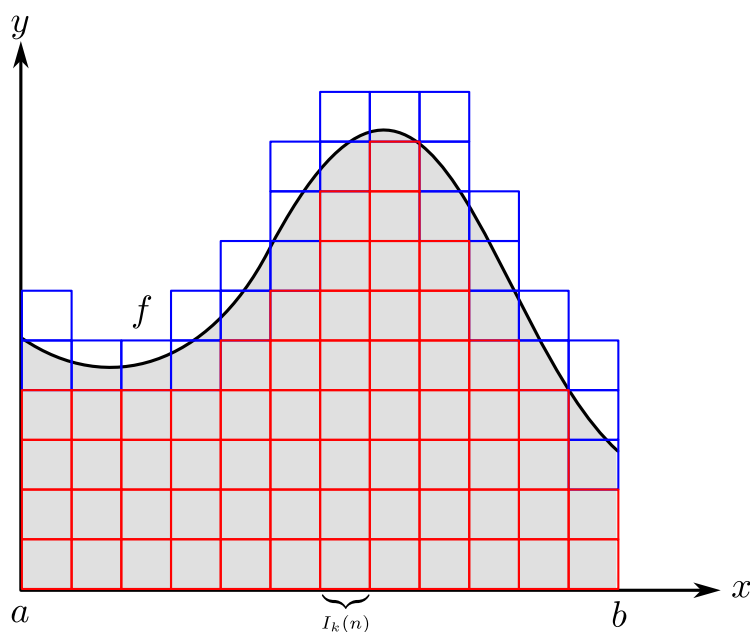
Um die Formeln einfach zu halten, gehen wir von $a = 0$ aus, was durch Verschieben des Bereichs stets problemlos erreicht werden kann, also keine inhaltliche, sondern nur eine formale Einschränkung darstellt. In der Mathematik ist an dieser Stelle die Formulierung „*Ohne Beschränkung der Allgemeinheit* nehmen wir $a = 0$ an“ üblich.

Wir überdecken den Bereich B wie im vorhergehenden Unterkapitel mit einem Raster aus Quadraten der Kantenlänge $l > 0$ und ermitteln die Anzahlen $i(l)$ und $a(l)$ der Quadrate in der inneren und in der äußeren Überdeckung.

Man kann sich relativ leicht überlegen (tun wir hier nicht), dass $\lim_{l \rightarrow 0} i(l) l^2$ genau dann existiert, wenn $\lim_{n \rightarrow \infty} i(l_n) l_n^2$ für irgendeine Folge $(l_n)_{n \in \mathbb{N}}$ mit $\lim_{n \rightarrow \infty} l_n = 0$ gilt. Außerdem sind im Falle der Existenz die beiden Grenzwerte gleich. Analoges gilt für $a(l) l^2$. Diese Feststellung vereinfacht das weitere Vorgehen, da wir statt dem Grenzwert $\lim_{l \rightarrow 0} i(l) l^2$ einer Funktion bzgl. l nur den Grenzwert $\lim_{n \rightarrow \infty} i(l_n) l_n^2$ einer Zahlenfolge betrachten müssen.

Wählen wir nun $l_n := \frac{b}{n}$, so gilt offensichtlich $\lim_{n \rightarrow \infty} l_n = 0$. Innere und äußere Überdeckung für ein n haben stets gleich viele Spalten, da $\lfloor \frac{b}{l_n} \rfloor = \lfloor n \rfloor = n = \lceil n \rceil = \lceil \frac{b}{l_n} \rceil$ gilt. Anders ausgedrückt: b liegt immer genau auf der Grenze zwischen zwei Quadratspalten. Entsprechend haben Unter- und Obersumme die gleiche Anzahl Summanden.

Beim Summieren der Quadratflächen gehen wir spaltenweise vor. Die Anzahl der Spalten ist n . Die Spalten seien von links nach rechts mit $1, 2, \dots, n$ nummeriert. Die Anzahl der



Die Oberkanten der äußeren Überdeckung (blau) liegen vollständig oberhalb des Funktionsgraphs. Die Oberkanten der inneren Überdeckung (rot) liegen vollständig unterhalb des Funktionsgraphs.

Quadrate in Spalte k hängt offensichtlich von den Funktionswerten von f im Bereich dieser Spalte ab. Sei also $I_k(n) := [(k-1)l, kl]$ das Intervall von x -Koordinaten zur k -ten Spalte. Dann ist die Anzahl Quadrate der inneren Überdeckung in Spalte k gerade

$$\left\lfloor \frac{\min_{x \in I_k(n)} f(x)}{l_n} \right\rfloor \quad (7.14)$$

und wir erhalten für den Grenzwert des Flächeninhalts der inneren Überdeckung

$$\lim_{n \rightarrow \infty} i(l_n) l_n^2 = \lim_{n \rightarrow \infty} \left(l_n \sum_{k=1}^n \min_{x \in I_k(n)} f(x) \right), \quad (7.15)$$

sofern der rechte Grenzwert existiert (Video 675, Notizen dazu). Der Ausdruck im rechten Grenzwert wird auch als **Riemannsche Untersumme** zur Unterteilungsfinheit n bezeichnet.

Völlig analog erhalten wir für die äußeren Überdeckungen

$$\lim_{n \rightarrow \infty} a(l_n) l_n^2 = \lim_{n \rightarrow \infty} \left(l_n \sum_{k=1}^n \max_{x \in I_k(n)} f(x) \right). \quad (7.16)$$

Der Ausdruck im zweiten Grenzwert wird auch als **Riemannsche Obersumme** zur Unterteilungseinheit n bezeichnet.

Machen wir die Verschiebung zu $a = 0$ rückgängig, so können wir folgende Definition formulieren:

Merke

Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und bezeichne $I_k(n) := [a + (k-1) \frac{b-a}{n}, a + k \frac{b-a}{n}]$ für $k = 1, \dots, n$ die Teilintervalle von $[a, b]$ bei Unterteilung in n gleich große Teile. Existieren die Grenzwerte

$$\lim_{n \rightarrow \infty} \left(\frac{b-a}{n} \sum_{k=1}^n \min_{x \in I_k(n)} f(x) \right) \quad (7.17)$$

und

$$\lim_{n \rightarrow \infty} \left(\frac{b-a}{n} \sum_{k=1}^n \max_{x \in I_k(n)} f(x) \right) \quad (7.18)$$

der Riemannschen Unter- und Obersummen und stimmen die beiden Grenzwerte überein, so heißt f **Riemann-integrierbar**. Der Wert der beiden (gleichen) Grenzwerte heißt dann **Riemann-Integral** von f über $[a, b]$ und entspricht gerade dem vorzeichenbehafteten Flächeninhalt zwischen dem Graph von f und der x -Achse.

Als Symbol für das Riemann-Integral wird aus historischen Gründen und auch aus Gründen, die später erhellt werden, das Symbol

$$\int_a^b f(x) \, dx \quad (7.19)$$

verwendet.

Mit “vorzeichenbehaftet” meinen wir hier, dass das Riemann-Integral negativ ist, wenn der Funktionsgraph vollständig unterhalb der x -Achse liegt. Schneidet die Funktion die x -Achse, so gehen Flächenanteile unterhalb der x -Achse negativ in das Riemann-Integral ein. Bei gleichen Anteilen ober- und unterhalb der x -Achse wird das Riemann-Integral also Null sein.

Mit Blick auf den Flächeninhalt von Normalbereichen (mit Untergrenze f und Obergrenze g) ist diese Vorzeicheneigenschaft des Riemann-Integrals sehr nützlich. Sie führt dazu, dass die Differenz

$$\int_a^b g(x) \, dx - \int_a^b f(x) \, dx, \quad (7.20)$$

also der Flächeninhalt des Normalbereichs, stets positiv ist (wenn g überall oberhalb

von f verläuft). Das gilt insbesondere auch, wenn der Normalbereich unterhalb der x -Achse liegt (also beide Integrale negativ sind), und auch, wenn die x -Achse durch den Normalbereich verläuft.

Hinweise zum Riemann-Integral

Man überlegt sich leicht anhand der Eigenschaften von Grenzwerten, dass

$$\int_a^b f(x) \, dx + \int_a^b g(x) \, dx = \int_a^b f(x) + g(x) \, dx \quad (7.21)$$

und

$$\lambda \int_a^b f(x) \, dx = \int_a^b \lambda f(x) \, dx \quad \text{für } \lambda \in \mathbb{R} \quad (7.22)$$

gelten. Außerdem gilt

$$\int_a^b f(x) \, dx + \int_b^c f(x) \, dx = \int_a^c f(x) \, dx \quad (7.23)$$

(ebenfalls leicht einzusehen).

Die praktische Berechnung von Riemann-Integralen klären wir im nächsten Unterkapitel. Obwohl der Weg über Unter- und Obersummen einfacher ist als über innere und äußere Überdeckungen aus Quadraten, ist der Weg immer noch zu mühsam.

Das Riemann-Integral wird oft einfach als Integral bezeichnet. Zur Unterscheidung von einem im nächsten Unterkapitel einzuführenden, eng verbundenen Begriff, spricht man auch vom **bestimmten Integral**.

Nebenbemerkung

Neben Riemann-Integralen gibt es noch eine Reihe anderer Integralbegriffe, die für uns aber keine Bedeutung haben, da sie nur in für uns unwesentlichen Nuancen vom hier eingeführten Riemann-Integral abweichen. Am verbreitetsten ist das Lebesgue-Integral, welches das Lebesgue-Maß zur Berechnung der Fläche unter dem Funktionsgraph nutzt.

