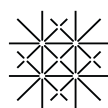


Frühjahrssemester 2022

Mathematik II für Naturwissenschaften

Christine Zehrt-Liebendörfer



Universität
Basel

Inhaltsverzeichnis

1	Beschreibende Statistik	1
1.1	Grundbegriffe	1
1.2	Häufigkeitsverteilung	3
1.3	Mittelwerte	5
1.4	Quantile und Boxplot	8
1.5	Empirische Varianz und Standardabweichung	13
1.6	Prozentrechnen	15
2	Korrelation und Regressionsgerade	18
2.1	Der Korrelationskoeffizient	18
2.2	Rangkorrelation	23
2.3	Die Regressionsgerade	24
2.4	Nichtlineare Regression	28
3	Wahrscheinlichkeitsrechnung	30
3.1	Zufallsexperimente und Ereignisse	30
3.2	Wahrscheinlichkeit	31
3.3	Bedingte Wahrscheinlichkeit	34
3.4	Unabhängige Ereignisse	37
4	Erwartungswert und Varianz von Zufallsgrößen	39
4.1	Zufallsgrösse und Erwartungswert	39
4.2	Varianz und Standardabweichung	40
4.3	Kombination von Zufallsgrößen	43
4.4	Schätzen von Erwartungswert und Varianz	45
5	Binomial- und Poissonverteilung	47
5.1	Die Binomialverteilung	47
5.2	Die Poissonverteilung	51
6	Die Normalverteilung	54
6.1	Eigenschaften der Glockenkurve	55
6.2	Approximation der Binomialverteilung	57
6.3	Normalverteilte Zufallsgrößen	59
6.4	Der zentrale Grenzwertsatz	63
7	Statistische Testverfahren	64
7.1	Testen von Hypothesen	64
7.2	Der t -Test für Mittelwerte	69
7.3	Der Varianzenquotiententest	74
7.4	Korrelationsanalyse	75
7.5	Der χ^2 -Test	78
7.6	Vertrauensintervall für eine Wahrscheinlichkeit	81

8	Lineare Abbildungen	84
8.1	Definition und Beispiele	84
8.2	Eigenschaften	89
8.3	Basiswechsel	90
8.4	Bedeutung der Determinante einer Darstellungsmatrix	93
9	Eigenwerte und Eigenvektoren	96
9.1	Bestimmung von Eigenwerten und Eigenvektoren	96
9.2	Diagonalisierung von Matrizen	100
9.3	Symmetrische Matrizen	107
9.4	Komplexe Matrizen	112
10	Differentialrechnung für Funktionen in mehreren Variablen	114
10.1	Graphische Darstellung	114
10.2	Partielle Ableitungen und Tangentialebenen	117
10.3	Richtungsableitung, Gradient und Hesse-Matrix	123
10.4	Lokale und globale Extrema	127
10.5	Extrema mit Nebenbedingung	132
11	Vektorfelder und Wegintegrale	136
11.1	Vektorfelder	136
11.2	Wege und Kurven	141
11.3	Wegintegrale	143
12	Integration in mehreren Variablen	149
12.1	Bereichsintegrale	149
12.2	Koordinatentransformationen	152
12.3	Flächenintegrale	154
12.4	Integralsätze	157

Dieses Skript basiert auf dem Vorlesungsskript von Hans Walser vom HS13/FS14, auf dem Vorlesungsskript von Hans-Christoph Im Hof und Hanspeter Kraft vom WS01/SS02 und auf eigenen Skripten von früheren Vorlesungen. Einige Graphiken wurden von Thomas Zehrt angefertigt. Einzelne Beispiele stammen von Büchern der Literaturliste.

16.02.2022

Christine Zehrt-Liebendörfer

Departement Mathematik und Informatik

Spiegelgasse 1, 4051 Basel

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

1 Beschreibende Statistik

In der beschreibenden Statistik geht es darum, grosse und unübersichtliche Datenmengen so aufzubereiten, dass wenige aussagekräftige Kenngrössen und Graphiken entstehen.

1.1 Grundbegriffe

In der Statistik nennt man Objekte, auf die sich eine statistische Untersuchung beziehen, *statistische Elemente* oder *Merkmalsträger*. Die Menge aller dieser Merkmalsträger heisst *Grundgesamtheit*. Wie der Name sagt, interessieren uns an den Merkmalsträgern gewisse Eigenschaften oder *Merkmale*. Die möglichen Werte, die ein Merkmal annehmen kann, heissen *Merkmalsausprägungen*.

Beispiele

<i>Grundgesamtheit</i>	<i>Merkmal</i>	<i>Merkmalsausprägungen</i>
Alle Studierenden der Vorlesung Mathematik II	Alter (in Jahren)	..., 19, 20, 21, ...
Bäume in der Schweiz	Baumart	Ahorn, Birke, Arve, ...
Arbeitslose in Basel-Stadt	Schulabschluss	Gymnasium, Sekundarschule, keiner, ...
Eingesammelte Bebbi-Säcke (35 l)	Gewicht (in kg)	..., 28, 35.5, 49.7, ...
Tage des Januars 2022	Durchschnittstemperatur (in °C)	..., -2, 4.5, 6, 11 ...

Bei Merkmalen unterscheidet man zwischen *qualitativen* und *quantitativen* Merkmalen.

Qualitative Merkmale

Dies sind Merkmale, die artmässig erfassbar sind und keine physikalische Masseinheit benötigen. Weiter wird hier unterschieden zwischen

- *nominalen Merkmalen*

Die Merkmalsausprägungen werden nur dem Namen nach unterschieden, ohne Wertung.

Beispiele: Baumart, Vorname, Studienfach, Nationalität

- *ordinalen Merkmalen*

Die Merkmalsausprägungen weisen eine natürliche Rangordnung auf.

Beispiele: Schulabschluss, Hausnummern, Lawinengefahrenskala

Quantitative Merkmale

Dies sind Merkmale, die durch Zahlen erfassbar sind und eine physikalische Masseinheit haben. Weiter wird hier unterschieden zwischen

- *diskreten Merkmalen*

Die Merkmalsausprägungen sind isolierte Zahlenwerte. Werte dazwischen können nicht angenommen werden.

Beispiele: Alter in Jahren, Anzahl Studierende pro Studienfach, Anzahl Einwohner

- *stetigen Merkmalen*

Diese Merkmale können (theoretisch) jeden Wert innerhalb eines Intervalls annehmen.

Beispiele: Gewicht, Durchschnittstemperatur, Grösse, Geschwindigkeit

Skalierung von Merkmalen

Man kann Merkmale auch hinsichtlich der Skala, auf der sie gemessen werden, unterscheiden. Von der Skala hängt ab, ob mit den Merkmalsausprägungen sinnvoll gerechnet werden kann.

- *Nominale Skala*

In einer nominalen Skala werden Zahlen als Namen ohne mathematische Bedeutung verwendet. Rechnen mit solchen Zahlen ist sinnlos.

Beispiele: Postleitzahlen, Codes

Zum Beispiel haben wir die Postleitzahlen

4051 Basel

8102 Oberengstringen (ZH)

Es ist $8102 = 2 \cdot 4051$, aber Oberengstringen ist nicht doppelt so gross wie Basel.

- *Ordinale Skala*

Die natürliche Ordnung der Zahlen ordnet die Objekte nach einem bestimmten Kriterium. Vergleiche sind sinnvoll, Differenzen und Verhältnisse jedoch nicht.

Beispiele: Prüfungsnoten, Hausnummern, Lawinengefahrenskala

Es gilt $|18 - 16| = |18 - 20| = 2$, doch die Distanz des Hauses mit der Nummer 18 zu den Häusern mit den Nummern 16 und 20 ist im Allgemeinen nicht gleich gross.

- *Intervallskala*

Der Nullpunkt ist willkürlich. Differenzen sind sinnvoll, Verhältnisse jedoch nicht.

Beispiele: Temperatur in °C und in °F, Höhe in m über Meer

Zum Beispiel hat die Aussage "Heute ist es doppelt so warm wie gestern" in °F eine andere Bedeutung als in °C.

- *Verhältnisskala*

Der Nullpunkt ist natürlich fixiert. Differenzen und Verhältnisse sind sinnvoll.

Beispiele: Geschwindigkeit, Gewicht, Masse, Volumen

Im Folgenden werden wir es meistens mit (quantitativen) Merkmalen auf einer Intervall- oder Verhältnisskala zu tun haben. Solche Merkmalsausprägungen entstehen durch Messungen.

Stichprobe

Eine *Stichprobe* ist eine zufällig ausgewählte endliche Teilmenge aus einer Grundgesamtheit. Hat diese Teilmenge n Elemente, so spricht man von einer Stichprobe *vom Umfang n* . Zum Beispiel werden 10 Studierende der Vorlesung Mathematik II zufällig ausgewählt. Dann sind diese 10 Studierenden eine Stichprobe vom Umfang 10 der Grundgesamtheit aller Studierenden der Vorlesung Mathematik II.

1.2 Häufigkeitsverteilung

Messdaten, das heisst Merkmalsausprägungen eines Merkmals, fallen zunächst ungeordnet in einer sogenannten *Urliste* an. Um einen Überblick über die Daten zu gewinnen, bestimmt man die Häufigkeitsverteilung des Merkmals (in der Stichprobe).

Das Merkmal X habe die k verschiedenen Merkmalsausprägungen $a_1, \dots, a_k$. Wir entnehmen eine Stichprobe vom Umfang n und notieren die Werte $x_1, \dots, x_n$ der Stichprobe (Urliste). Nun zählen wir, wie oft jede Merkmalsausprägung a_j in der Stichprobe auftritt. Diese Anzahl h_j nennt man *absolute Häufigkeit* von a_j , für $j = 1, \dots, k$:

$$h_j = \text{Anzahl der } x_i \text{ mit der Ausprägung } a_j$$

Die *relative Häufigkeit* f_j von a_j , für $j = 1, \dots, k$, ist gegeben durch

$$f_j = \frac{h_j}{n}.$$

Es gilt

$$0 \leq h_j \leq n \quad \text{und} \quad h_1 + \dots + h_k = n,$$

und Division durch n ergibt

$$0 \leq f_j \leq 1 \quad \text{und} \quad f_1 + \dots + f_k = 1.$$

Die Menge der Paare

$$\{ (a_j, h_j) \mid j = 1, \dots, k \} \quad \text{bzw.} \quad \{ (a_j, f_j) \mid j = 1, \dots, k \}$$

nennt man *Häufigkeitsverteilung* des Merkmals X in der Stichprobe. Sie kann mit Hilfe einer *Häufigkeitstabelle* bestimmt und graphisch durch ein *Stab-* oder *Balkendiagramm* dargestellt werden. Auf der waagrechten Achse werden die Merkmalsausprägungen $a_1, \dots, a_k$ abgetragen und darüber je ein Stab oder Balken, dessen Höhe der absoluten, bzw. relativen Häufigkeit entspricht.

Beispiel

Bei einer Befragung gaben 20 Personen Auskunft über die Anzahl Zimmer in ihrer Wohnung. Dies ergab die folgende Urliste:

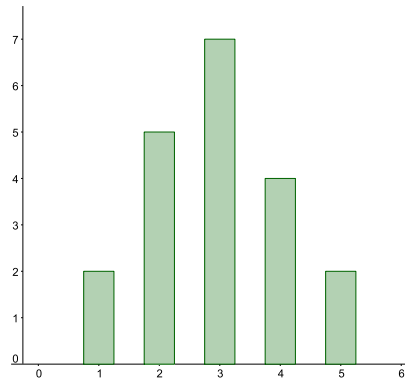
2, 4, 3, 4, 2, 3, 4, 5, 2, 1, 3, 2, 5, 3, 3, 4, 1, 2, 3, 3

Es ist also $n = 20$ und das Merkmal $X = (\text{Anzahl Zimmer})$ hat die Ausprägungen $a_1 = 1$, $a_2 = 2$, $a_3 = 3$, $a_4 = 4$, $a_5 = 5$.

Die Häufigkeitstabelle sieht so aus:

Anzahl Zimmer a_j	Strichliste	Häufigkeiten	
		absolut h_j	relativ f_j
1			
2			
3			
4			
5			
Summe			

Stabdiagramm mit absoluten Häufigkeiten:



Werden im Stabdiagramm die relativen anstatt die absoluten Häufigkeiten abgetragen, ändert sich nur die Beschriftung der senkrechten Achse. Allerdings ist dann der Umfang der Stichprobe nicht mehr ersichtlich.

Stetige Merkmale

Ist das (quantitative) Merkmal X stetig oder die Anzahl k der Merkmalsausprägungen von X viel grösser als der Stichprobenumfang n , dann ist das vorher beschriebene Vorgehen nicht sinnvoll, da die Häufigkeiten h_j sehr klein sind, bzw. viele h_j gleich Null sind. In diesem Fall fassen wir die Merkmalsausprägungen zu Klassen zusammen.

Seien wieder $x_1, \dots, x_n$ die Werte der Stichprobe und nehmen wir an, sie liegen im Intervall $[a, b)$. Dann unterteilen wir das Intervall $[a, b)$ in m (halboffene) Teilintervalle

$$[a_1, a_2), [a_2, a_3), [a_3, a_4), \dots, [a_m, a_{m+1})$$

mit $a = a_1 < a_2 < a_3 < \dots < a_m < a_{m+1} = b$. Das Intervall $[a_j, a_{j+1})$ nennt man *j-te Klasse*. Nun zählt man, wie viele der Stichprobenwerte $x_1, \dots, x_n$ in die einzelnen Klassen fallen.

Die *absolute Häufigkeit* h_j der *j-ten Klasse* ist gegeben durch

$$h_j = \text{Anzahl der } x_i \text{ mit } x_i \in [a_j, a_{j+1})$$

und die *relative Häufigkeit* f_j der *j-ten Klasse* ist

$$f_j = \frac{h_j}{n}.$$

Die Menge der Klassen mit ihren Häufigkeiten heisst *klassierte Häufigkeitsverteilung*.

Bei der Klassenbildung geht natürlich Information verloren. Die Verteilung der Werte innerhalb einer Klasse ist nicht mehr erkennbar. Viele Klassen bedeuten einen geringen Informationsverlust, aber wenige Klassen eine bessere Übersicht. Bei der Suche nach einem Kompromiss helfen die folgenden Faustregeln:

- $m \leq \sqrt{n}$ und $5 \leq m \leq 20$ für die Anzahl Klassen m und den Stichprobenumfang n
- Die Klassenbreiten (d.h. Intervalllängen) sollten alle gleich sein.

Graphisch stellt man eine klassierte Häufigkeitsverteilung mit Hilfe eines *Histogramms* dar. Die Intervallgrenzen $a_1, \dots, a_{m+1}$ werden auf der waagrechten Achse abgetragen und über jeder Klasse ein Rechteck gezeichnet, dessen Fläche proportional zur Häufigkeit der jeweiligen Klasse ist.

Beispiel

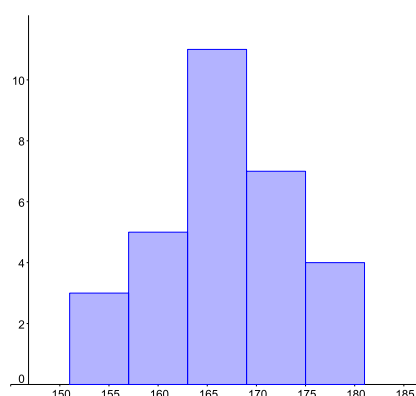
Von 30 (fiktiven) Studentinnen der Pharmazie wurden die Körperlängen (in cm) gemessen:

166, 168, 178, 177, 173, 163, 164, 167, 165, 162, 156, 163, 174, 165, 171, 169, 169,
159, 151, 163, 180, 170, 157, 170, 163, 160, 154, 178, 167, 161

Es ist also $n = 30$ und das Merkmal $X = \text{Körperlänge}$ nimmt in der Stichprobe Ausprägungen im Intervall $[151, 181)$ an. Wir wählen $m = 5$ Klassen. Dies ergibt die folgende Häufigkeitstabelle:

j	Klasse j in cm	Strichliste	Häufigkeiten	
			absolut h_j	relativ f_j
1	[151, 157)		3	0,1
2	[157, 163)		5	0,167
3	[163, 169)		11	0,367
4	[169, 175)		7	0,233
5	[175, 181)		4	0,133
Summe			30	1

Histogramm:



1.3 Mittelwerte

Gegeben seien n Zahlen $x_1, \dots, x_n$, die Merkmalsausprägungen eines quantitativen Merkmals X (der Grundgesamtheit oder einer Stichprobe davon) sind. Gesucht ist eine einzige Zahl, welche die "Mitte" der n Zahlen angibt, um die herum sich die gegebenen Zahlen häufen. In den meisten Fällen wird das arithmetische Mittel verwendet. In manchen Situationen ist jedoch die Angabe des sogenannten Medians besser geeignet.

Das arithmetische Mittel

Definition Das *arithmetische Mittel* $\bar{x}$ der Zahlen $x_1, \dots, x_n$ ist definiert durch

$$\bar{x} = \frac{1}{n}(x_1 + \dots + x_n) = \frac{1}{n} \sum_{i=1}^n x_i.$$

Oft nennt man das arithmetische Mittel auch *Durchschnitt* oder einfach *Mittelwert*.

Beispiel

i	1	2	3	4	5
x_i	9	4	3	6	3

Wir berechnen den Durchschnitt (d.h. das arithmetische Mittel):

Eigenschaften

1. Die Summe der Quadrate der Abstände vom arithmetischen Mittel $\bar{x}$ zu den einzelnen Messwerten $x_1, \dots, x_n$ ist minimal; das heisst, die Funktion

$$f(x) = \sum_{i=1}^n (x_i - x)^2$$

ist minimal für $x = \bar{x}$.

Im obigen Beispiel könnten wir also auch einfach das Minimum der Funktion

$$f(x) = \sum_{i=1}^5 (x_i - x)^2 = (9 - x)^2 + (4 - x)^2 + (3 - x)^2 + (6 - x)^2 + (3 - x)^2$$

bestimmen. Dies ist schnell gemacht. Durch Ausmultiplizieren erhalten wir

$$f(x) = 5x^2 - 50x + 151.$$

Das Minimum von f finden wir durch Nullsetzen der Ableitung:

Dass allgemein die Funktion $f(x) = \sum_{i=1}^n (x_i - x)^2$ ein Minimum in $x = \bar{x}$ hat, zeigt man ebenso durch Nullsetzen der Ableitung.

2. Das arithmetische Mittel hat weiter die Eigenschaft, dass die Summe aller Abweichungen “links” von $\bar{x}$ gleich der Summe aller Abweichungen “rechts” davon ist:

$$\sum_{x_i < \bar{x}} (\bar{x} - x_i) = \sum_{x_i > \bar{x}} (x_i - \bar{x})$$

Physikalisch interpretiert ist das gerade die Gleichgewichtsbedingung: Denkt man sich die x -Achse als langen masselosen Stab und darauf an den Positionen x_i jeweils eine konstante punktförmige Masse angebracht, so befindet sich der Stab genau dann im Gleichgewicht, wenn er im Punkt $\bar{x}$ gehalten wird.

Hier der Nachweis dieser Eigenschaft:

$$\sum_{x_i > \bar{x}} (x_i - \bar{x}) - \sum_{x_i < \bar{x}} (\bar{x} - x_i) = \sum_{i=1}^n (x_i - \bar{x}) = \sum_{i=1}^n x_i - \sum_{i=1}^n \bar{x} = \sum_{i=1}^n x_i - n\bar{x} = 0.$$

3. Das arithmetische Mittel ist empfindlich gegenüber *Ausreissern*. Eine Zahl in einer Datenreihe nennt man Ausreisser, wenn sie von den anderen Daten weit weg liegt (im nächsten Abschnitt wird dies noch präzisiert). In vielen Fällen entsteht ein Ausreisser aufgrund eines Schreib- oder Messfehlers.

Beispiel

In einem kleinen Dorf wohnen 20 Handwerker und ein Manager. Nehmen wir an, die Handwerker verdienen etwa 3000 CHF pro Monat und der Manager 40000 CHF. Ist der Durchschnitt ein guter Repräsentant für die Einkommen in diesem Dorf?

Nein, der Durchschnitt $\bar{x}$ wird durch das Einkommen des Managers dermassen in die Höhe gezogen, dass die Einkommen der Handwerker nicht erkennbar sind.

Für Situationen wie im Beispiel brauchen wir eine andere Zahl als den Durchschnitt, um die “Mitte” einer Datenreihe angeben zu können.

Der Median oder Zentralwert

Definition Der *Median* oder *Zentralwert* $\tilde{x}$ der Zahlen $x_1, \dots, x_n$ ist der mittlere Wert der nach der Grösse geordneten Zahlen $x_1, \dots, x_n$.

Dies bedeutet: Die Zahlen $x_1, \dots, x_n$ werden zuerst der Grösse nach geordnet. Ist die Anzahl n der Werte ungerade, so gibt es einen mittleren Wert $\tilde{x}$. Ist n gerade, so sind zwei Zahlen in der Mitte und $\tilde{x}$ kann zwischen diesen Zahlen gewählt werden. Üblich ist in diesem Fall, $\tilde{x}$ als arithmetisches Mittel der beiden Zahlen zu wählen, was auch wir hier tun werden. Dies ist jedoch nicht einheitlich festgelegt.

Beispiele

1. Im obigen Beispiel schreiben wir die Einkommen der Grösse nach geordnet hin. Das Einkommen des Managers ist (mit Abstand) der grösste Wert. Der mittlere Wert ist eines der 20 Einkommen der Handwerker. Also ist der Median in diesem Beispiel ein sinnvoller Repräsentant der Einkommen.

2.

i	1	2	3	4	5
x_i	9	4	3	6	3

Wir berechnen den Median:

Nehmen wir in diesem Beispiel eine weitere Zahl $x_6 = 10$ hinzu:

i	1	2	3	4	5	6
x_i	9	4	3	6	3	10

Median:

Eigenschaften

1. Die Summe der Abstände vom Median zu den einzelnen Zahlen $x_1, \dots, x_n$ ist minimal; das heisst, die Funktion

$$f(x) = \sum_{i=1}^n |x_i - x|$$

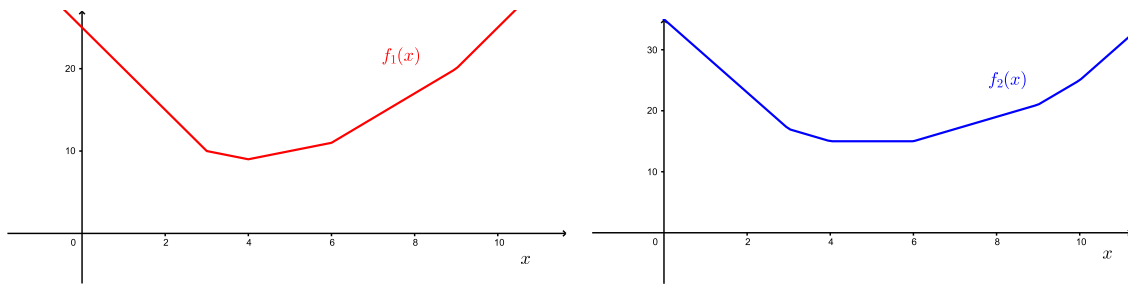
ist minimal für $x = \tilde{x}$.

Im 2. Beispiel oben könnten wir also auch einfach die Minima der Funktionen

$$f_1(x) = \sum_{i=1}^5 |x_i - x| = |9 - x| + |4 - x| + |3 - x| + |6 - x| + |3 - x|$$

$$f_2(x) = \sum_{i=1}^6 |x_i - x| = |9 - x| + |4 - x| + |3 - x| + |6 - x| + |3 - x| + |10 - x|$$

bestimmen. Nullsetzen der Ableitungen funktioniert nun aber nicht, da weder f_1 noch f_2 differenzierbar ist. Schauen wir uns stattdessen die Graphen von f_1 und f_2 an:



Die Funktion $f_1(x)$ hat wie erwartet ein Minimum in $x = \tilde{x} = 4$. Die Funktion $f_2(x)$ hat ein Minimum in $x = \tilde{x} = 5$ (aber auch jedes andere x zwischen 4 und 6 ist eine Minimalstelle).

2. Eine wichtige Eigenschaft des Medians ist, dass er unempfindlich gegenüber Ausreissern ist.

3. Der Median wird auch mit $\tilde{x} = \tilde{x}_{0,5}$ bezeichnet. Dies, weil höchstens die Hälfte aller Zahlen kleiner als $\tilde{x}$ und höchstens die Hälfte aller Zahlen grösser als $\tilde{x}$ ist.