Uneigentliche Riemann-Integrale

Gelegentlich trifft man auf Flächen, die unbeschränkt sind, aber in gewissem Sinne einen endlichen Flächeninhalt besitzen. Solchen Flächen werden wir beispielsweise in der Wahrscheinlichkeitsrechnung begegnen. Praktisch sollten solche Flächen eigentlich nicht

relevant sein, da die vom Menschen wahrnehmbare Welt eine beschränkte Menge ist; somit sind auch alle Flächen und Körper darin beschränkt. Bei der vereinfachenden Modellierung gewisser Sachverhalte tauchen dann aber doch solche unbeschränkten Flächen auf. Eine ähnliche Situation hatten wir schon bei Polstellen beschrieben (vgl. Beispiel (Massepunkt)).

Wir beschränken uns hier auf den einfachsten Fall, nämlich die Situation, dass der Definitionsbereich einer Funktion f , die zusammen mit der x -Achse eine Fläche begrenzt, unbeschränkt ist. Eine andere Variante unbeschränkter Flächen mit endlichem Inhalt tritt bei Polstellen am Rand oder im Inneren des Definitionsbereichs einer Funktion auf. Diesen Fall kann man völlig analog handhaben, sodass wir nicht näher darauf eingehen.

Einseitig unbeschränkter Definitionsbereich Wie so oft in der Mathematik, führt man eine komplexere auf eine einfachere Situation zurück. Den Flächeninhalt komplexerer Flächen hatten wir beispielsweise über den Grenzwert der Flächeninhalte von aus Quadraten zusammengesetzten Flächen erklärt. Den Flächeninhalt unbeschränkter Flächen erklären wir nun wiederum als Grenzwert des Inhalts gewisser beschränkter Flächen:

Merke

Sei $f : (-\infty, b) \rightarrow \mathbb{R}$ stetig. Dann setzen wir

$$\int_{-\infty}^b f(x) \, dx := \lim_{a \rightarrow -\infty} \int_a^b f(x) \, dx, \quad (7.24)$$

sofern dieser Grenzwert existiert, und nennen diesen Wert **uneigentliches Riemann-Integral** zu f . Für $f : [a, \infty) \rightarrow \mathbb{R}$ setzen wir analog

$$\int_a^{\infty} f(x) \, dx := \lim_{b \rightarrow \infty} \int_a^b f(x) \, dx. \quad (7.25)$$

Es werden also die (vorzeichenbehafteten) Inhalte immer größer werdender Teilflächen berechnet. Konvergiert diese Folge von Flächeninhalten, so kann man auch der unbeschränkten Gesamtfläche sinnvoll einen Flächeninhalt zuordnen. Bei bestimmter Divergenz kann man die Werte ∞ bzw. $-\infty$ als Flächeninhalt ansehen. Bei unbestimmter Divergenz ist überhaupt kein sinnvoller Inhaltsbegriff mehr denkbar.

Beispiele für uneigentliche Integrale schauen wir uns erst im nächsten Unterkapitel an, wenn wir ein einfaches Verfahren zur Berechnung von Integralen kennengelernt haben.

Beidseitig unbeschränkter Definitionsbereich Ist der Definitionsbereich beidseitig unbeschränkt, also ganz $\mathbb{R}$, so gibt es zwei verschiedene Möglichkeiten zur Definition eines

Flächeninhalts. Entweder man teilt die Fläche in zwei Teilflächen und betrachtet jede Teilfläche für sich und addiert die beiden Inhalte dann, oder man führt den Grenzübergang symmetrisch an beiden Enden des Definitionsbereichs durch. Da beide Varianten gebräuchlich sind, führen wir hier beide ein.

Merke

Sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stetig. Existieren die beiden uneigentlichen Riemann-Integrale auf den Intervallen $(-\infty, 0]$ und $[0, \infty)$, so heißt

$$\int_{-\infty}^{\infty} f(x) \, dx := \int_{-\infty}^0 f(x) \, dx + \int_0^{\infty} f(x) \, dx \quad (7.26)$$

uneigentliches Riemann-Integral zu f .

Außerdem setzen wir

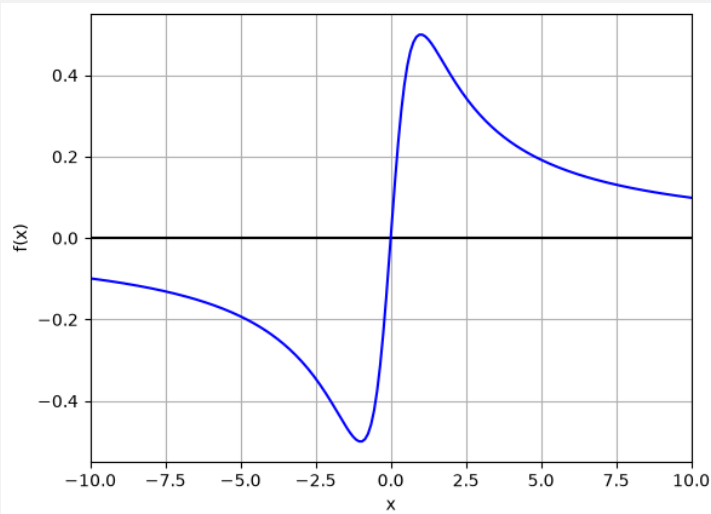
$$\text{CH} \int_{-\infty}^{\infty} f(x) \, dx := \lim_{c \rightarrow \infty} \int_{-c}^c f(x) \, dx, \quad (7.27)$$

sofern dieser Grenzwert existiert, und nennen diese Zahl den **Cauchyschen Hauptwert** zu f .

Existiert das uneigentliche Riemann-Integral, so existiert auch der Cauchysche Hauptwert und beide stimmen überein. Umgekehrt muss dies nicht gelten.

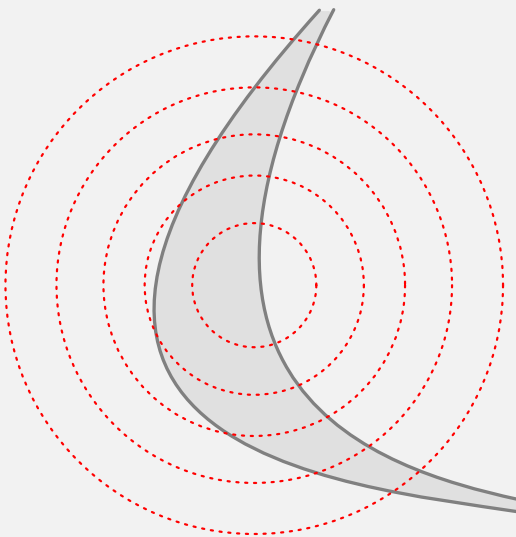
Beispiel

Für die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) : \frac{x}{1+x^2}$ ist aus Symmetriegründen der Cauchysche Hauptwert offensichtlich Null. Jedoch ist die Funktion nicht uneigentlich Riemann-Integrierbar, da man zeigen kann, dass die beiden einseitigen uneigentlichen Integrale (bestimmt) divergieren.



Nebenbemerkung

Für die Definition des Inhalts allgemeiner unbeschränkter Flächen kann man den Gedanke des Cauchyschen Hauptwerts nutzen: Man bilde den Durchschnitt der Fläche mit größer werdenden Kreisscheiben und betrachte den Grenzwert der Flächeninhalte der (beschränkten) Durchschnittsflächen.



In Formeln: Bezeichnet $K(r) \subseteq \mathbb{R}^2$ die Kreisscheibe mit Radius r um den Ursprung, so kann der Inhalt $A(F)$ einer unbeschränkten Fläche $F \subseteq \mathbb{R}^2$ als

$$A(F) := \lim_{r \rightarrow \infty} A(F \cap K(r)) \quad (7.28)$$

erklärt werden, sofern alle Durchschnitte $F \cap K(r)$ Jordan-messbar sind und der Grenzwert existiert.

Berechnet man diesen Grenzwert für eine beschränkte Fläche, so erhält man den korrekten Flächeninhalt, da für hinreichend großes r die Fläche vollständig in allen Kreisscheiben enthalten ist.

7.3 Praktische Integration

Wie kaum anders zu erwarten, werden Flächeninhalte und die dafür benötigten bestimmten Integrale heute mit dem Computer berechnet. Bevor wir darauf eingehen, schauen wir uns zunächst den klassischen manuellen Weg an, da dieser einen auch an anderen Stellen gelegentlich sehr nützlichen Zusammenhang zutage fördert (z.B. beim in dieser Veranstaltung nicht betrachteten Lösen von Differentialgleichungen).

Manuelles Integrieren

Die Definition des Riemann-Integrals taugt nicht für das praktische Ausrechnen von Integralen. Stattdessen werden wir ein zunächst gar nicht offensichtliches Resultat der Mathematik nutzen, welches Flächeninhalte mit Ableitungen von Funktionen verbindet.

Hauptsatz der Differential- und Integralrechnung Die folgende Erkenntnis, auf deren Herleitung wir hier verzichten, vereinfacht das Berechnen von Integralen erheblich:

Merke

Ist der Integrand in einem Riemann-Integral die Ableitung einer stetig differenzierbaren Funktion $f : [a, b] \rightarrow \mathbb{R}$, so gilt

$$\int_a^b f'(x) \, dx = f(b) - f(a). \quad (7.29)$$

Bei differenzierbaren Funktionen hatten wir als Definitionsbereich stets offene Intervalle (a, b) betrachtet, damit wir uns keine Gedanken über die Definition von Differenzierbarkeit an Randpunkten machen müssen. Bei Integralen verwenden wir hingegen abgeschlossene Intervalle, damit die entstehenden Flächen abgeschlossen sind, also einen Rand haben. Mit Differenzierbarkeit auf $[a, b]$ in obiger Aussage meinen wir, dass es ein offenes Intervall (c, d) gibt, auf dem die Funktion f differenzierbar ist und dass das Intervall $[a, b]$ enthält, also $[a, b] \subseteq (c, d)$.