Der Median ist also die Schnittstelle, wenn man die der Grösse nach geordneten Zahlen in zwei gleich grosse Haufen teilt.

1.4 Quantile und Boxplot

Es ist nützlich, die dritte Eigenschaft des Medians wie folgt zu verallgemeinern. Die der Grösse nach geordneten Zahlen $x_1, \dots, x_n$ werden in zwei Haufen geteilt, doch soll der erste Haufen (mit den kleineren Zahlen) zum Beispiel nur $\frac{1}{10} = 0,1$ aller Zahlen umfassen. An der Schnittstelle ist dann das sogenannte Quantil $\tilde{x}_{0,1}$.

Definition Sei α eine Zahl mit $0 \leq \alpha \leq 1$. Dann ist das *Quantil* $\tilde{x}_\alpha$ durch die folgende Bedingung definiert: Der Anteil der $x_i < \tilde{x}_\alpha$ ist $\leq \alpha$, der Anteil der $x_i > \tilde{x}_\alpha$ ist $\leq 1 - \alpha$.

Speziell nennt man das Quantil $\tilde{x}_{0,25}$ das *erste Quartil*, das Quantil $\tilde{x}_{0,75}$ das *dritte Quartil* und entsprechend ist der Median $\tilde{x}_{0,5}$ auch das zweite Quartil.

Wird α in Zehnteln angegeben, spricht man von *Dezilen*, bei Hundertsteln von *Perzentilen*. Der Median ist also auch das fünfte Dezil oder das fünfzigste Perzentil.

Beispiele

1. Gesucht ist das erste Quartil $\tilde{x}_{0,25}$ der Zahlen

i	1	2	3	4	5	6	7	8
x_i	2	3	7	13	13	18	21	24

Ein Viertel von 8 Messwerten sind 2 Messwerte, also liegt das Quartil $\tilde{x}_{0,25}$ zwischen der zweiten und der dritten Zahl, das heisst zwischen 3 und 7. Wie beim Median ist es üblich, für $\tilde{x}_{0,25}$ den Mittelwert der beiden Zahlen zu nehmen:

2. Gesucht ist das erste Quartil $\tilde{x}_{0,25}$ der Zahlen

i	1	2	3	4	5	6
x_i	2	3	7	13	13	18

Nun ist Vierteln der Messwerte nicht mehr möglich. Wir müssen also die Definition für das Quantil (für $\alpha = 0,25$) anwenden: Der Anteil der $x_i < \tilde{x}_{0,25}$ ist $\leq 0,25$, der Anteil der $x_i > \tilde{x}_{0,25}$ ist $\leq 1 - 0,25 = 0,75$.

Ein Anteil von 0,25 von 6 Messwerten ist gleich $0,25 \cdot 6 = 1,5$ Messwerte. Die Aussage “der Anteil der $x_i < \tilde{x}_{0,25}$ ist $\leq 0,25$ ” bedeutet also, dass es höchstens 1,5 Messwerte x_i mit $x_i < \tilde{x}_{0,25}$ gibt; das heisst, es gibt höchstens einen solchen Messwert x_i .

Die Aussage “der Anteil der $x_i > \tilde{x}_{0,25}$ ist $\leq 1 - 0,25 = 0,75$ ” bedeutet dementsprechend, dass es höchstens $0,75 \cdot 6 = 4,5$ Messwerte mit $x_i > \tilde{x}_{0,25}$ gibt; das heisst es gibt höchstens 4 solche Messwerte.

Wir sehen nun, dass nur $\tilde{x}_{0,25} = x_2 = 3$ diese Bedingungen erfüllt.

Satz 1.1 Gegeben seien der Grösse nach geordnete Zahlen $x_1, \dots, x_n$ und $0 \leq \alpha \leq 1$.

- Ist $n\alpha$ eine ganze Zahl (wie im 1. Beispiel), dann liegt $\tilde{x}_\alpha$ zwischen zwei der gegebenen Zahlen. Es gilt

$$\tilde{x}_\alpha = \frac{1}{2}(x_{n\alpha} + x_{n\alpha+1}).$$

- Ist $n\alpha$ keine ganze Zahl (wie im 2. Beispiel), dann ist $\tilde{x}_\alpha$ eine der gegebenen Zahlen. Es gilt

$$\tilde{x}_\alpha = x_{[n\alpha]}$$

wobei $[n\alpha]$ bedeutet, dass $n\alpha$ auf eine ganze Zahl aufgerundet wird, also zum Beispiel ist $[3,271] = 4$.

Beispiele

1. Wir untersuchen die Messwerte der Lymphozytenanzahl X pro Blutvolumeneinheit von 84 Ratten:

968, 1090, 1489, 1208, 828, 1030, 1727, 2019, 944, 1296, 1734, 1089, 686, 949, 1031, 1699, 692, 719, 750, 924, 715, 1383, 718, 894, 921, 1249, 1334, 806, 1304, 1537, 1878, 605, 778, 1510, 723, 872, 1336, 1855, 928, 1447, 1505, 787, 1539, 934, 1650, 727, 899, 930, 1629, 878, 1140, 1952, 2211, 1165, 1368, 676, 813, 849, 1081, 1342, 1425, 1597, 727, 1859, 1197, 761, 1019, 1978, 647, 795, 1050, 1573, 2188, 650, 1523, 1461, 1691, 2013, 1030, 850, 945, 736, 915, 1521.

Gesucht sind der Median, die beiden Quartile sowie das Dezil $\tilde{x}_{0,1}$. Also müssen wir die Messwerte zuerst der Grösse nach ordnen. Das erledigt zum Beispiel Excel für uns.

i	x_i	i	x_i	i	x_i	i	x_i
1	605	22	849	43	1081	64	1521
2	647	23	850	44	1089	65	1523
3	650	24	872	45	1090	66	1537
4	676	25	878	46	1140	67	1539
5	686	26	894	47	1165	68	1573
6	692	27	899	48	1197	69	1597
7	715	28	915	49	1208	70	1629
8	718	29	921	50	1249	71	1650
9	719	30	924	51	1296	72	1691
10	723	31	928	52	1304	73	1699
11	727	32	930	53	1334	74	1727
12	727	33	934	54	1336	75	1734
13	736	34	944	55	1342	76	1855
14	750	35	945	56	1368	77	1859
15	761	36	949	57	1383	78	1878
16	778	37	968	58	1425	79	1952
17	787	38	1019	59	1447	80	1978
18	795	39	1030	60	1461	81	2013
19	806	40	1030	61	1489	82	2019
20	813	41	1031	62	1505	83	2188
21	828	42	1050	63	1510	84	2211

Die Quartile können wir nun ablesen.

erstes Quartil:

Median:

drittes Quartil:

Nun berechnen wir noch das Dezil $\tilde{x}_{0,1}$ mit Hilfe von Satz 1.1.

2. Wir betrachten die erzielten Punkte an der Prüfung Mathematik I vom 28.01.22. Es gab 209 Prüfungsteilnehmer*innen, wir haben also 209 ungeordnete Zahlen $x_1, \dots, x_{209}$, wobei jede Zahl x_i die Anzahl der erzielten Punkte der Person i angibt. Um die Quartile zu berechnen, ordnen wir zuerst diese 209 Zahlen der Grösse nach.

Für den Median $\tilde{x} = \tilde{x}_{0,5}$ rechnen wir $209 \cdot 0,5 = 104,5$. Der zweite Punkt von Satz 1.1 sagt nun, dass der Median gleich der 105. geordneten Zahl ist (da $\lceil 104,5 \rceil = 105$). Diese Zahl ist gleich 31,5. Es gilt also $\tilde{x} = 31,5$ (Punkte).

Für das erste Quartil rechnen wir $209 \cdot 0,25 = 52,25$. Wieder der zweite Punkt von Satz 1.1 sagt, dass $\tilde{x}_{0,25}$ gleich der 53. geordneten Zahl ist (da $\lceil 52,25 \rceil = 53$). Diese Zahl ist gleich 21, das heisst, $\tilde{x}_{0,25} = 21$ (Punkte).

Für das dritte Quartil rechnen wir $209 \cdot 0,75 = 156,75$. Analog zum ersten Quartil ist $\tilde{x}_{0,75}$ gleich der 157. geordneten Zahl. Diese Zahl ist gleich 38,5. Wir erhalten also $\tilde{x}_{0,75} = 38,5$ (Punkte).

Wir sehen in diesem Beispiel, dass das erste und das dritte Quartil etwas über die Streuung der Daten aussagt. Nämlich die Hälfte der Zahlen (die "mittlere Hälfte") liegt zwischen $\tilde{x}_{0,25} = 21$ und $\tilde{x}_{0,75} = 38,5$. Und da der Median $\tilde{x} = 31,5$ deutlich näher bei $\tilde{x}_{0,75}$ als bei $\tilde{x}_{0,25}$ liegt, ist die Streuung "gegen unten" deutlich grösser. Für ein aussagekräftiges Gesamtbild interessiert allenfalls noch die kleinste Zahl $x_{\min} = 0,5$ und die grösste Zahl $x_{\max} = 50$.

Für eine bessere Übersicht werden die Quartile durch einen Boxplot graphisch dargestellt.

Boxplot

Der Boxplot eines Datensatzes stellt die Lage des Medians, des ersten und dritten Quartils, der Extremwerte und der Ausreisser graphisch dar.

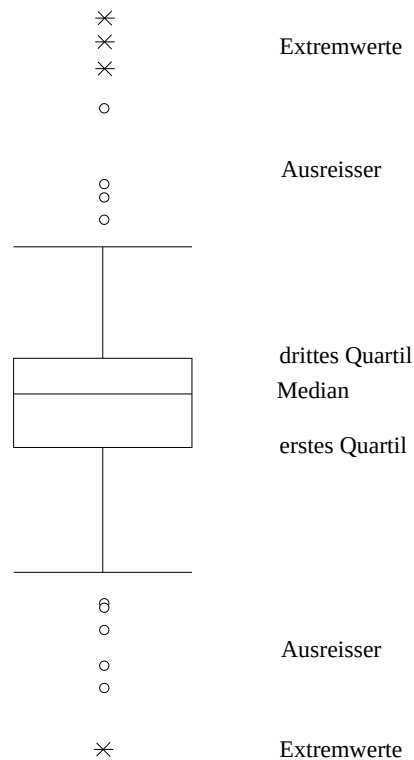
• innerhalb der Box

untere Boxgrenze	$\tilde{x}_{0,25}$
obere Boxgrenze	$\tilde{x}_{0,75}$
Linie in der Box	$\tilde{x}_{0,5}$

Die Höhe der Box wird als *Interquartilsabstand* bezeichnet. Dieser Teil umfasst also die Hälfte aller Daten.

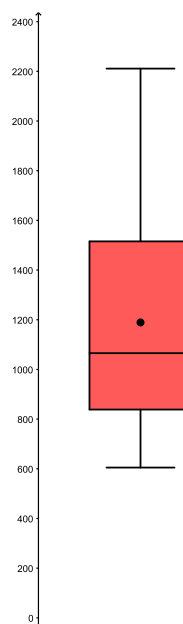
• ausserhalb der Box

- *Extremwerte*: mehr als 3 Boxlängen vom unteren bzw. oberen Boxrand entfernt, wiedergegeben durch „*“
- *Ausreisser*: zwischen $1\frac{1}{2}$ und 3 Boxlängen vom oberen bzw. unteren Boxrand entfernt, wiedergegeben durch „o“
- Der kleinste und der grösste Wert, der jeweils nicht als Ausreisser eingestuft wird, ist durch eine horizontale Strecke darzustellen.

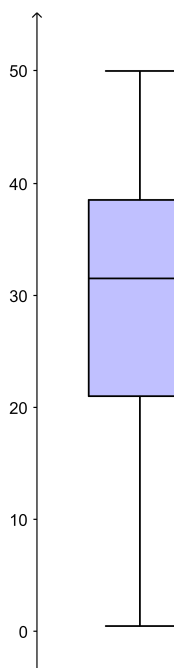


Beispiele

1. Im ersten Beispiel der Seite 10 haben wir die folgenden Quartile berechnet: $\tilde{x}_{0,25} = 838,5$, $\tilde{x}_{0,5} = 1065,5$, $\tilde{x}_{0,75} = 1515,5$. Es gibt weder Ausreisser noch Extremwerte. Der kleinste Wert ist 605 und der grösste Wert 2211. Der Boxplot sieht wie folgt aus, wobei hier der schwarze Punkt in der Box die Lage des Mittelwerts $\bar{x} = 1189,18$ beschreibt.