Stammfunktionen Die Frage, ob eine zu integrierende Funktion überhaupt Ableitung einer anderen Funktion ist, ist zunächst offen. Nur für solche Funktionen hilft uns der Hauptsatz der Differential- und Integralrechnung bei der Berechnung von Riemann-Integralen. Zur weiteren Klärung führen wir den Begriff der Stammfunktion ein.

Merke

Jede differenzierbare Funktion, deren Ableitung gerade einer gegebenen Funktion f gleicht, heißt **Stammfunktion** von f .

Sind F_1 und F_2 zwei Stammfunktionen von f , so gilt $F_1(x) = F_2(x) + c$ für ein $c \in \mathbb{R}$ und alle x aus dem Definitionsbereich von f . Ist umgekehrt F eine Stammfunktion von f , so ist auch die durch $\tilde{F}(x) := F(x) + c$ gegebene Funktion $\tilde{F}$ für beliebiges $c \in \mathbb{R}$ eine Stammfunktion von f .

Wenn eine Funktion f eine Stammfunktion besitzt, dann besitzt sie also automatisch unendlich viele Stammfunktionen, die sich nur um eine additive Konstante unterscheiden.

Die meisten praktisch relevanten Funktionen besitzen eine Stammfunktion:

Merke

Sei $f : (a, b) \rightarrow \mathbb{R}$ stetig. Dann ist die durch

$$F(x) := \int_a^x f(\tilde{x}) \, d\tilde{x} \quad (7.30)$$

definierte Funktion $F : (a, b) \rightarrow \mathbb{R}$ eine Stammfunktion von f .

Aus praktischer Sicht viel relevanter ist die Frage, ob man Stammfunktionen auch explizit als Formel angeben kann. Hier sind Sachlage und Vorgehen ähnlich wie bei der Frage der Differenzierbarkeit:

- Für viele elementare Funktionen kann man eine Stammfunktion als Formel angeben.
- Für aus elementaren Funktionen zusammengesetzte Funktionen kann man oft mittels diverser “Integrationsregeln” eine Stammfunktion aus den Stammfunktionen der entsprechenden elementaren Funktionen zusammensetzen.
- In der Masse aller stetigen Funktionen sind die Funktionen, für die man eine Stammfunktion explizit als Formel angeben kann, in der Minderheit.

Wir werden im Folgenden einen kurzen Blick auf Stammfunktionen einiger elementarer Funktionen werfen, die Arbeit mit Integrationsregeln aber an das Ende des zweiten Semesters verschieben, da diese für Ingenieure kaum noch praktische Relevanz hat.

Bevor wir zu den elementaren Funktionen kommen, noch ein Hinweis auf eine verbreitete Schreibweise:

Merke

Für eine Stammfunktion einer Funktion f verwendet man häufig das Symbol

$$\int f(x) \, dx. \quad (7.31)$$

Man bezeichnet Stammfunktionen entsprechend auch als **unbestimmte Integrale**. Mit dieser Schreibweise kann der Hauptsatz der Differential- und Integralrechnung als

$$\int_a^b f(x) \, dx = F(b) - F(a) \quad \text{mit} \quad F(x) := \int f(x) \, dx \quad (7.32)$$

formuliert werden.

Beachte bei dieser Schreibweise, dass $\int f(x) \, dx$ eine Funktion ist, $\int_a^b f(x) \, dx$ aber eine Zahl. Welche Stammfunktion man wählt, spielt im Hauptsatz der Differential- und Integralrechnung keine Rolle, da durch die Differenz der beiden Stammfunktionswerte additive Konstanten wieder entfallen. Gelegentlich meint man mit dem Symbol $\int f(x) \, dx$ auch die Menge aller Stammfunktionen von f , statt einer konkreten Stammfunktion.

Stammfunktionen elementarer Funktionen In Formelsammlungen gibt es lange Listen mit Stammfunktionen für häufig benötigte Funktionen. Hier nur eine kleine Auswahl:

$$\int x^c \, dx = \frac{1}{c+1} x^{c+1} \quad \text{für } c \neq -1, \quad (7.33)$$

$$\int x^{-1} \, dx = \ln |x|, \quad (7.34)$$

$$\int \cos x \, dx = \sin x, \quad (7.35)$$

$$\int \sin x \, dx = -\cos x, \quad (7.36)$$

$$\int e^x \, dx = e^x. \quad (7.37)$$

Beachte, dass stets nur eine Stammfunktion angegeben ist. Addition beliebiger Konstanten liefert alle weiteren Stammfunktionen. Dass eine Funktion Stammfunktion einer anderen Funktion ist, kann man leicht durch Ableiten der Stammfunktion überprüfen.

Aus den Eigenschaften von Ableitungen sieht man außerdem sofort

$$\int f(x) + g(x) \, dx = \int f(x) \, dx + \int g(x) \, dx \quad (7.38)$$

und

$$\int \lambda f(x) \, dx = \lambda \int f(x) \, dx, \quad (7.39)$$

wobei diese Formeln für beliebige Auswahlen konkreter Stammfunktion gelten.

Beispiele Hier einige konkrete Beispiele zur manuellen Berechnung von Flächeninhalten mittels Riemann-Integralen.

Beispiel (Normalbereich)

Wir berechnen den Flächeninhalt des Normalbereichs

$$\left\{ (x, y) : 1 \leq x \leq 3, \frac{1}{2}x + 1 \leq y \leq x + 2 \right\} \quad (7.40)$$

aus einem früheren Beispiel zu Normalbereichen als Differenz

$$\int_1^3 x + 2 \, dx - \int_1^3 \frac{1}{2}x + 1 \, dx = \int_1^3 \frac{1}{2}x + 1 \, dx. \quad (7.41)$$

Wegen

$$\int \frac{1}{2}x + 1 \, dx = \frac{1}{2} \int x^1 \, dx + \int x^0 \, dx = \frac{1}{4}x^2 + x \quad (7.42)$$

erhalten wir als Flächeninhalt

$$\int_1^3 \frac{1}{2}x + 1 \, dx = \frac{1}{4} \cdot 3^2 + 3 - \left(\frac{1}{4} \cdot 1^2 + 1 \right) = 4 \quad (7.43)$$

Beispiel (unbeschränkte Fläche)

Wir berechnen den Flächeninhalt der unbeschränkten Fläche unter

$$f : [0, \infty) \rightarrow \mathbb{R}, \quad f(x) := \frac{x}{x^2 + 1} \quad (7.44)$$

(vgl. früheres Beispiel zu uneigentlichen Integralen) als uneigentliches Integral

$$\int_0^\infty \frac{x}{x^2 + 1} \, dx = \lim_{b \rightarrow \infty} \int_0^b \frac{x}{x^2 + 1} \, dx. \quad (7.45)$$

Aus einer Formelsammlung oder mittels einer am Ende des zweiten Semesters noch näher zu betrachtenden Integrationsregel (“Integration durch Substitution”)

findet man die Stammfunktion

$$\int \frac{x}{x^2 + 1} dx = \ln(x^2 + 1) \quad (7.46)$$

(leicht überprüfbar durch Ableiten!). Somit

$$\int_0^b \frac{x}{x^2 + 1} dx = \ln(b^2 + 1) - \ln(0^2 + 1) = \ln(b^2 + 1), \quad (7.47)$$

also

$$\lim_{b \rightarrow \infty} \int_0^b \frac{x}{x^2 + 1} dx = \infty. \quad (7.48)$$

Das uneigentliche Integral ist bestimmt divergent, der Flächeninhalt also nicht endlich.

Beispiel (keine explizite Stammfunktion)

In der Wahrscheinlichkeitsrechnung spielt die sogenannte Gaußsche Glockenkurve

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} \quad (7.49)$$

eine wichtige Rolle. Insbesondere Integrale der Form

$$\int_a^b \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} dx \quad (7.50)$$

werden benötigt. Da die Funktion stetig ist, besitzt sie eine Stammfunktion. Diese lässt sich jedoch nicht explizit als Formel darstellen. Während man früher die Werte der Integrale für $a = -\infty$ und viele verschiedene b in langen Tabellen erfasst hat, berechnet man diese heute mit dem Computer mittels numerischer Integration.

SymPy

SymPy kann Stammfunktionen bilden und auch bestimmte Integrale (über den Hauptsatz der Differential- und Integralrechnung) berechnen.

```
import sympy
```

```
x = sympy.symbols('x')

# Stammfunktion
F = sympy.integrate(x ** 2 * sympy.sin(x), x)
display(F)

# bestimmtes Integral
i = sympy.integrate(x ** 2 * sympy.sin(x), (x, 0, sympy.pi))
display(i)

# uneigentliches Integral
i = sympy.integrate(x / (x ** 2 + 1), (x, 0, sympy.oo))
display(i)
```

```
-x**2*cos(x) + 2*x*sin(x) + 2*cos(x)
```

```
-4 + pi**2
```

```
oo
```

Numerische Integration

Verfahren zur numerischen Integration sind in allen mathematischen Software-Bibliotheken vorhanden, z.B. auch in SciPy (als Ergänzung zu NumPy). Mit diesen Verfahren können nur bestimmte Integrale berechnet werden, keine Stammfunktionen. Grundidee ist das Annähern der zu berechnenden Fläche durch einfache Teilflächen (Dreiecke, Rechtecke, Trapeze,...), deren Flächeninhalt aus expliziten Formeln bestimmt werden kann. Der Ansatz ist also sehr ähnlich zur Definition von Flächeninhalten mittels Überdeckungen aus Quadraten, nur dass die Teilflächen je nach konkretem Verfahren vielfältiger sein können. Insbesondere kommen gelegentlich auch krummlinig begrenzte Teilflächen (z.B. Parabelstücke) zum Einsatz.

Hier ein Beispiel mit `scipy.integrate.quad`:

```
import numpy as np
from scipy.integrate import quad

# zu integrierende Funktion als Python -Funktion
f = lambda x: x ** 2 * np.sin(x)

# alternativ:
#def f(x):
#    return x ** 2 * np.sin(x)
```

```
# Integration von 0 bis pi  
quad(f, 0, np.pi)
```

```
(5.869604401089359, 6.516569951530559e -14)
```

Der erste zurückgegebene Wert ist der berechnete Wert des Integrals. Der zweite Wert ist eine Schätzung (!) des Fehlers im Vergleich zum exakten Wert des Integrals. Der Fehler liegt hier also sehr wahrscheinlich im Bereich von 10^{-13} , also erst in der 13. Nachkommastelle.

Im Gegensatz zur numerischen Differentiation arbeitet die numerische Integration sehr zuverlässig. Für stetige Funktion erhält man mit den Standardverfahren der Bibliotheken praktisch immer korrekte Ergebnisse. Vorsicht ist nur bei Polstellen oder anderweitigen Besonderheiten geboten.

8 Wahrscheinlichkeitsrechnung

Grundlagen der Wahrscheinlichkeitsrechnung und ausgewählte Wahrscheinlichkeitsverteilungen.

9 Statistik

Dieses Kapitel wird in den nächsten Wochen veröffentlicht...

10 Übung 1 (Grundlegendes)

10.1 Logik, Argumentation, Begriffsbildung

Aufgabe 1.1

Beantworten Sie die in Beispiel mit dem Hase und der Schildkröte gestellten Fragen.

Aufgabe 1.2

Diskutieren Sie die in Beispiel zum Begriff “gleich viel” gestellten Fragen. Geben Sie eine für plausible Definition für “zwei Mengen enthalten gleich viele Elemente” an, mit der sich die Fragen beantworten lassen. Hinweis: Hier gibt es verschiedene Antwortmöglichkeiten, die sich ggf. widersprechen, aber dennoch als richtig angesehen werden können.

Aufgabe 1.3

Diskutieren Sie die folgende Argumentation. Wo liegt der Fehler?

In einen Koffer passen unendlich viele Briefmarken, denn würden nur endlich viele hinein passen, so gäbe es also eine maximale Anzahl $n \in \mathbb{N}$ an Briefmarken, die in den Koffer passen. Jedoch würde man auch dann, wenn schon n Briefmarken im Koffer sind, noch irgendwo eine weitere Briefmarke einbringen können (kein Koffer kann so voll sein, dass keine Briefmarke mehr rein geht). Somit sind nun $n + 1$ Briefmarken im Koffer. Die Zahl n kann also doch nicht die Maximalanzahl sein. Die Annahme der endlichen Anzahl führt somit zum Widerspruch. Folglich passen unendlich viele Briefmarken in den Koffer.

10.2 Wiederholung Schulmathematik

Aufgabe 1.4

Lösen Sie den Eingangstest MINTFIT von OMB+. Dazu ist eine kostenlose und datensparsame Registrierung nötig.

Aufgabe 1.5

Gleich Sie die im Eingangstest identifizierten Defizite mit der Section 2.2, ab. Erstellen Sie einen Zeitplan für das selbstständige Aufarbeiten der relevanten Punkte im Laufe des Semesters.

11 Übung 2 (Python und Jupyter)

11.1 Jupyter

Aufgabe 2.1

Starten Sie durch Klick auf den Link ein browser-basiertes JupyterLab

Erstellen Sie ein neues Dokument (“Notebook”), berechnen Sie darin $123 * 456$ und speichern Sie das Dokument.

Das gespeicherte Dokument liegt nun im internen Speicherbereich Ihres Browsers. Um es bei Ihren Dateien zu speichern, können Sie den Download-Button verwenden.

Aufgabe 2.2

Starten Sie ein JupyterLab auf mybinder.org. Laden Sie die in der vorhergehenden Aufgabe gespeicherte Datei dort hoch und öffnen Sie sie.

Aufgabe 2.3

Machen Sie sich mit Jupyter-Notebooks vertraut. Fügen Sie neue Zellen ein. Probieren Sie auch den Zelltyp “Markdown” aus für Text, der nicht von Python verarbeitet werden soll.

Aufgabe 2.4 (optional)

Installieren Sie JupyterLab lokal auf Ihrem Computer (siehe z.B. Install Jupyter Locally in Data Science and Artificial Intelligence for Undergraduates).

11.2 Python

Aufgabe 2.5

Werten Sie das Polynom $x^3 + 2x^2 + 1$ an den Stellen 1, 3, 4, 6, 8 aus. Ergänzen Sie dazu den folgenden Code:

```
stellen = [1, 3, 4, 6, 8]
for x in stellen:
    # hier ergänzen
```

Aufgabe 2.6

Geben Sie alle durch 4 teilbaren Zahlen zwischen 9 und 99 aus. Verwenden Sie dazu eine `for`-Schleife.

```
# Ihre Lösung
```

Aufgabe 2.7

Zählen Sie die geraden Zahlen in einer gegebenen Liste.

```
liste = [2, 5, 4, 3, 3, 2, 0, 6, 7, 5, 6, 7]
# weitere Listen zum Testen:
#liste = [2, 4]
#liste = [1, 3]
#liste = []

# Ihre Lösung
```

12 Übung 3 (Aussagen, Abbildungen, Folgen)

12.1 Aussagen

Aufgabe 3.1 (KTA)

Bilden Sie zu folgenden Aussagen die genau gegenteiligen Aussagen. Dabei spielt es keine Rolle, ob die gegebenen Aussagen wahr oder falsch sind!

- (a) Alle Studierenden bestehen die Mathematikprüfung.
- (b) Alle Studierenden bestehen alle Prüfungen dieses Semesters.
- (c) Es gibt Studierende, die dieses Semester durch mindestens eine Prüfung fallen.
- (d) Wenn alle Studierenden die Prüfung bestanden haben, sind alle Studierenden glücklich.

Aufgabe 3.2 (KTA)

Skizzieren Sie die Menge

$$\{(x, y) \in \mathbb{R}^2 : y \geq 0, (x - 1)^2 + y^2 \leq 1, 2x + y \leq 2\}. \quad (12.1)$$

12.2 Abbildungen

Aufgabe 3.3 (KTA)

Entscheiden Sie für jede der folgenden Abbildungen, welche Abbildung injektiv/surjektiv/bijektiv ist. Ist keine Grundmenge für den Wertebereich angegeben, so ist die kanonische Menge zu wählen (bei (a) z.B. $\mathbb{N}$).

- (a) Jedem Insekt wird die Anzahl seiner Beine zugeordnet.
- (b) Jedem Land wird seine Nationalflagge zugeordnet.
- (c) Jedem Bundestagsabgeordneten wird ein Sitzplatz im Parlamentssaal zugeordnet.

12.3 Zahlenfolgen

Aufgabe 3.4 (KTA)

Zeigen Sie durch Berechnen der folgenden Grenzwerte, dass Summen der Form " $\infty - \infty$ " verschiedene Werte annehmen können und somit stets einer Einzelfallbetrachtung bedürfen.

12 Übung 3 (Aussagen, Abbildungen, Folgen)

- (a) $\lim_{n \rightarrow \infty} (n^2 - n)$
- (b) $\lim_{n \rightarrow \infty} (n^2 - n^3)$
- (c) $\lim_{n \rightarrow \infty} \left(n^2 - \frac{n^3}{n + \frac{1}{n}} \right)$

Aufgabe 3.5

Stellen Sie durch numerisches Auswerten der Folgenglieder

$$a_n = \left(\frac{(n-1)^2}{(2^3)^4 n + 1} \right)^{\frac{1}{50}} \quad (12.2)$$

eine Vermutung zum Konvergenzverhalten der Folge $(a_n)_{n \in \mathbb{N}}$ an. Überprüfen Sie Ihre Vermutung durch Berechnung des Grenzwerts mit SymPy.

Ihre Lösung

Aufgabe 3.6

Durch

$$\begin{aligned} x_1 &:= 0, \\ y_1 &:= 0, \\ x_{n+1} &:= x_n^2 - y_n^2 + \frac{1}{5}, \quad n \in \mathbb{N}, \\ y_{n+1} &:= 2x_n y_n + \frac{1}{2}, \quad n \in \mathbb{N}, \end{aligned} \quad (12.3)$$

ist eine Folge $((x_n, y_n))_{n \in \mathbb{N}}$ von Punkten in der Ebene gegeben.

- (a) Stellen Sie die ersten 100 Elemente der Punktfolge grafisch dar.
- (b) Stellen Sie eine Vermutung zum Konvergenzverhalten an, wenn wir unter Konvergenz einer Punktfolge die Konvergenz beider Koordinatenfolgen verstehen.
- (c) Formulieren Sie anhand der Definition des Konvergenzbegriffs für Zahlenfolgen die Bedingungen, die für den Nachweis der Konvergenz zu überprüfen wären.