2. Im zweiten Beispiel auf Seite 11 haben wir die folgenden Quartile erhalten: $\tilde{x}_{0,25} = 21$, $\tilde{x}_{0,5} = 31,5$, $\tilde{x}_{0,75} = 38,5$. Auch hier gibt es weder Ausreisser noch Extremwerte. Die grösste Zahl ist 50 und 0,5 ist die kleinste Zahl.



1.5 Empirische Varianz und Standardabweichung

Mittelwerte und Quantile alleine genügen nicht für die Beschreibung eines Datensatzes.

Beispiel

Zwei Studenten der Geowissenschaften, nennen wir sie A und B, haben bei acht Examen die folgenden Noten erzielt. Student A: 4, 4, 4, 3, 5, 4, 4, 4. Student B: 2, 6, 2, 6, 2, 6, 2, 6. Beide Studenten haben einen Notendurchschnitt von einer 4 und auch der Median ist bei beiden 4 (bei B ist $\tilde{x} = \tilde{x}_{0,5}$ das arithmetische Mittel von 2 und 6, also 4). Dabei unterscheiden sich A und B völlig in der Konstanz ihrer Leistungen. Die Quartile geben einen Hinweis auf die grössere Streuung der Noten von B, doch sie sagen nichts aus über die einzelnen Abweichungen vom arithmetischen Mittel.

Zusätzlich zu den Mittelwerten und Quantilen benötigen wir deshalb Masszahlen, die etwas über die Abweichung der Einzeldaten vom arithmetischen Mittel aussagen: die Varianz und die Standardabweichung.

Definition Die (*empirische*) Varianz der Daten $x_1, \dots, x_n$ ist definiert durch

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2.$$

Die *Standardabweichung* ist die positive Quadratwurzel aus der Varianz,

$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}.$$

Die empirische Varianz ist also fast die mittlere quadratische Abweichung vom Mittelwert. Warum wir nicht den Faktor $\frac{1}{n}$, sondern den Faktor $\frac{1}{n-1}$ nehmen, werden wir erst später einsehen. Tatsächlich wird die Varianz oft auch mit dem Faktor $\frac{1}{n}$ definiert.

Beispiel

Für den Studenten A mit den Noten 4, 4, 4, 3, 5, 4, 4, 4 und dem Mittelwert $\bar{x} = 4$ gilt:

Für den Studenten B mit den Noten 2, 6, 2, 6, 2, 6, 2, 6 und dem Mittelwert $\bar{x} = 4$ gilt:

Die Formel für die empirische Varianz kann umgeformt werden:

Satz 1.2 *Es gilt*

$$s^2 = \frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - n\bar{x}^2 \right) .$$

Für konkrete Berechnungen ist diese Formel oft praktischer als die Definition.

Wann welche Masszahlen?

Um für eine Datenreihe die Lage auf der Zahlengeraden und die Streuung der Daten zu beschreiben, haben wir also das arithmetische Mittel und die Standardabweichung sowie den Median und die Quartile zur Verfügung.

Sind die Daten Merkmalsausprägungen eines Merkmals, das auf einer ordinalen Skala gemessen wird, dann können wir nur den Median und die Quartile gebrauchen (das arithmetische Mittel und die Standardabweichung sind sinnlos).

Wird das Merkmal hingegen auf einer Intervall- oder Verhältnisskala gemessen, haben wir die Wahl zwischen arithmetischem Mittel mit der Standardabweichung und dem Median mit den Quartilen. In den meisten Fällen wird das arithmetische Mittel mit der Standardabweichung verwendet. Weist die Datenreihe jedoch Ausreisser auf, ist im Allgemeinen der Median mit den Quartilen die bessere Wahl. Allerdings können diese Masszahlen auch missbraucht werden, um unerwünschte Ausreisser unter den Teppich zu kehren.

1.6 Prozentrechnen

Prozentrechnen ist lediglich Bruchrechnen, denn

$$1\% = \frac{1}{100} = 0,01 .$$

Beispiele

1. Wieviel ist 4 % von 200 ?

2. In der Prüfung Mathematik I vom HS21 haben 78 von den 209 Teilnehmer*innen die Note 5, 5.5 oder 6 erzielt. Wieviel Prozent sind das?

3. Eine Eisenbahngesellschaft hat die Billet-Preise seit 2007 zweimal erhöht, nämlich um 8,2 und um 11,8 Prozent. Das macht zusammen 20 Prozent. Stimmt diese Rechnung?

Absolut und relativ

Bei Statistiken können absolute Zahlenangaben andere Resultate liefern als Angaben in Prozenten.

Beispiele

1. Wir vergleichen die Altersverteilung in der Schweiz in den Jahren 1900 und 2000 (Quelle: Bundesamt für Statistik).

Schweiz	1900		2000	
	absolut	relativ	absolut	relativ
65 und mehr Jahre	193 266	6 %	1 109 416	23 %
20 – 64 Jahre	1 778 227	54 %	4 430 460	62 %
0 – 19 Jahre	1 343 950	40 %	1 664 124	15 %
Total	3 315 443	100 %	7 204 000	100 %

Betrachten wir den Anteil der Jugendlichen. In absoluten Zahlen wuchs der Anteil der Jugendlichen zwischen 1900 und 2000 (nämlich um 320 174 Jugendliche). Der relative Anteil nahm jedoch ab, und zwar um 25 Prozentpunkte (von 40 % auf 15 %).

2. Aus dem Erfundenland stammt die folgende Statistik:

Altersstufe	Landesbürger			Ausländer		
	total pro Altersstufe	davon kriminell absolut	relativ	total pro Altersstufe	davon kriminell absolut	relativ
0 – 19	4 Mio.	40 000	1 %	1 Mio.	2000	0,2 %
20 – 39	4 Mio.	400 000	10 %	6 Mio.	560 000	9,33 %
40 – 59	6 Mio.	60 000	1 %	1 Mio.	2000	0,2 %
60 – 79	4 Mio.	40 000	1 %	0,2 Mio.	1000	0,5 %
80 – 99	1 Mio.	1000	0,1 %	-	-	-

Die Partei A fasst dies so zusammen: Obwohl es viel mehr Landesbürger als Ausländer gibt (nämlich 19 Mio. Landesbürger und 8,2 Mio. Ausländer) gibt es mehr kriminelle Ausländer als kriminelle Landesbürger; nämlich 565 000 Ausländer sind kriminell im Gegensatz zu 541 000 kriminellen Landesbürgern.

Die Partei B kontert: In jeder Altersstufe stellen die Ausländer prozentual weniger Kriminelle als die Landesbürger.

3. Sie sind krank und der Arzt empfiehlt Ihnen, entweder Medikament A oder Medikament B einzunehmen.

Der Arzt sagt, dass Sie mit Medikament A schneller gesund werden als mit Medikament B, aber das Risiko einer gravierenden Nebenwirkung sei bei Medikament A um 100 Prozent grösser als bei Medikament B.

In absoluten Zahlen sieht es so aus: Bei Medikament A treten bei durchschnittlich 2 von 10 000 Patienten gravierende Nebenwirkungen auf, bei Medikament B lediglich bei 1 von 10 000 Patienten.

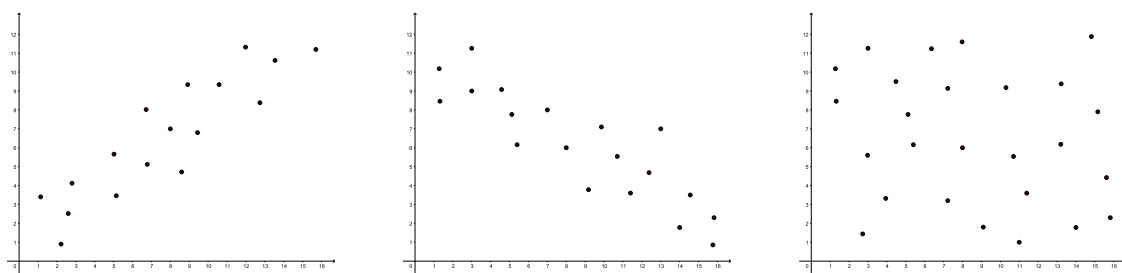
2 Korrelation und Regressionsgerade

Oft untersucht man nicht nur eine, sondern zwei Datenreihen und fragt sich, ob ein Zusammenhang zwischen den beiden Datenreihen besteht. Auskunft über einen linearen Zusammenhang gibt der sogenannte Korrelationskoeffizient.

2.1 Der Korrelationskoeffizient

Von einer Menge von Merkmalsträgern (Grundgesamtheit) betrachten wir zwei quantitative Merkmale X und Y , gemessen auf einer Intervall- oder Verhältnisskala. Hat ein Merkmalsträger i die Merkmalsausprägungen x_i von X und y_i von Y , dann notieren wir dies als Wertepaar (x_i, y_i) . Wir nehmen eine Stichprobe vom Umfang n und erhalten demnach n Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$. Zum Beispiel untersuchen wir die Merkmale $X = \text{Körpergrösse}$ und $Y = \text{Gewicht}$ von allen Studierenden der Universität Basel.

In diesem Beispiel vermutet man einen Zusammenhang zwischen den Merkmalen: Je grösser ein*e Studierende*r, desto grösser sein*ihr Gewicht. Um allgemein bei gegebenen Wertepaaren einen allfälligen Zusammenhang abschätzen zu können, zeichnet man die Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$ als Punkte im Koordinatensystem ein. Dies ergibt eine Punktwolke, die man *Streudiagramm* nennt. Hier drei Beispiele:



Im ersten Streudiagramm erkennt man einen Zusammenhang: Je grösser x_i , desto grösser y_i . Im zweiten Streudiagramm ist der Zusammenhang umgekehrt: Je grösser x_i , desto kleiner y_i . Und im dritten Streudiagramm ist kein Zusammenhang zwischen den x_i und den y_i erkennbar.

Wir sind hier auf der Suche nach einem *linearen* Zusammenhang, das heisst, wir fragen uns, ob die Wertepaare (ungefähr) auf einer Geraden liegen. Eine Antwort darauf liefert der Korrelationskoeffizient r_{xy} , der ein Mass sowohl für die Stärke des linearen Zusammenhangs als auch die Richtung im Falle eines Zusammenhangs ist. Im Korrelationskoeffizienten r_{xy} steckt die sogenannte Kovarianz c_{xy} , welche die Richtung eines allfälligen Zusammenhangs anzeigt.

Definition Die (*empirische*) Kovarianz der Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$ ist definiert durch

$$c_{xy} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}).$$

Mit denselben Rechenumformungen wie auf Seite 14 für die empirische Varianz finden wir die für Berechnungen praktischere Formel

$$c_{xy} = \frac{1}{n-1} \left(\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y} \right).$$

Ist $c_{xy} > 0$ (bzw. $c_{xy} < 0$), dann liegen die Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$, im Falle eines linearen Zusammenhangs, auf einer Geraden mit positiver (bzw. negativer) Steigung. Die Kovarianz kann jedoch beliebig grosse und beliebig kleine Werte annehmen und sie hängt von den Einheiten ab, mit denen die Merkmalsausprägungen x_i und y_i gemessen werden. Um eine Masszahl für die Stärke eines linearen Zusammenhangs zu erhalten, wird die Kovarianz deshalb durch die Standardabweichungen

$$s_x = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2} \quad \text{und} \quad s_y = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})^2}$$

der Zahlen $x_1, \dots, x_n$, bzw. $y_1, \dots, y_n$, dividiert.

Definition Gegeben seien die n Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$, wobei nicht alle x_i gleich sind und nicht alle y_i gleich sind. Der (*empirische*) *Korrelationskoeffizient* ist definiert durch

$$r_{xy} = \frac{c_{xy}}{s_x s_y} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}} = \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sqrt{\sum_{i=1}^n x_i^2 - n \bar{x}^2} \sqrt{\sum_{i=1}^n y_i^2 - n \bar{y}^2}}.$$

Der Korrelationskoeffizient r_{xy} wurde vom britischen Mathematiker KARL PEARSON (1857 – 1936) eingeführt. Die Interpretation von r_{xy} zeigt der folgende Satz.