Ihre Lösung

Aufgabe 3.7 (KTA)

Wird ein Kapital von 1000000 Euro nach einem Jahr mit 1 Prozent Zinsen vergütet, so stehen nach diesem Jahr

$$1000000 (1 + 0.01) = 1010000 \quad (12.4)$$

Euro zur Verfügung. Werden hingegen monatlich Zinsen zu $\frac{1}{12}$ des jährlichen Zinssatzes gezahlt, so sind nach Jahresende

$$1000000 \left(1 + \frac{0.01}{12}\right)^{12} \approx 1010045.96 \quad (12.5)$$

Euro vorhanden. Bei täglichlicher Verzinsung würde man

$$1000000 \left(1 + \frac{0.01}{365}\right)^{365} \approx 1010050.03 \quad (12.6)$$

Euro erhalten.

Lässt sich durch immer feinere Unterteilung des Jahres in Zinsabschnitte (Stunden, Minuten,...) das Endkapital beliebig erhöhen? Begründen Sie Ihre Antwort!

13 Übung 4 (Funktionen)

13.1 Definitionsbereich

Aufgabe 4.1 (KTA)

Bestimmen Sie die größtmögliche Menge $D \subseteq \mathbb{R}$, sodass die folgenden Zuordnungsvorschriften für alle $x \in D$ einen sinnvollen Wert liefern. Skizzieren Sie jeweils die Menge D .

- (a) $x \mapsto \frac{5}{(x-1)^2}$,
- (b) $x \mapsto \frac{x}{x}$,
- (c) $x \mapsto \sqrt{9 - |x|}$,
- (d) $x \mapsto \ln(x^2 - 2x + 2)$,
- (e) $x \mapsto \sqrt{4 - x} + \sqrt{x - 3}$.

Aufgabe 4.2 (KTA)

Bestimmen Sie die größtmögliche Menge $D \subseteq \mathbb{R}^2$, sodass die folgenden Zuordnungsvorschriften für alle $(x, y) \in D$ einen sinnvollen Wert liefern. Skizzieren Sie jeweils die Menge D .

- (a) $(x, y) \mapsto \ln(x^2 - y^2)$,
- (b) $(x, y) \mapsto \sqrt{y - 2x}$,
- (c) $(x, y) \mapsto (9 - x^2 - y^2)^{-2}$,
- (d) $(x, y) \mapsto \frac{\sqrt{y-x}}{y}$,
- (e) $(x, y) \mapsto e^{-x^2}$,
- (f) $(x, y) \mapsto \tan\left(\frac{1}{x+y}\right)$.

13.2 Visualisierung

Nebenbemerkung

Auch wenn es für das korrekte Lösen der folgenden Übungsaufgaben nicht zwingend nötig ist, so sei doch die Beschäftigung mit weiteren Funktionen des Matplotlib-Pakets empfohlen (Achseneinteilung, Gitterlinien, Farben, Linienarten, Legende,...). Einen Überblick über die Möglichkeiten bietet beispielsweise das Kapitel Matplotlib Basics in DSAI.

Aufgabe 4.3

Visualisieren Sie die Funktion $f : [-2\pi, 2\pi] \rightarrow \mathbb{R}$, $f(x) := x^2 \cos(10x)$.

Ihre Lösung

Aufgabe 4.4

Visualisieren Sie die Funktion $f : [-\frac{\pi}{8}, \frac{\pi}{8}] \rightarrow \mathbb{R}$, $f(x) := \sin \frac{1}{x}$.

Ihre Lösung

Aufgabe 4.5

Visualisieren Sie die durch

$$x(t) := \cos(10t), \quad y(t) = \sin(10t), \quad t \in [0, 2\pi] \quad (13.1)$$

gegebene Kurve im Raum als 3D-Plot.

Ihre Lösung

13.3 Grenzwerte und Stetigkeit**Aufgabe 4.6 (KTA)**

Zeigen Sie ausgehend von der Definition des Grenzwerts bei Funktionen, dass der Grenzwert

$$\lim_{x \rightarrow 0} \sin \frac{1}{x} \quad (13.2)$$

nicht existiert.

Aufgabe 4.7

Ermitteln Sie mit SymPy, ob folgende Grenzwerte existieren und welchen Wert sie ggf. haben.

Hinweis: Das Berechnen von Funktionsgrenzwerten mit SymPy funktioniert genauso wie das Berechnen von Folgengrenzwerten. Allerdings liefert SymPy stets nur den rechtsseitigen Grenzwert. Mit `sympy.limit(..., x, x0, '-')` erhält man den linksseitigen Grenzwert. Die direkte Berechnung des beidseitigen Grenzwerts wird nicht unterstützt.

- (a) $\lim_{x \rightarrow 2} \frac{(x-2)^2}{x^2 - 4}$,

- (b) $\lim_{x \rightarrow 0} \frac{x}{|x|}$,
- (c) $\lim_{x \rightarrow 0} \sin \frac{1}{x}$.

Aufgabe 4.8

Implementieren Sie die Funktion

$$f : \mathbb{R} \rightarrow \mathbb{R}, \quad f(x) := \begin{cases} -1 - x^2, & \text{falls } x < 0, \\ 0, & \text{falls } x = 0, \\ 1 + x^2, & \text{falls } x > 0 \end{cases} \quad (13.3)$$

als Python-Funktion. Testen Sie Ihre Funktion mit dem folgenden Code:

```
def f(x):
    # Hier Code einfügen

print(f'f( -2) = {f( -2)}')
print(f'f( -1) = {f( -1)}')
print(f'f(0) = {f(0)}')
print(f'f(1) = {f(1)}')
print(f'f(2) = {f(2)}')
print(f'f(0) = {f(0.1 ** 2 - 0.01)}')
```

Warum wird trotz korrekter Implementierung ein falscher Funktionswert in der letzten Zeile ausgegeben? Kann dies für stetige Funktionen ebenfalls passieren?

Zusatz: Visualisieren Sie die Funktion f . Achten Sie insbesondere auf die korrekte Darstellung für $x \approx 0$.

```
# Ihre Lösung
```

14 Übung 5 (Matrizen, Vektorgeometrie)

14.1 Rechnen mit Matrizen

Aufgabe 5.1 (KTA)

Geben Sie zu folgenden Ausdrücken die Größe der Ergebnismatrix an:

- (a) $A B C$, $A \in \mathbb{R}^{3 \times 7}$, $B \in \mathbb{R}^{7 \times 5}$, $C \in \mathbb{R}^{5 \times 4}$,
- (b) $(A B^T)^T C$, $A \in \mathbb{R}^{3 \times 4}$, $B \in \mathbb{R}^{5 \times 4}$, $C \in \mathbb{R}^{3 \times 4}$,
- (c) $x^T y$ $x \in \mathbb{R}^n$, $y \in \mathbb{R}^n$,
- (d) $x y^T$ $x \in \mathbb{R}^n$, $y \in \mathbb{R}^n$,
- (e) $x^T A x$ $x \in \mathbb{R}^n$, $A \in \mathbb{R}^{n \times n}$.

Aufgabe 5.2 (KTA)

Berechnen Sie die folgenden Produkte und Summen:

- (a) $\begin{bmatrix} 2 & 3 & 1 \\ 3 & 2 & 1 \end{bmatrix} \begin{bmatrix} 4 & 5 \\ 1 & 1 \\ 2 & 0 \end{bmatrix}$,
- (b) $\begin{bmatrix} 2 & 3 & 1 \\ 3 & 2 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 2 \\ 0 \end{bmatrix} + \begin{bmatrix} 2 & 3 & 1 \\ 3 & 2 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \\ 2 \end{bmatrix}$,
- (c) $\begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix} [4 \ 5 \ 6]$.

Aufgabe 5.3 (KTA)

Rechnen Sie die Beziehung $(A B)^T = B^T A^T$ für die beiden Matrizen

$$A = \begin{bmatrix} 2 & 3 & 1 \\ 3 & 2 & 1 \end{bmatrix}, \quad B = \begin{bmatrix} 4 & 5 \\ 1 & 1 \\ 2 & 0 \end{bmatrix} \quad (14.1)$$

nach.

Aufgabe 5.4 (KTA)

Berechnen Sie die Determinanten:

- (a) $\begin{vmatrix} 1 & 2 \\ 3 & 4 \end{vmatrix},$
- (b) $\begin{vmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{vmatrix},$
- (c) $\begin{vmatrix} 2 & 0 & 1 \\ 3 & 2 & 0 \\ 1 & 1 & 1 \end{vmatrix}.$

14.2 Rechnen mit Vektoren

Aufgabe 5.5 (KTA)

Die Eckpunkte eines Dreiecks im Raum sind $A = (2, 3, 4)$, $B = (3, 1, 7)$, $C = (1, 8, 2)$. Berechnen Sie die Länge der Seite $\overline{AB}$ sowie den Winkel am Punkt B .

Aufgabe 5.6 (KTA)

Im Raum sind die drei Punkte $A(1, -2, 3)$, $B = (3, 2, 1)$, $C = (6, 4, 4)$ gegeben. Ermitteln Sie eine Punkt D , sodass die Punktfolge A, B, C, D ein Parallelogramm beschreibt.

Aufgabe 5.7 (KTA)

Seien $a, b \in \mathbb{R}^3$ zwei linear unabhängige Vektoren. Geben Sie eine Formel zur Berechnung eines Vektors $w \in \mathbb{R}^3$ entlang der Winkelhalbierenden zwischen a und b an.

14.3 Skalarprodukt und orthogonale Projektion

Aufgabe 5.8 (KTA)

Finden Sie alle Vektoren, die orthogonal zum Vektor $\begin{bmatrix} 3 & -1 & 2 \end{bmatrix}^T$ sind.

Aufgabe 5.9 (KTA)

Berechnen Sie $\mathbb{R}^5$ die orthogonale Projektion von $\begin{bmatrix} 1 & 0 & 1 & -1 & 2 \end{bmatrix}^T$ auf $\begin{bmatrix} 0 & 2 & -1 & -2 & 0 \end{bmatrix}^T$.

Aufgabe 5.10 (KTA)

Von einem Eckpunkt eines Quadrates werden Geraden durch die Mittelpunkte der beiden dem Eckpunkt gegenüberliegenden Seiten gezogen. Wie groß ist der Winkel zwischen den

beiden Geraden?

Aufgabe 5.11 (KTA)

Eine Kraft vom Betrag 60 N, die in Richtung des Vektors $\begin{bmatrix} -1 & 3 & 2 \end{bmatrix}^T$ wirkt, bewegt einen Körper vom Punkt $(2, -2, 1)$ geradlinig zum Punkt $(3, 4, 5)$ (Ortsangaben in Meter). Welche Arbeit wird dabei verrichtet?

15 Übung 6 (lineare Abbildungen, Basen)

15.1 Lineare Abbildungen

Aufgabe 6.1 (KTA)

Sind folgende Abbildungen linear? Begründen Sie Ihre Entscheidung.

- (a) $f : \mathbb{R}^3 \rightarrow \mathbb{R}, x \mapsto x_1$,
- (b) $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2, (x, y) \mapsto (-y, x)$,
- (c) $f : \mathbb{R}^2 \rightarrow \mathbb{R}, (x, y) \mapsto x y$,
- (d) $f : \mathbb{R}^2 \rightarrow \mathbb{R}, (x, y) \mapsto x + y$,
- (e) $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2, (x, y) \mapsto (x + 1, y - 1)$.

Aufgabe 6.2 (KTA)

Geben Sie zu folgenden linearen Abbildungen die Matrixdarstellung an:

- (a) $f : \mathbb{R}^3 \rightarrow \mathbb{R}^3$, 90°-Drehung um die z -Achse gegen die Uhr (bei Betrachtung vom Punkt $(0, 0, 1)$ aus),
- (b) $f : \mathbb{R}^3 \rightarrow \mathbb{R}$, Bilden der Summe aller Komponenten,
- (c) $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, Drehung um $\frac{\pi}{4}$ im Uhrzeigersinn,
- (d) $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, orthogonale Projektion auf den Vektor $\begin{bmatrix} 1 & 1 \end{bmatrix}^T$.

15.2 Unterräume und Basen

Aufgabe 6.3 (KTA)

Zeigen Sie, dass die Menge

$$U := \{x \in \mathbb{R}^4 : x_3 = 2x_1 - x_2\} \quad (15.1)$$

ein linearer Unterraum von $\mathbb{R}^4$ ist.

Welche Dimension hat U ? Geben Sie eine Basis von U an.

Aufgabe 6.4 (KTA)

Begründen Sie anhand der Definition der linearen Unabhängigkeit, dass die Menge

$$\left\{ \begin{bmatrix} 2 \\ 0 \\ 4 \end{bmatrix}, \begin{bmatrix} 1 \\ -1 \\ 2 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix} \right\} \quad (15.2)$$

linear abhängig ist.

Aufgabe 6.5 (KTA)

Sind die folgenden Aussagen wahr oder falsch? Begründen Sie jeweils kurz Ihre Antwort.

- (a) Jede ONB ist ein ONS.
- (b) Jedes ONS ist eine Basis in seiner Aufspannung.
- (c) Jede Basis in $\mathbb{R}^2$ ist eine ONB.
- (d) Ist $\{b_1, b_2\}$ eine ONB in einem zweidimensionalen Unterraum des $\mathbb{R}^n$, so ist die orthogonale Projektion von b_1 auf b_2 stets der Nullvektor.

16 Übung 7 (lineare Gleichungssysteme)

16.1 Formulieren von LGS

Aufgabe 7.1 (KTA)

Bringen Sie das folgende lineare Gleichungssystem in Matrix-Vektor-Schreibweise:

$$\begin{aligned} 3x_1 + x_3 &= 2x_2 + 2 \\ 2x_3 + x_1 + x_2 &= 1 + x_1 \\ 5 - x_2 - x_1 &= 5 + x_1. \end{aligned} \tag{16.1}$$

Aufgabe 7.2 (KTA)

An der Innenseite einer Hauswand liegt die Oberflächentemperatur bei $U_i = 22$ [°C]. Die Temperatur an der Außenseite beträgt $U_a = -15$ [°C]. Zur (näherungsweisen) Berechnung des Temperaturverlaufs im Inneren der Wand wird die Wand in n Schichten unterteilt (parallel zur Wandoberfläche). Anhand physikalischen Gesetzmäßigkeiten kann man sich überlegen, dass die Temperatur u_k in der k -ten Schicht dann gerade der Mittelwert aus den Temperaturen u_{k-1} und u_{k+1} der beiden benachbarten Schichten ist. Stellen Sie ein LGS zur Berechnung der Temperaturen $u_2, \dots, u_{n-1}$ für $n = 7$ auf. Die Temperaturen in den äußeren Schichten entsprechen den Oberflächentemperaturen, also $u_1 = U_i$ und $u_n = U_a$, sind also bekannt.

Aufgabe 7.3

Lösen Sie das LGS aus der vorhergehenden Aufgabe mit NumPy und visualisieren Sie das Ergebnis. Gestalten Sie den Quellcode so, dass Sie die Schichtanzahl n leicht verändern können.

Probieren Sie verschiedene Werte für n aus. Testen Sie Ihren Code auch mit verschiedenen Oberflächentemperaturen. Leiten Sie aus der Simulation eine Vermutung über die allgemeine Form des Temperaturverlaufs in der Wand ab.

Ihre Lösung

16.2 Lineare Unabhängigkeit

Aufgabe 7.4

Sind die Vektoren

$$\begin{bmatrix} 1 \\ 7 \\ 2 \end{bmatrix}, \begin{bmatrix} 3 \\ 1 \\ 2 \end{bmatrix}, \begin{bmatrix} 2 \\ -1 \\ 1 \end{bmatrix} \quad (16.2)$$

linear unabhängig? Formulieren Sie zur Klärung dieser Frage ein geeignetes LGS und lösen Sie dieses mit dem Computer.

Ihre Lösung

Aufgabe 7.5

Lösen Sie die folgenden Teilaufgaben mit dem Computer.

- (a) Zeigen Sie, dass

$$u := \begin{bmatrix} 2 \\ 1 \\ -1 \\ 0 \end{bmatrix}, \quad v := \begin{bmatrix} 0 \\ -3 \\ 2 \\ -4 \end{bmatrix}, \quad w := \begin{bmatrix} 1 \\ -2 \\ 4 \\ 1 \end{bmatrix} \quad (16.3)$$

linear unabhängig sind.

- (b) Welcher der Vektoren

$$x := \begin{bmatrix} 1 \\ 0 \\ -3 \\ -5 \end{bmatrix}, \quad y := \begin{bmatrix} 2 \\ -1 \\ 6 \\ 6 \end{bmatrix}, \quad z := \begin{bmatrix} -1 \\ 1 \\ 0 \\ 5 \end{bmatrix} \quad (16.4)$$

ist geeignet, zusammen mit u, v, w eine Basis in $\mathbb{R}^4$ zu bilden?

- (c) Ermitteln Sie die Koordinaten des Vektors $\begin{bmatrix} 0 \\ 1 \\ -3 \\ 0 \end{bmatrix}$ bzgl. der Basis aus (b).

Ihre Lösung

16.3 Geraden und Ebenen

Aufgabe 7.6 (KTA)

Bestimmen Sie den Mittelpunkt und den Radius des Kreises, der bei Rotation des Punktes $(-2, 1, 1)$ um die Gerade durch

$$\begin{bmatrix} x \\ y \\ z \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} + t \begin{bmatrix} -1 \\ 1 \\ 1 \end{bmatrix}, \quad t \in \mathbb{R}, \quad (16.5)$$

beschriebene Gerade erzeugt wird.

Aufgabe 7.7 (KTA)

Bestimmen Sie die Koordinaten des Spiegelpunktes von $(-7, 9, 1)$ bezüglich der Ebene durch $5x - 4y - z = 12$ beschriebenen Ebene.

17 Übung 8 (Koordinatentransformationen)

17.1 Aufgabe 8.1 (KTA)

Die zwei Vektoren $a := \begin{bmatrix} 1 & 1 & 0 \end{bmatrix}^T$ und $b := \begin{bmatrix} 2 & 1 & 1 \end{bmatrix}^T$ spannen eine Ebene im $\mathbb{R}^3$ auf. Finden Sie eine ONB $\{u, v, w\}$, sodass u und v in dieser Ebene liegen und u die gleiche Richtung hat wie a .

Hinweis: Nutzen Sie das Kreuzprodukt von Vektoren. Finden Sie zunächst u , dann w und dann v .