Satz 2.1 *Der Korrelationskoeffizient nimmt nur Werte zwischen -1 und $+1$ an. Insbesondere gilt:*

$$\begin{aligned} r_{xy} = +1 &\iff y_i = ax_i + b \quad \text{mit } a > 0 \\ r_{xy} = -1 &\iff y_i = ax_i + b \quad \text{mit } a < 0. \end{aligned}$$

Die Wertepaare (x_i, y_i) liegen also exakt auf einer Geraden genau dann, wenn $r_{xy} = \pm 1$.

Woher kommen diese Eigenschaften von r_{xy} und wie sind die Werte von r_{xy} zwischen -1 und 1 zu interpretieren? Zur Beantwortung dieser Fragen definieren wir die beiden Vektoren in $\mathbb{R}^n$

$$\vec{x} = \begin{pmatrix} x_1 - \bar{x} \\ \vdots \\ x_n - \bar{x} \end{pmatrix} \quad \text{und} \quad \vec{y} = \begin{pmatrix} y_1 - \bar{y} \\ \vdots \\ y_n - \bar{y} \end{pmatrix}.$$

Dann gilt

$$r_{xy} = \frac{\vec{x} \cdot \vec{y}}{\|\vec{x}\| \|\vec{y}\|}$$

und die sogenannte Ungleichung von Cauchy-Schwarz sagt aus, dass die rechte Seite eine reelle Zahl zwischen -1 und 1 ist. Also gilt $-1 \leq r_{xy} \leq 1$.

In $\mathbb{R}^2$ und $\mathbb{R}^3$ gilt

$$\frac{\vec{x} \cdot \vec{y}}{\|\vec{x}\| \|\vec{y}\|} = \cos \varphi$$

für den Winkel φ zwischen den Vektoren $\vec{x}$ und $\vec{y}$. In $\mathbb{R}^n$ für $n > 3$ definiert man den Winkel φ zwischen $\vec{x}$ und $\vec{y}$ durch diese Gleichung. Es gilt also allgemein

$$r_{xy} = \cos \varphi$$

für den Zwischenwinkel φ der Vektoren $\vec{x}$ und $\vec{y}$.

Nehmen wir nun an, dass $r_{xy} \approx 1$ oder $r_{xy} \approx -1$. Dies bedeutet, dass der Zwischenwinkel φ von $\vec{x}$ und $\vec{y}$ nahe bei 0° , bzw. 180° ist. Die beiden Vektoren $\vec{x}$ und $\vec{y}$ sind also (beinahe) parallel, das heisst, $\vec{y} \approx a\vec{x}$ für eine reelle Zahl $a > 0$, bzw. $a < 0$:

Für die Komponenten gilt in diesem Fall

Wir können demnach folgern:

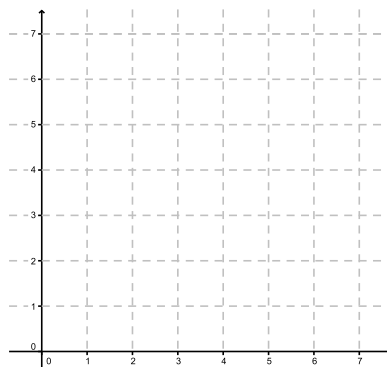
- Ist r_{xy} nahe bei 1, so gilt $y_i \approx ax_i + b$ für ein $a > 0$, das heisst, es besteht (beinahe) ein linearer Zusammenhang zwischen den Wertepaaren. Man spricht in diesem Fall von einer *starken positiven* Korrelation.
- Ist r_{xy} nahe bei -1 , so gilt $y_i \approx ax_i + b$ für ein $a < 0$, das heisst, es besteht (beinahe) ein linearer Zusammenhang zwischen den Wertepaaren. Man spricht in diesem Fall von einer *starken negativen* Korrelation.
- Ist r_{xy} nahe bei 0, so bedeutet dies, dass φ nahe bei 90° ist. Die beiden Vektoren $\vec{x}$ und $\vec{y}$ sind also fast orthogonal. Die Wertepaare korrelieren in diesem Fall nicht.

Beispiele

1. Gegeben sind die folgenden Wertepaare:

x_i	5	3	4	6	2
y_i	1	4	2	1	7

Streudiagramm:



Berechnungen:

i	x_i	y_i	$x_i y_i$	x_i^2	y_i^2
1	5	1			
2	3	4			
3	4	2			
4	6	1			
5	2	7			
Summe					

Mittelwerte:

Empirische Kovarianz (mit Hilfe der Formel nach der Definition):

Empirische Varianzen (mit Hilfe von Satz 1.2):

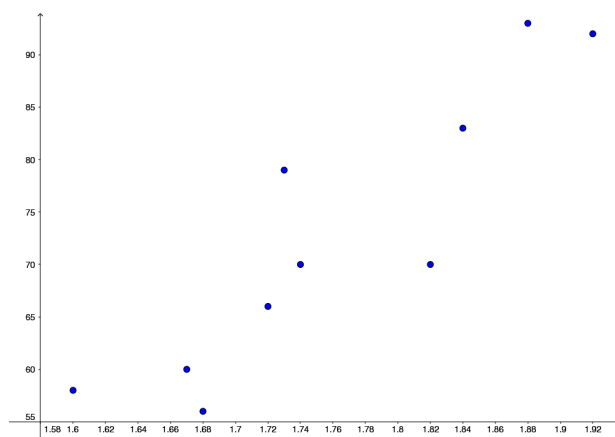
Korrelationskoeffizient:

Wir haben also eine starke negative Korrelation.

2. Gibt es einen linearen Zusammenhang zwischen der Körpergrösse und dem Gewicht eines Menschen? Gemessen wurden die Körpergrösse x_i (in cm) und das Gewicht y_i (in kg) von 10 Personen (von Schweizer Medaillengewinner*innen der Olympischen Winterspiele im Februar 2022):

x_i	173	184	172	160	168	174	182	167	192	188
y_i	79	83	66	58	56	70	70	60	92	93

Streudiagramm:



Wir finden (z.B. mit Excel, GeoGebra oder R)

$$r_{xy} = 0,901.$$

Wir haben eine starke positive Korrelation.

Bemerkungen zur Interpretation von r_{xy}

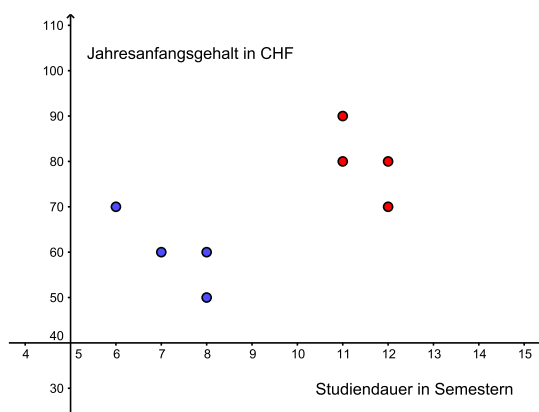
- Ist $r_{xy} \approx 0$, dann sagt dies nur, dass die beiden Datensätze keinen *linearen* Zusammenhang haben. Eventuell hängen sie jedoch quadratisch, exponentiell oder durch eine trigonometrische Funktion voneinander ab (vgl. Abschnitt 2.4).
- Falls r_{xy} nahe bei 1 oder -1 liegt, folgt lediglich, dass die Datensätze stark korrelieren. Man darf jedoch *nicht* daraus schliessen, dass zwischen den Datensätzen ein *kausaler* Zusammenhang besteht (d.h. dass der eine Datensatz Ursache für den anderen Datensatz ist). Es könnte so sein, es könnte aber auch eine gemeinsame Ursache im Hintergrund geben oder die Korrelation zufällig sein. Weiter muss ein Datensatz allenfalls in Teildatensätze unterteilt werden, um nicht eine der Erwartungen entgegengesetzte Korrelation zu erhalten (dieses Phänomen ist bekannt als *Simpson-Paradoxon*).

Beispiel

Wir betrachten die Jahresanfangsgehälter y_i (in 1000 CHF) von acht Universitätsabgänger*innen in Abhängigkeit von deren Studiendauer x_i (in Anzahl Semestern):

x_i	6	7	8	8	11	12	12	11
y_i	70	60	50	60	80	70	80	90

Der Korrelationskoeffizient $r_{xy} = 0,640$ weist auf eine positive Korrelation hin, also je länger die Studiendauer, desto höher das Anfangsgehalt. Doch das ist für Studierende zu schön, um wahr zu sein. Tatsächlich haben die ersten vier Studienabgänger*innen das gleiche Fach studiert und die restlichen vier ein anderes gemeinsames Fach (das mehr Zeit in Anspruch nimmt als das erste Fach). Im folgenden Streudiagramm sind die ersten vier Wertepaare blau und die restlichen vier rot eingezeichnet.



Betrachtet man die Fächer separat, so findet man für das erste Fach den Korrelationskoeffizienten $r_{xy} = -0,853$ und für das zweite Fach $r_{xy} = -0,707$. Studiendauer und Anfangsgehalt sind also doch negativ korreliert!

2.2 Rangkorrelation

Der Korrelationskoeffizient r_{xy} ist nicht sinnvoll, wenn eines der beiden Merkmale X und Y nicht auf einer Intervall- oder Verhältnisskala gemessen wird. Werden beide Merkmale zumindest auf einer Ordinalskala gemessen, dann kann der sogenannte Rangkorrelationskoeffizient gebildet werden.

Gegeben seien also die Merkmalsausprägungen $x_1, \dots, x_n$ und $y_1, \dots, y_n$ von zwei ordinalskalierten Merkmalen X , bzw. Y . Das heisst, den Daten können Ränge zugeordnet werden. Haben zwei oder mehr Daten denselben Rang (man nennt dies eine *Bindung*), so wird als Rang dieser Daten das arithmetische Mittel der zu vergebenden Ränge gewählt. Anschliessend bildet man von diesen Rängen r_{x_i} und r_{y_i} die Differenzen $d_i = r_{x_i} - r_{y_i}$. Das heisst, jedem Wertepaar (x_i, y_i) ordnet man die Rangdifferenz d_i zu.

Definition Gegeben seien die n Wertepaare $(x_1, y_1), \dots, (x_n, y_n)$ mit den Rangdifferenzen $d_1, \dots, d_n$. Der *Rangkorrelationskoeffizient* ist definiert durch

$$r_S = 1 - \frac{6}{n(n^2 - 1)} \sum_{i=1}^n d_i^2.$$

Der Rangkorrelationskoeffizient geht auf den britischen Psychologen CHARLES SPEARMAN (1863 - 1945) zurück.

Der Rangkorrelationskoeffizient r_S nimmt Werte zwischen -1 und 1 an und er wird analog zu r_{xy} interpretiert. Stimmen die Rangreihenfolgen für die beiden Datensätze überein, dann sind alle Rangdifferenzen d_i Null und $r_S = 1$. Bei genau umgekehrten Rangreihenfolgen für die beiden Datensätze führt der Faktor $\frac{6}{n(n^2-1)}$ zu $r_S = -1$.

Beispiel

In einem erdbebengefährdeten Gebiet fanden im vergangenen Jahr 7 Erdbeben statt. In der Tabelle sind die Stärke (gemäss Richterskala) sowie die Schadensumme (in Mio. CHF) von jedem Erdbeben aufgelistet.

Stärke	Schaden	Ränge		Rangdifferenz	
x_i	y_i	r_{x_i}	r_{y_i}	$d_i = r_{x_i} - r_{y_i}$	d_i^2
3,8	42				
2,6	33				
2,4	20				
3,7	40				
5,4	49				
6,2	45				
3,8	33				
				Summe	

Wir erhalten

Wir haben also eine starke positive Rangkorrelation.

2.3 Die Regressionsgerade

Wie im ersten Abschnitt dieses Kapitels betrachten wir von einer Grundgesamtheit zwei quantitative Merkmale X und Y . Anders als zuvor gehen wir jedoch davon aus, dass Y von X abhängt und wir fragen uns, *wie* Y von X abhängt. Wir nehmen eine Stichprobe von Wertepaaren $(x_1, y_1), \dots, (x_n, y_n)$ und suchen also eine Funktion f , so dass $y_i \approx f(x_i)$.