17.2 Aufgabe 8.2 (KTA)

Im $\mathbb{R}^2$ wird ein neues Koordinatensystem durch Drehen des Standardkoordinatensystems um 30° gegen den Uhrzeigersinn und anschließendes Verschieben des Ursprungs in den Punkt $(2, 1)$ erzeugt.

- (a) Geben Sie die Transformationvorschrift zum Erhalt der neuen Koordinaten aus den ursprünglichen Koordinaten an.
- (b) Welche Koordinaten hat der Punkt $(1, 1)$ bzgl. des neuen Koordinatensystems?

17.3 Aufgabe 8.3 (KTA)

Im $\mathbb{R}^2$ wird ein neues Koordinatensystem durch Drehen des Standardkoordinatensystems um 60° im Uhrzeigersinn und anschließendes Verschieben des Ursprungs in den Punkt $(1, 3)$ erzeugt. Geben Sie die Transformationsvorschrift für die Berechnung der neuen Koordinaten aus den Koordinaten bzgl. des Koordinatensystems aus der vorhergehenden Aufgabe an.

18 Übung 9 (Differentialrechnung)

18.1 Aufgabe 9.1 (KTA)

Ermitteln Sie die Ableitungen folgender Funktionen manuell.

- (a) $f(x) = 5x^4 - 2x^3 + x + 4$,
- (b) $(x) = (\sin x) \cos x$,
- (c) $f(x) = \frac{\sin x}{x^2+2}$,
- (d) $f(x) = (\cos x)^2$,
- (e) $f(x) = (\sin x)^2 (x^3 + x)$.

18.2 Aufgabe 9.2

Überprüfen Sie Ihre Ergebnisse in der vorhergehenden Aufgabe mit SymPy.

18.3 Aufgabe 9.3 (KTA)

Stellen Sie das Taylor-Polynom vom Grad 3 für die Cosinus-Funktion an der Entwicklungsstelle $x_0 = 0$ auf.

Geben Sie auch die maximale Abweichung zwischen der Cosinus-Funktion und dem gefundenen Taylor-Polynom auf dem Intervall $[0, \frac{\pi}{4}]$ an. (Lösung: rund 0.191)

18.4 Aufgabe 9.4 (KTA)

Zur Bestimmung des Durchmessers einer Kugel wurde zunächst deren Volumen bestimmt, indem ihre Verdrängung in einem Wasserbehälter gemessen wurde. Die Messung ergab ein Volumen von 1.45 dm^3 bei einer Messunsicherheit von $\pm 0.005 \text{ dm}^3$. Aus diesem Wert wurde anschließend der Kugeldurchmesser mit der Formel

$$V = \frac{\pi}{6} r^3 \quad (18.1)$$

berechnet. Wie groß ist der zu erwartende relative Fehler des berechneten Durchmessers? (Lösung: rund 0.11%)

18.5 Aufgabe 9.5 (KTA)

Sollen geografische Breite und Länge in Koordinaten im $\mathbb{R}^3$ umgerechnet werden, so ergibt sich die z -Koordinate (in Kilometer) aus der Breite ζ (in Grad) mittels

$$z = 6371 \cos\left(\frac{\pi}{2} - \frac{\pi}{180} \zeta\right). \quad (18.2)$$

(vgl. Section 5.6). Die Breite eines zu transformierenden Punktes liegt bei $\zeta_0 = 53.1$, wobei der exakte Wert auf eine Nachkommastelle gerundet wurde. Wie groß ist der durch das Runden zu erwartende absolute Fehler in der z -Koordinate? (Lösung: rund 3.34 km)

19 Übung 10 (Integralrechnung)

19.1 Stammfunktionen

Aufgabe 10.1 (KTA)

Finden Sie alle (!) Stammfunktionen zu folgenden Funktionen.

- (a) $f(x) = 3x^4 - 4x^2 + 2$,
- (b) $f(x) = \sqrt{x}$, $x > 0$,
- (c) $f(x) = \frac{1}{x^2}$, $x > 0$,
- (d) $f(x) = 5 \sin x - \cos x$,
- (e) $f(x) = e^{2x}$.

Aufgabe 10.2 (KTA)

Für welche $c \in \mathbb{R}$ konvergiert das uneigentliche Integral

$$\int_1^\infty x^c dx? \quad (19.1)$$

Aufgabe 10.3 (KTA)

Geben Sie eine Stammfunktion zu

$$f(x) = \frac{1}{(x^2 + 1)^3} \quad (19.2)$$

an. Nutzen Sie dazu eine Formelsammlung Ihrer Wahl, z.B. die Tabelle von Ableitungs- und Stammfunktionen bei Wikipedia.

Aufgabe 10.4 (KTA)

Finden Sie eine Stammfunktion zu

$$f(x) = x^2 + \frac{1}{x}, \quad x > 0, \quad (19.3)$$

die durch den Punkt $(1, 2)$ verläuft.

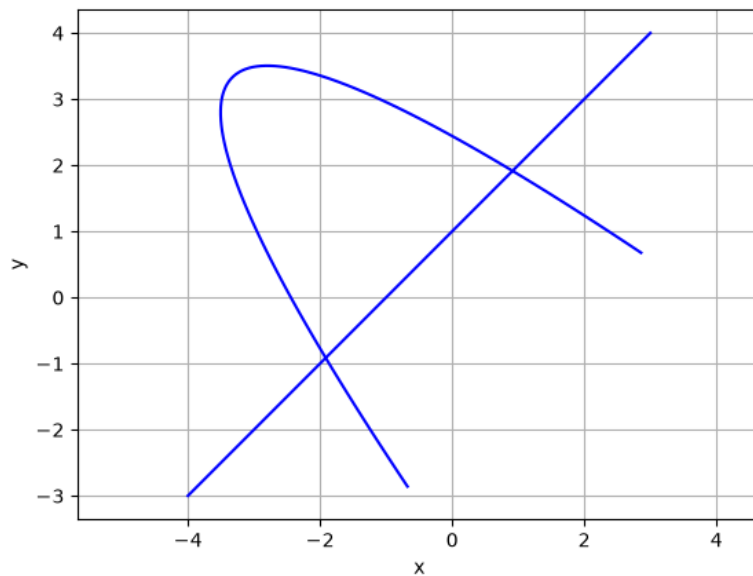
19.2 Flächeninhalte

Aufgabe 10.5 (KTA)

Die durch

$$(x + y)^2 + \sqrt{2}(y - x) - 8 - \sqrt{2} = 0 \quad (19.4)$$

beschriebene Kurve in der Ebene begrenzt gemeinsam mit der Gerade $y = x + 1$ eine Fläche.



- (a) Überzeugen Sie sich, dass die Kurve symmetrisch zur Gerade $y = -x$ verläuft, indem Sie die Kurve an dieser Gerade spiegeln (Koordinatentransformation!) und die Formel in den neuen Koordinaten mit der ursprünglichen Formel vergleichen.
- (b) Drehen Sie die Fläche um 45° im Uhrzeigersinn (Koordinatentransformation!). Die Formeln zur Beschreibung der Fläche sollten nun einfacher sein.
- (c) Beschreiben Sie die Fläche als Normalbereich und berechnen Sie den Flächeninhalt. (Lösung: rund 4.5523)

19.3 Aufgabe 10.6

Berechnen Sie die Fläche zwischen dem Graph der Funktion

$$f : [0, 2] \rightarrow \mathbb{R}, \quad f(x) := \sqrt{x} + x^2 \quad (19.5)$$

und der x -Achse auf folgende Arten und vergleichen sie die Ergebnisse.

- (a) manuelles Rechnen mittels Hauptsatz der Differential- und Integralrechnung,

19 Übung 10 (Integralrechnung)

- (b) symbolisches Rechnen mit SymPy,
- (c) numerisches Rechnen SciPy.

Ihre Lösung

20 Download

Zum Offline-Lesen steht das Vorlesungsskript hier als PDF-Datei zur Verfügung. Da dieses Vorlesungsskript dem **Online-First-Prinzip** folgt, wird die PDF-Datei nur nachrangig gepflegt. Die Inhalte sind zwar stets identisch mit den Online-Inhalten, allerdings gibt es aus technischen Gründen einige eher formale Abweichungen:

- Animationen sind nicht enthalten, sondern fehlen ohne ersichtlichen Hinweis auf das Fehlen.
- Grafiken sind teils an anderen Stellen im Text als in der Online-Version (plus minus eine Seite).
- Einige Links, z.B. zu den Videos und den zugehörigen Notizen, funktionieren nicht.
- Links zum Glossar der Abkürzungen funktionieren nicht.
- Das Layout enthält diverse kleinere Fehler.

Download PDF-Datei

21 Organisatorisches

Modulbeschreibung in Modulux

21.1 Ablauf der Veranstaltung

Die Veranstaltung findet im Wesentlichen im **Selbststudium** statt. Zur Unterstützung des Selbststudiums wird es Konsultationseinheiten geben:

- zwei “große” Präsenzveranstaltungen in Dresden,
- eine “große” Online-Veranstaltung,
- sechs “kleine” Online-Veranstaltungen.

Die genauen Termine finden Sie unten im Section 21.4. Den Link zum BBB-Raum für die Online-Veranstaltungen finden Sie im OPAL-Kurs. Organisatorische Fragen können Sie jederzeit per E-Mail an den Vorlesenden richten.

Merke

Die Konsultationseinheiten dienen der Klärung von Fragen zu Stoff und Übungsaufgaben. In den Einheiten wird keine Vorlesung stattfinden, also insbesondere auch kein neuer Stoff vermittelt werden. Für eine effiziente und zielführende Nutzung der Zeit ist die **gute Vorbereitung Ihrerseits** essentiell!