Beispiel

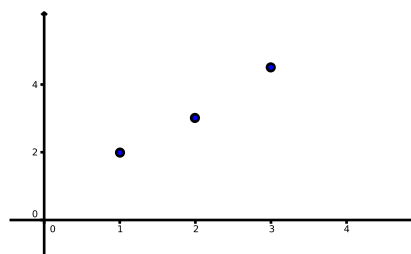
Der Umsatz einer Apotheke gibt einen wichtigen Hinweis auf ihre Wirtschaftlichkeit. Kann dieser Umsatz beispielsweise durch die Anzahl Kunden pro Tag abgeschätzt werden?

Bei drei Apotheken, für welche der Jahresumsatz bekannt ist, werden die Anzahl Kunden pro Tag gezählt. Man erhält die folgenden drei Messwertpaare (x_i, y_i) , $i = 1, 2, 3$,

i	1	2	3
x_i	1	2	3
y_i	2	3	4,5

wobei $x_i \cdot 100$ die Anzahl Kunden pro Tag in der Apotheke i sind und y_i der Jahresumsatz der Apotheke i in Millionen CHF ist. Wenn es eine Funktion f gibt, so dass $y_i \approx f(x_i)$, für $i = 1, 2, 3$, dann könnte für jede weitere Apotheke die Anzahl Kunden x gezählt werden und mit Hilfe der Funktionsgleichung $y = f(x)$ der Jahresumsatz y der Apotheke geschätzt werden.

Um eine passende Funktion f zu finden, zeichnen wir das Streudiagramm der Messwertpaare:



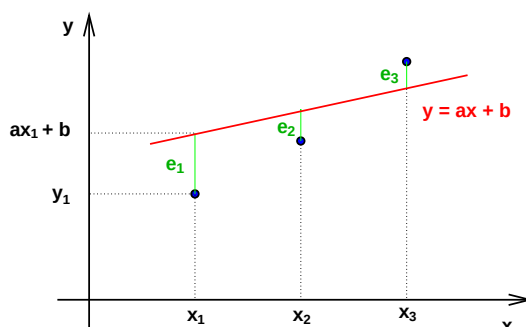
Die drei Punkte liegen fast auf einer Geraden. Es könnte also sein, dass ein linearer Zusammenhang zwischen den Messwerten x_1, x_2, x_3 und y_1, y_2, y_3 besteht, der jedoch durch verschiedene Einflüsse verfälscht wurde.

Wir machen deshalb den Ansatz

$$y = f(x) = ax + b$$

und versuchen, a und b so zu bestimmen, dass der Graph von f (eine Gerade) die drei Messwertpaare am besten approximiert. Setzen wir im Ansatz für x die Messwerte x_1, x_2, x_3 ein, dann sollen die Abweichungen $f(x_1)$ von y_1 , $f(x_2)$ von y_2 , $f(x_3)$ von y_3 möglichst klein sein. Im Beispiel sind dies die Abweichungen

$$\begin{aligned} e_1 &= y_1 - f(x_1) = 2 - (a + b) = 2 - a - b \\ e_2 &= y_2 - f(x_2) = 3 - (2a + b) = 3 - 2a - b \\ e_3 &= y_3 - f(x_3) = 4,5 - (3a + b) = 4,5 - 3a - b \end{aligned}$$



Wie beim arithmetischen Mittel soll die Summe der Quadrate der Abweichungen minimal sein, das heisst, wir suchen das Minimum der Funktion

$$\sum_{i=1}^3 e_i^2 = \sum_{i=1}^3 (y_i - f(x_i))^2 = \sum_{i=1}^3 (y_i - (ax_i + b))^2 = F(a, b).$$

Dies ist eine Funktion in zwei Variablen, nämlich in den Variablen a und b :

$$\begin{aligned} F(a, b) &= (2 - a - b)^2 + (3 - 2a - b)^2 + (4,5 - 3a - b)^2 \\ &= 14a^2 + 3b^2 + 12ab - 43a - 19b + 33,25 \end{aligned}$$

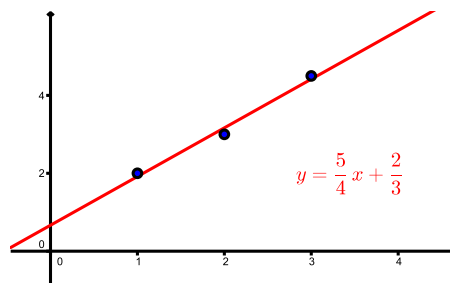
Wir werden im dritten Teil dieses Semesters lernen, dass eine notwendige Bedingung für ein Minimum das Verschwinden der Ableitungen von $F(a, b)$ nach den Variablen a und b ist:

$$(\text{Ableitung von } F(a, b) \text{ nach } a) = \frac{\partial}{\partial a} F(a, b) = 0$$

$$(\text{Ableitung von } F(a, b) \text{ nach } b) = \frac{\partial}{\partial b} F(a, b) = 0$$

Für unser Beispiel ergibt sich

Dies ist ein lineares Gleichungssystem in a und b mit der eindeutigen Lösung $a = \frac{5}{4}$ und $b = \frac{2}{3}$. Die gesuchte Gerade ist also $y = \frac{5}{4}x + \frac{2}{3}$. Die Graphik zeigt, dass $(a, b) = (\frac{5}{4}, \frac{2}{3})$ tatsächlich ein Minimum und nicht ein Maximum von $F(a, b)$ ist.



Zählen wir also in einer weiteren Apotheke beispielsweise 270 Kunden pro Tag, dann können wir den Jahresumsatz dieser Apotheke auf $y = f(2,7) \approx 4,04$ Millionen CHF schätzen.

Wir könnten das vorherige Problem auch mit einer anderen Methode lösen. Wir tun so, wie wenn die drei Messwertpaare auf einer Geraden $y = mx + q$ liegen würden. Wir setzen die drei Messwertpaare ein und erhalten also

$$\begin{aligned} 2 &= m + q \\ 3 &= 2m + q \\ 4,5 &= 3m + q. \end{aligned}$$

Dies ist nun ein lineares Gleichungssystem in m und q . Da die drei Messwertpaare nicht auf einer Geraden liegen, hat dieses Gleichungssystem natürlich keine Lösung. Wir können aber eine Näherungslösung bestimmen, und zwar nach der Methode von Abschnitt 9.5 vom letzten Semester. Das lineare System kann man schreiben als

$$A \begin{pmatrix} m \\ q \end{pmatrix} = \vec{b} \quad \text{mit } A = \begin{pmatrix} 1 & 1 \\ 2 & 1 \\ 3 & 1 \end{pmatrix}, \quad \vec{b} = \begin{pmatrix} 2 \\ 3 \\ 4,5 \end{pmatrix}.$$

Satz 9.14 sagt, dass eine Näherungslösung $\begin{pmatrix} m \\ q \end{pmatrix}$ gegeben ist durch

$$\begin{pmatrix} m \\ q \end{pmatrix} = (A^T A)^{-1} (A^T \vec{b}) = \frac{1}{6} \begin{pmatrix} 3 & -6 \\ -6 & 14 \end{pmatrix} \begin{pmatrix} 21,5 \\ 9,5 \end{pmatrix} = \begin{pmatrix} \frac{5}{4} \\ \frac{2}{3} \end{pmatrix}.$$

Wir erhalten also dieselbe Gerade wie mit der vorherigen Methode!

Dies überrascht eigentlich nicht, denn wir haben in Abschnitt 9.5 ja eine Summe von Quadraten minimiert (die Länge des "Fehlervektors"), genau wie bei der Minimierung von $F(a, b)$. Wie im Abschnitt 9.5 nennt man das Minimieren von $F(a, b)$ *Methode der kleinsten Quadrate*. Sie geht auf den Mathematiker CARL FRIEDRICH GAUSS (1777 – 1855) zurück.

Allgemeine Methode

Allgemein sind nun n Messwertpaare (x_i, y_i) , für $i = 1, \dots, n$, gegeben. Wir vermuten einen linearen Zusammenhang

$$y = f(x) = ax + b \quad \text{mit } a \neq 0$$

und bestimmen a und b so, dass die Summe der Quadrate der Abweichungen $e_i = y_i - (ax_i + b)$ minimal ist, das heisst, wir suchen die Minimalstelle (a, b) (es gibt tatsächlich genau eine) der Funktion

$$\sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - f(x_i))^2 = \sum_{i=1}^n (y_i - ax_i - b)^2 = F(a, b).$$

Die Gerade $y = ax + b$ mit dieser Minimalitätseigenschaft heisst *Regressionsgerade*.

Wie im Beispiel müssen wir zur Bestimmung der Minimalstelle die Ableitungen von $F(a, b)$ nach a und nach b Null setzen. Da $F(a, b)$ quadratisch in a und b ist, sind diese Ableitungen linear in a und b . Wir erhalten (wie im Beispiel) das folgende lineare Gleichungssystem in a und b :

$$\begin{aligned} \sum_{i=1}^n x_i y_i &= a \sum_{i=1}^n x_i^2 + b n \bar{x} \\ \bar{y} &= a \bar{x} + b \end{aligned}$$

Die zweite Gleichung zeigt, dass der Punkt $(\bar{x}, \bar{y})$ auf der Geraden liegt. Durch Auflösen der zweiten Gleichung nach b und Einsetzen in die erste Gleichung erhalten wir

$$b = \bar{y} - a\bar{x} \quad \text{und} \quad a = \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} = \frac{(n-1)c_{xy}}{(n-1)s_x^2} = \frac{c_{xy}}{s_x^2}$$

für die Standardabweichung s_x der Messwerte $x_1, \dots, x_n$ und die Kovarianz der Messwertpaare $(x_1, y_1), \dots, (x_n, y_n)$. Für die Umformung von a haben wir Satz 1.2 und die Formel für die Kovarianz auf Seite 18 benutzt.

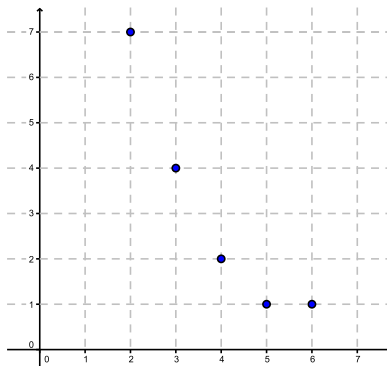
Satz 2.2 Die Regressionsgerade zu den Wertepaaren $(x_1, y_1), \dots, (x_n, y_n)$ hat die Gleichung $y = ax + b$ mit

$$a = \frac{c_{xy}}{s_x^2} \quad \text{und} \quad b = \bar{y} - a\bar{x}.$$

Der Koeffizient a wird auch als *erster Regressionskoeffizient* oder *Regressionskoeffizient bezüglich x* bezeichnet. Man beachte, dass er nicht symmetrisch in x und y ist.

Beispiele

1. Betrachten wir nochmals das 1. Beispiel von Seite 20 mit dem folgenden Streudiagramm:



Die fünf Punkte liegen fast auf einer Geraden, bzw. der Korrelationskoeffizient $r_{xy} = -0,93$ deutet auf einen linearen Zusammenhang der Wertepaare hin. Welche Gleichung hat die Regressionsgerade? Auf Seite 21 haben wir schon berechnet:

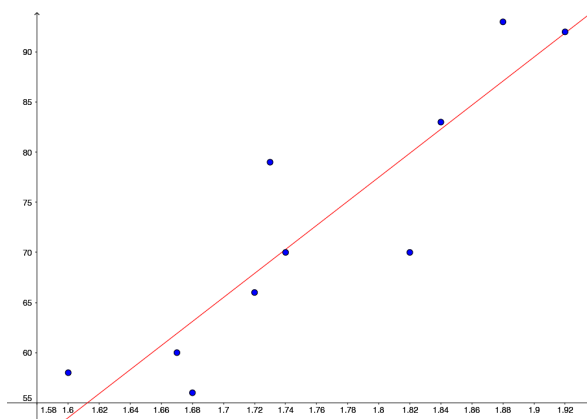
$$\bar{x} = 4, \quad \bar{y} = 3, \quad c_{xy} = \frac{-15}{4}, \quad s_x^2 = \frac{10}{4}$$

Damit erhalten wir die Steigung a und den y -Achsenabschnitt b der Regressionsgeraden

und die Gleichung der Regressionsgeraden lautet:

2. Im 2. Beispiel von Seite 21 deutet das Streudiagramm und der Korrelationskoeffizient $r_{xy} = 0,901$ darauf hin, dass das Gewicht von einer Person (zumindest eines*r Schweizer Olympia Medaillengewinners*in) von deren Körpergrösse linear abhängt. Es ist also sinnvoll, die Regressionsgerade zu berechnen:

$$y = 1,20x - 138,31$$



2.4 Nichtlineare Regression

In vielen Fällen legt das Streudiagramm von zwei Datensätzen einen nichtlinearen Ansatz nahe, zum Beispiel eine Polynomfunktion oder eine Exponentialfunktion f . Auch in diesen Fällen kann die Methode der kleinsten Quadrate verwendet werden; man minimiert die Summe über die Abweichungen im Quadrat $(y_i - f(x_i))^2$.

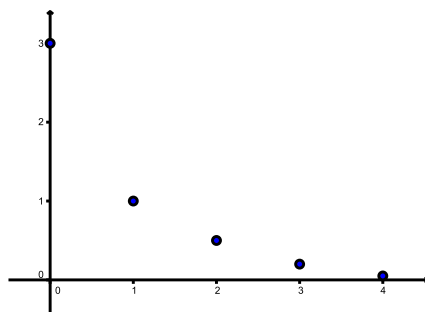
Im Fall einer Exponentialfunktion kann dieses Minimierungsproblem auf eine lineare Regression zurückgeführt werden.

Beispiel

Gegeben sind die folgenden Wertepaare

x_i	0	1	2	3	4
y_i	3	1	0,5	0,2	0,05

Streudiagramm:



Das Streudiagramm zeigt, dass die Daten y_i exponentiell von den Daten x_i abhängen könnten. Wir machen also den Ansatz

$$y = f(x) = c e^{ax}.$$

Anstatt nun die Summe über die Abweichungen im Quadrat $(y_i - f(x_i))^2$ zu minimieren, logarithmieren wir diesen Ansatz:

Das heisst, wenn zwischen den Wertepaaren (x_i, y_i) ein exponentieller Zusammenhang besteht, dann besteht zwischen den Wertepaaren $(x_i, \ln(y_i))$ ein linearer Zusammenhang. Wir können also die Regressionsgerade bestimmen für die Wertepaare

x_i	0	1	2	3	4
$\ln(y_i)$	1,099	0	-0,693	-1,609	-2,996

Wir erhalten die Regressionsgerade

$$\ln(y) = -0,9798x + 1,1197.$$

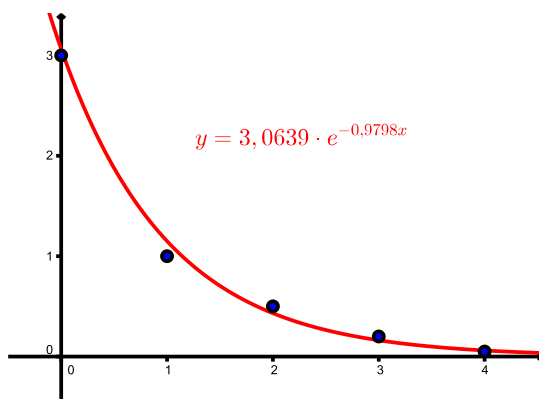
Es ist also $a = -0,9798$ und für c finden wir

$$\ln(c) = 1,1197 \implies c = e^{1,1197} = 3,0639.$$

Der exponentielle Zusammenhang zwischen den Wertepaaren kann also näherungsweise durch die Funktion

$$y = f(x) = 3,0639 \cdot e^{-0,9798x}$$

beschrieben werden.



3 Wahrscheinlichkeitsrechnung

Das Hauptziel der Stochastik ist, Modelle zur mathematischen Beschreibung von sogenannten Zufallsexperimenten (wie zum Beispiel das Würfeln, die Grösse von Messfehlern, die Qualität eines Laptops, langfristige Wettervorhersage oder die Ausbreitung einer Krankheit) zu entwickeln.

3.1 Zufallsexperimente und Ereignisse

Wenn wir eine Münze werfen, so bestimmt der Zufall, ob das Ergebnis “Kopf” oder “Zahl” sein wird. Es ist nicht vorhersagbar, wie oft in den kommenden 100 Jahren im Februar in Basel Schnee liegen wird. In beiden Fällen handelt es sich um ein Zufallsexperiment.

Definition Ein *Zufallsexperiment* ist ein Vorgang, der

- beliebig oft unter den gleichen Bedingungen wiederholt werden kann und
- dessen Ergebnis nicht mit Sicherheit vorhergesagt werden kann.

Die Menge aller möglichen (sich gegenseitig ausschliessenden) Ergebnisse des Zufallsexperiments wird *Ergebnisraum* genannt und mit Ω bezeichnet.

Eine Teilmenge $A \subseteq \Omega$ heisst *Ereignis*. Es ist eingetreten, wenn das Ergebnis des Experiments ein Element von A ist. Ein Ergebnis $\omega \in \Omega$ heisst auch *Elementarereignis*.

Beispiele

1. Werfen einer Münze:

$$\Omega = \{ \text{Kopf, Zahl} \} = \{ K, Z \}$$

2. Werfen eines Würfels:

$$\Omega = \{ 1, 2, 3, 4, 5, 6 \}$$

Ereignis A = Wurf einer geraden Zahl

Ereignis B = Wurf einer Zahl < 3

3. Werfen von zwei Münzen:

$$\Omega = \{ KK, KZ, ZK, ZZ \}$$

Ereignis A = Wurf von genau einer Zahl

Ereignis B = Wurf von mindestens einem Kopf

4. Messung der Körpergrösse eines zufällig ausgewählten Chemiestudenten:

$$\Omega = (0, \infty)$$

Ereignis A = die Körpergrösse ist grösser als 160 cm und kleiner als 180 cm

Definition Seien $A, B \subseteq \Omega$ Ereignisse.

- Das Ereignis A und B entspricht dem Durchschnitt $A \cap B$.
- Das Ereignis A oder B entspricht der Vereinigung $A \cup B$.
- Das *Gegenereignis* von A ist jenes Ereignis, das eintritt, wenn A nicht eintritt. Es wird mit $\bar{A}$ bezeichnet und entspricht dem Komplement $\bar{A} = \Omega \setminus A$.
- Zwei Ereignisse A und B heissen *unvereinbar*, wenn $A \cap B = \emptyset$ (die leere Menge), das heisst, A und B können nicht gleichzeitig eintreten.

Beispiel

Wir bestimmen $A \cap B$, $A \cup B$ und $\bar{A}$ für das 2. Beispiel oben.

3.2 Wahrscheinlichkeit

Nun ordnen wir den Ereignissen Wahrscheinlichkeiten zu. Das heisst, wir suchen eine Funktion P , die jedem Element (bzw. jeder Teilmenge) des Ereignisraums Ω eine reelle Zahl zuordnet. Diese Zahl soll der Wahrscheinlichkeit entsprechen, mit der das Ergebnis (bzw. das Ereignis) eintritt. Die Funktion P muss dabei gewissen Mindestanforderungen genügen.

Definition (Axiome von Kolmogorow) Eine Funktion P , die jedem Ereignis A von Ω eine reelle Zahl $P(A)$ zuordnet, heisst *Wahrscheinlichkeitsverteilung*, wenn sie die folgenden drei Eigenschaften erfüllt:

1. Für jedes $A \subseteq \Omega$ gilt $0 \leq P(A) \leq 1$.
2. Für das sichere Ereignis Ω gilt $P(\Omega) = 1$.
3. Für zwei unvereinbare Ereignisse A und B (d.h. falls $A \cap B = \emptyset$) gilt

$$P(A \cup B) = P(A) + P(B).$$

Setzen wir im dritten Punkt $A = \Omega$ und $B = \emptyset$, so folgt für das unmögliche Ereignis $\emptyset$, dass

$$P(\emptyset) = 0.$$

Weiter folgt aus der dritten Eigenschaft, dass zur Bestimmung der Wahrscheinlichkeit $P(A)$ eines Ereignisses A über die Wahrscheinlichkeiten $P(\omega)$ der einzelnen Ergebnisse ω von A summiert werden kann. Dabei gehen wir davon aus, dass Ω eine nicht-leere endliche oder abzählbar unendliche Menge ist (das heisst, die Elemente können durchnummeriert werden). Man nennt in diesem Fall das Paar (Ω, P) einen *diskreten Wahrscheinlichkeitsraum*.

Aber wie bestimmen wir nun $P(A)$ für ein Ereignis A ? Nun, der Ausgang eines einzelnen Zufallsexperiments ist völlig offen. Wiederholt man jedoch ein Zufallsexperiment oft (n Mal) und zählt dabei, wie oft ein bestimmtes Ereignis A eintritt (k Mal), so scheint sich die relative

Häufigkeit $\frac{k}{n}$ um einen festen Wert p zu “stabilisieren”. Dieser Wert p kann als Näherung für die Wahrscheinlichkeit $P(A)$ verwendet werden.

Beispiel

Nehmen wir einen Würfel, von dem wir nicht wissen, ob er gezinkt ist. Wir wollen herausfinden, wie gross die Wahrscheinlichkeit ist, die Augenzahl 6 zu würfeln. Dazu würfeln wir n Mal und zählen die Anzahl k der Augenzahl 6. Hier ist also $\Omega = \{1, 2, 3, 4, 5, 6\}$ und $A = \{6\}$.

n	k	relative Häufigkeit $\frac{k}{n}$
100	16	0,16
200	34	0,17
300	49	0,163
400	62	0,155

Unser Experiment zeigt, dass $P(A) \approx 0,155$.

Wäre der Würfel nicht gezinkt, dann könnten wir davon ausgehen, dass alle Augenzahlen gleich wahrscheinlich sind. Man nennt einen solchen Würfel *fair* oder *ideal*. Die Bestimmung von $P(A)$ ist in diesem Fall viel einfacher. Aus der Bedingung $P(\Omega) = 1$ folgt direkt $P(A) = \frac{1}{6}$, da 6 verschiedene Augenzahlen gewürfelt werden können und jede Augenzahl gleich wahrscheinlich ist.

Definition Ein *Laplace-Experiment* ist ein Zufallsexperiment mit den folgenden Eigenschaften:

1. Das Zufallsexperiment hat nur endlich viele mögliche Ergebnisse.
2. Jedes dieser Ergebnisse ist gleich wahrscheinlich.

Zum Beispiel sind (wie oben erwähnt) beim Wurf eines fairen Würfels alle Augenzahlen gleich wahrscheinlich. Oder bei der zufälligen Entnahme einer Stichprobe einer Warenlieferung haben alle Artikel dieselbe Wahrscheinlichkeit, gezogen zu werden.

Für eine Menge M bezeichnen wir mit $|M|$ die Anzahl Elemente dieser Menge.

Satz 3.1 Bei einem Laplace-Experiment hat jedes Ergebnis $\omega \in \Omega$ die Wahrscheinlichkeit

$$P(\omega) = \frac{1}{|\Omega|}.$$

Für jedes Ereignis $A \subset \Omega$ folgt

$$P(A) = \sum_{\omega \in A} P(\omega) = \sum_{\omega \in A} \frac{1}{|\Omega|} = \frac{|A|}{|\Omega|} = \frac{\text{Anzahl der für } A \text{ günstigen Fälle}}{\text{Anzahl der möglichen Fälle}}.$$

Beispiele

1. Es wird ein fairer Würfel geworfen. Wie gross sind die Wahrscheinlichkeiten $P(A)$ und $P(B)$ für $A = \{\text{gerade Augenzahl}\}$ und $B = \{\text{Augenzahl durch 3 teilbar}\}$?

2. Aline (A) und Beat (B) spielen wiederholt ein faires Spiel, bei dem beide die gleiche Gewinnchance haben. Sie setzen je 50 CHF ein und wer zuerst sechs Runden gewonnen hat, erhält den gesamten Einsatz von 100 CHF. Leider muss das Spiel beim Stand von 5:3 für Aline abgebrochen werden. Wie soll nun der Einsatz gerecht aufgeteilt werden? Eine Möglichkeit wäre, im Verhältnis 5:3, also Aline erhält 62,50 CHF und Beat 37,50 CHF. Dies entspricht jedoch nicht den einzelnen Gewinnwahrscheinlichkeiten, die wir wie folgt berechnen können. Würde das Spiel weitergeführt, gäbe es vier verschiedene mögliche Spielausgänge:

Spielausgang				
Gewinnreihenfolge				
Wahrscheinlichkeit				

Nur in einem der vier Spielausgänge gewinnt Beat, doch die vier Spielausgänge sind nicht gleich wahrscheinlich, also ist auch die Aufteilung 75 CHF für Aline und 25 CHF für Beat nicht sinnvoll. Die Berechnung in der Tabelle zeigt, dass 87,50 CHF für Aline und 12,50 CHF für Beat wohl am gerechtesten wären.