Hinweise zur Vorbereitung auf die Konsultationseinheiten:

- Stoff und Übungsaufgaben entsprechend Arbeitsplan **vor (!) den Konsultationseinheiten vorbereiten.**
- **Kleinere Fragen** zu Stoff und Übungsaufgaben können Sie jederzeit gern im Forum des OPAL-Kurses stellen. Dort können sowohl der Vorlesende als auch Mitstudierende antworten. Per E-Mail bitte nur Fragen stellen, die für andere Studierende nicht von Interesse sind oder aus Datenschutzgründen nicht für andere Studierende sichtbar sein sollen.
- **Umfangreichere Fragen** zu Stoff und Übungsaufgaben bitte möglichst schon vor der nächsten Konsultationseinheit per E-Mail an den Vorlesenden schicken, damit dieser den Ablauf der Konsultationseinheit entsprechend den zu klärenden Fragen strukturieren kann.
- Bis 72 Stunden vor der nächsten Konsultationseinheit können Sie gern Ihre **Lösungen und Lösungswege per E-Mail** an den Vorlesenden schicken (nur PDF-Dateien). Dieser wird Ihre Einsendungen korrigieren und mit Kommentaren versehen. Später eingegangene Lösungen oder Lösungen zu Übungsaufgaben, die laut Arbeitsplan schon zu einer früheren Konsultationseinheit hätten bearbeitet sein

sollen, werden nicht korrigiert! Auch Zusendungen von JPEG- und anderen Nicht-PDF-Dateien werden ignoriert. Dieses Angebot können Sie ganz ungezwungen auf freiwilliger Basis nutzen (oder auch nicht); für die Modulnote sind ausschließlich die Klausurergebnisse relevant.

Einige Übungsaufgaben sind mit dem Computer zu bearbeiten. Dazu ist keine spezielle Software nötig, sondern nur ein Webbrowser. Genaueres klären wir Übung 2.

Merke

Für den erfolgreichen Abschluss dieses Moduls bei durchschnittlicher Begabung ist eine Arbeitszeit von ca. **10 Stunden pro Woche** (bei 15 Vorlesungswochen) und zusätzlich ca. 30 Stunden vor der Prüfung anzusetzen (vgl. Modulbeschreibung). Etwa die Hälfte dieser Zeit sollte in die Erarbeitung neuen Stoffes fließen, die andere Hälfte in die Bearbeitung von Übungsaufgaben.

Behalten Sie bitte stets im Hinterkopf: Die Tätigkeit, für die Sie sich entschieden haben, heißt **“Studieren”**, nicht **“studiert werden”**. Das ist also eine Tätigkeit, die Sie (!) aktiv ausführen. Der Vorlesende wird Sie dabei so gut wie möglich unterstützen, kann aber keine Wunder vollbringen, wenn es an Eigeninitiative und Fleiß fehlt.

21.2 Hinweise zum Skript

Das Skript enthält den gesamten für die Prüfung relevanten Stoff. Wird auf externe Inhalte verlinkt, so sind diese Inhalte nicht relevant für die Prüfung. Abschnitte des Skripts, die nicht für die Prüfung relevant sind, sind entsprechend gekennzeichnet.

Das Skript steht auch als PDF-Datei zum Download zur Verfügung. Falls Sie **kapitelweise PDF-Dateien** nutzen möchten (z.B. zum Anbringen von Notizen), können Sie die Druckfunktion Ihres Webbrowsers nutzen. Firefox schneidet allerdings Textzeilen an den Seitenumbrüchen ab; bei Chromium gibt es dieses Problem nicht.

Diverse Rechnungen und Herleitungen fehlen im Skript und werden jeweils als **Video** zur Verfügung gestellt, das an entsprechender Stelle verlinkt ist.

Die Übungsaufgaben stehen zusätzlich als PDF-Aufgabenblatt und als **Jupyter-Notebook** zur Verfügung; siehe Button oben rechts bei jeder Übung. Die Jupyter-Notebooks erlauben das Lösen der Aufgaben direkt im Dokument (vgl. Übung 2).

Merke

Das Skript enthält viel Text. **Einmaliges Lesen ist nicht genug**. Lesen Sie einen Abschnitt zunächst vollständig, damit Sie einen Überblick über das aktuelle Thema bekommen. Beginnen Sie dann nochmals von vorn und durchdenken Sie

jeden Satz und jede Formel! Wenn Verständnisprobleme auftreten: weiter darüber nachdenken! Wenn das nicht hilft: Frage im Forum stellen.

Die mathematische Sprache (in Text und Formeln) kann komplexe Sachverhalte sehr komprimiert darstellen. Da kann es schonmal vorkommen, dass man an einem unscheinbaren Dreizeiler eine halbe Stunde oder mehr zu “knabbern” hat. Anders ausgedrückt: **Eine Mathebuch ist kein Rosamunde-Pilcher-Roman!**

Nebenbemerkung

Das Skript liegt zum Semesterbeginn noch nicht vollständig vor, wird also im Laufe des Semesters noch ergänzt werden. Vorgenommene Änderungen können Sie unter Section 21.5 nachvollziehen.

21.3 Prüfung

Die Prüfung findet als 90-minütige Klausur in der ersten Präsenzeinheit des zweiten Semesters statt (Mitte April 2027).

Übungsaufgaben, die von Typ und Schwierigkeit den Klausuraufgaben entsprechen, sind mit “KTA” (= klausurtypische Aufgabe) gekennzeichnet.

Eine Beispielklausur wird im Laufe des Semesters hier veröffentlicht.

Zugelassene Hilfsmittel zur Klausur: Papier, Stifte (kein Rot), einfacher Taschenrechner, also ohne CAS und ohne Grafik (die Aufgaben sind auch ohne Taschenrechner gut lösbar).

21.4 Arbeitsplan

Hier finden Sie alle Konsultationstermine und die bis zum jeweiligen Termin zu bearbeitenden Inhalte.

- **KW 42 (Termin wird noch präzisiert), online:**

- Schulmathematik
- Motivation
- Übung 1 (Grundlegendes)
- Computereinsatz
- Übung 2 (Python und Jupyter)
- Elementare Begriffe: Aussagen
- Elementare Begriffe: Abbildungen

- **KW 44 (Termin wird noch präzisiert), online:**
 - Elementare Begriffe: Zahlenfolgen
 - Übung 3 (Aussagen, Abbildungen, Folgen)
 - Elementare Begriffe: Funktionen
 - Übung 4 (Funktionen)
- **12.11.2026-14.11.2026 (Termin wird noch präzisiert), Dresden:**
 - Vektoren und Matrizen: Formalismus
 - Vektoren und Matrizen: Vektorgeometrie
 - Übung 5 (Matrizen, Vektorgeometrie)
 - Vektoren und Matrizen: Lineare Abbildungen
 - Vektoren und Matrizen: Lineare Unterräume und Untermannigfaltigkeiten
 - Übung 6 (lineare Abbildungen, Basen)
- **KW 48 (Termin wird noch präzisiert), online:**
 - Vektoren und Matrizen: Lineare Gleichungssysteme
 - Übung 7 (lineare Gleichungssysteme)
 - Vektoren und Matrizen: Koordinatentransformationen
 - Übung 8 (Koordinatentransformationen)
 - Vektoren und Matrizen: Ausblick: Eigenwerte und Eigenvektoren
- **KW 50 (Termin wird noch präzisiert), online:**
 - Differentialrechnung
 - Übung 9 (Differentialrechnung)
 - Integralrechnung
 - Übung 10 (Integralrechnung)
- **18.12.2026-19.12.2026 (Termin wird noch präzisiert), online:**
 - Wiederholen Sie bitte allen bisherigen Stoff und sammeln Sie Unklarheiten und Fragen.
- **KW 1 (Termin wird noch präzisiert), online:**
 - Zu bearbeitende Inhalte werden später bekanntgegeben...
- **KW 3 (Termin wird noch präzisiert), online:**
 - Zu bearbeitende Inhalte werden später bekanntgegeben...

- **04.02.2027-06.02.2027 (Termin wird noch präzisiert), Dresden:**
 - Zu bearbeitende Inhalte werden später bekanntgegeben...

21.5 Änderungshistorie

Da dieses Skript für das Wintersemester 2026/2027 vollständig neu strukturiert und formuliert wurde, werden im Laufe des Semesters sehr wahrscheinlich Korrekturen und Ergänzungen nötig werden. Der besseren Nachvollziehbarkeit halber werden die inhaltlich relevanten Anpassungen nach Semesterbeginn hier dokumentiert:

- 16.09.2026:
 - Kapitel zu Differential- und Integralrechnung hinzugefügt.
 - Ergebnis in Beispiel bei Zahlenfolgen ergänzt.
 - Abschnitt Section 4.3 eingefügt.
 - Prüfungstermin korrigiert.
- 28.08.2026:
 - Link zu Video 130 repariert.
 - Kapitel VI von OMB+ in die Liste der als aus der Schulmathematik bekannt vorausgesetzten Themen verschoben.
 - Nummerierung der Aufgaben in Übung 1 und Übung 2 korrigiert.
 - Nebenbemerkung zu Symbolen für logische Verknüpfungen eingefügt.
 - Unterscheidung zwischen Periode und Grundperiode bei Zahlenfolgen eingefügt.
 - Unterscheidung zwischen Section 4.4 bei Funktionen eingefügt.
 - Diverse sprachliche Korrekturen.