Die folgenden Eigenschaften, die direkt aus den drei Bedingungen an eine Wahrscheinlichkeitsverteilung folgen, sind sehr nützlich zur Bestimmung von Wahrscheinlichkeiten.

Satz 3.2 Für $A, B \subseteq \Omega$ gilt:

- (a) $P(\overline{A}) = 1 - P(A)$
- (b) $P(A \setminus B) = P(A) - P(A \cap B)$
- (c) $P(A \cup B) = P(A) + P(B) - P(A \cap B)$
- (d) $A \subseteq B \implies P(A) \leq P(B)$
- (e) $P(A) = P(A \cap B) + P(A \cap \overline{B})$

Beispiel

In einem Restaurant essen gewöhnlich 20% der Gäste Vorspeise (V) und Nachtisch (N), 45 % nehmen Vorspeise oder Nachtisch und 65% nehmen keine Vorspeise. Man bestimme den Prozentsatz der Gäste, die wie folgt wählen:

- (a) Vorspeise und keinen Nachtisch
- (b) einen Nachtisch

3.3 Bedingte Wahrscheinlichkeit

Oft ist die Wahrscheinlichkeit eines Ereignisses B unter der Bedingung (bzw. dem Wissen), dass ein Ereignis A bereits eingetreten ist, gesucht. Man bezeichnet diese Wahrscheinlichkeit mit $P(B|A)$.

Beispiel

Zwei faire Würfel werden geworfen. Wie gross ist die Wahrscheinlichkeit, die Augensumme 5 zu werfen unter der Bedingung, dass wenigstens einmal die Augenzahl 1 geworfen wird?

Bei Laplace-Experimenten kann man stets so wie im Beispiel vorgehen. Für beliebige Zufallsexperimente definieren wir die Wahrscheinlichkeit $P(B|A)$ durch die eben gefundene Formel.

Definition Die Wahrscheinlichkeit des Ereignisses B unter der Bedingung, dass Ereignis A eingetreten ist, ist definiert als

$$P(B|A) = \frac{P(A \cap B)}{P(A)}.$$

Man spricht von der *bedingten Wahrscheinlichkeit* $P(B|A)$.

Der ursprüngliche Ergebnisraum Ω reduziert sich also auf A , und von B sind nur jene Ergebnisse zu zählen, die auch in A liegen.

Formt man die Gleichung in der Definition um, erhält man eine nützliche Formel für die Wahrscheinlichkeit $P(A \cap B)$.

Satz 3.3 (Multiplikationssatz) *Gegeben sind Ereignisse A und B mit Wahrscheinlichkeiten ungleich Null. Dann gilt*

$$P(A \cap B) = P(A)P(B|A) = P(B)P(A|B).$$

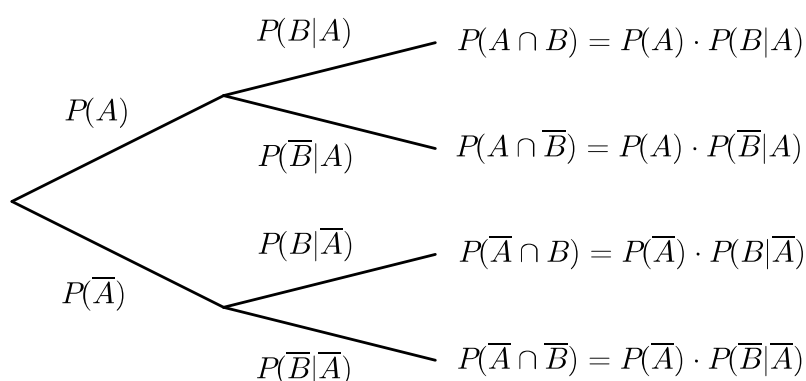
Die zweite Gleichheit im Satz folgt, indem wir die Rollen von A und B in der Definition vertauschen.

Beispiel

In einer Urne befinden sich 5 rote und 10 blaue Kugeln. Wir entnehmen nun hintereinander zufällig zwei Kugeln ohne Zurücklegen. Wie gross ist die Wahrscheinlichkeit,

- (a) zuerst eine rote und dann eine blaue Kugel und
- (b) zwei rote oder zwei blaue Kugeln zu ziehen?

Oft ist es hilfreich, die Wahrscheinlichkeiten mit Hilfe eines *Wahrscheinlichkeitsbaums* zu veranschaulichen:



Dabei sind A und B zwei beliebige Ereignisse.

Beispiele

1. Die Studentin Maja wohnt im Studentenheim Basilea. Dieses besitzt eine Brandmeldeanlage, welche bei Feuerausbruch mit einer Wahrscheinlichkeit von 99 % Alarm gibt. Manchmal gibt die Anlage einen Fehlalarm, und zwar in etwa 2 % aller Nächte. Schliesslich ist die Wahrscheinlichkeit, dass in einer bestimmten Nacht Feuer ausbricht, gleich 0,05 %.

- (a) Mit welcher Wahrscheinlichkeit kann Maja diese Nacht ruhig schlafen?
- (b) Mit welcher Wahrscheinlichkeit geht diese Nacht die Alarmanlage los?
- (c) Maja hört den Feueralarm. Mit welcher Wahrscheinlichkeit brennt es wirklich?

Wir zeichnen dazu einen Wahrscheinlichkeitsbaum:

Wir finden damit die folgenden Antworten zu den Fragen im Beispiel:

(a) $P(\text{weder Feuer noch Alarm}) =$

(b) $P(\text{Alarm}) =$

(c) $P(\text{Feuer}|\text{Alarm}) =$

Die Wahrscheinlichkeit, dass es bei Alarm auch wirklich brennt, ist also sehr klein, nur 2,4 %. Dies im Gegensatz zur Wahrscheinlichkeit von 99 %, dass bei Feuer der Alarm auch losgeht. Man darf Ereignis und Bedingung also nicht verwechseln.

2. In einem Land seien 0,01 % der Bevölkerung HIV positiv. Ein HIV-Test reagiert bei HIV positiven Personen mit 99,9 % Wahrscheinlichkeit positiv. Bei HIV negativen Personen gibt er mit 0,01 % Wahrscheinlichkeit irrtümlicherweise auch ein positives Resultat.

Eine Person wird getestet und es ergibt sich ein positives Resultat. Mit welcher Wahrscheinlichkeit ist die Person wirklich HIV positiv?

3. Alle Personen eines Landes werden auf Tuberkulose getestet. Dabei erhalten 32 % ein positives Testresultat. Welcher Anteil der Bevölkerung ist tatsächlich mit Tuberkulose infiziert, wenn der Test bei infizierten Personen mit 90 % Wahrscheinlichkeit und bei nicht infizierten Personen mit 30 % Wahrscheinlichkeit ein positives Resultat gibt?

Der Trick hier ist, für die gesuchte Wahrscheinlichkeit p zu setzen, und dann wie vorher den Wahrscheinlichkeitsbaum zu zeichnen:

Damit erhalten wir die folgende Gleichung für p :

3.4 Unabhängige Ereignisse

In vielen Fällen ist die Wahrscheinlichkeit, dass ein Ereignis B eintritt, völlig unabhängig davon, ob ein anderes Ereignis A eintritt, das heisst $P(B|A) = P(B)$. Der Multiplikationssatz vereinfacht sich dadurch.

Definition Zwei Ereignisse A und B heissen (*stochastisch*) *unabhängig*, wenn gilt

$$P(A \cap B) = P(A) \cdot P(B) .$$

Äquivalent dazu heissen zwei Ereignisse A und B unabhängig, wenn

$$\begin{aligned} P(B|A) &= P(B) && \text{mit } P(A) > 0 \text{ bzw.} \\ P(A|B) &= P(A) && \text{mit } P(B) > 0 . \end{aligned}$$

Es ist nicht immer intuitiv erkennbar, ob zwei Ereignisse A und B unabhängig sind oder nicht. Die stochastische Unabhängigkeit von zwei Ereignissen A und B besagt, dass A und B

im *wahrscheinlichkeitstheoretischen Sinn* keinen Einfluss aufeinander haben. Es kann vorkommen, dass zwei Ereignisse A und B stochastisch unabhängig sind, obwohl real das Eintreten von B davon abhängt, ob A eintritt.

Beispiele

1. Ein Würfel wird zweimal geworfen. Wie gross ist die Wahrscheinlichkeit, beim ersten Wurf die Augenzahl 1 und beim zweiten Wurf die Augenzahl 2 zu würfeln?

2. Wieder werfen wir einen Würfel zweimal. Dabei sei A das Ereignis, dass die Augenzahl des ersten Wurfes gerade ist und B sei das Ereignis, dass die Summe der beiden geworfenen Augenzahlen gerade ist. Sicher entscheidet hier das Ereignis A mit, ob B eintritt. Sind A und B stochastisch unabhängig?

4 Erwartungswert und Varianz von Zufallsgrössen

Bei vielen Zufallsexperimenten (wie beispielsweise beim Würfeln oder bei Messfehlern) geht es um die Wahrscheinlichkeit, dass eine bestimmte Zahl auftritt. Bei anderen Zufallsexperimenten kann jedem Ergebnis eine Zahl *zugeordnet* werden (zum Beispiel ein Geldbetrag bei einem Glücksspiel oder das Gewicht einer zufällig aus einer Packung entnommenen Tablette). In beiden Fällen interessiert uns, welche Zahl durchschnittlich auftritt, wenn das Zufallsexperiment oft wiederholt wird. Diese Zahl nennt man Erwartungswert.

4.1 Zufallsgrösse und Erwartungswert

Wir beginnen mit einem **Beispiel**.

Der Händler A verkauft ein Laptop ohne Garantie für 500 CHF. Der Händler B verkauft dasselbe Modell mit einer Garantie von einem Jahr für 550 CHF. Bei einem Schaden des Laptops wird dieses kostenlos repariert oder durch ein neues ersetzt. Die Wahrscheinlichkeit, dass ein Laptop dieses Modells innerhalb des ersten Jahres aussteigt, beträgt 5 %.

Die Situation beim Händler A sieht so aus:

ω	gutes Laptop	schlechtes Laptop
$P(\omega)$	0,95	0,05
Kosten	500	1000

Welche (durchschnittlichen) Kosten sind bei Händler A zu erwarten?

Die Kosten sind eine sogenannte Zufallsgrösse. Die zu erwartenden Kosten nennt man den Erwartungswert der Zufallsgrösse.

Definition Sei Ω ein Ereignisraum. Eine *Zufallsgrösse* (oder *Zufallsvariable*) ist eine Funktion, die jedem Ergebnis ω aus Ω eine reelle Zahl zuordnet, $X : \Omega \rightarrow \mathbb{R}$, $\omega \mapsto X(\omega)$.

Eine Zufallsgrösse heisst *diskret*, wenn sie nur endlich viele oder abzählbar unendlich viele verschiedene Werte $x_1, x_2, x_3, \dots$ annehmen kann.

Wir gehen in diesem Kapitel stets davon aus, dass die Zufallsgrösse diskret ist.

Sei x_k der Wert der Zufallsgrösse für das Ergebnis ω_k , also $x_k = X(\omega_k)$. Dann bezeichnen wir mit p_k die Wahrscheinlichkeit, dass die Zufallsgrösse X den Wert x_k annimmt, also $p_k = P(X(\omega) = x_k)$.

Definition Der *Erwartungswert* einer diskreten Zufallsgrösse X ist definiert durch

$$\mu = E(X) = p_1x_1 + p_2x_2 + \cdots + p_nx_n = \sum_{k=1}^n p_kx_k.$$

Beispiele

1. Wir werfen einen Würfel. Beim Werfen der Augenzahl 5 gewinnt man 5 CHF, in allen anderen Fällen muss man 1 CHF bezahlen. Mit welchem durchschnittlichen Gewinn oder Verlust muss man rechnen?

$\omega = \text{Augenzahl}$	1	2	3	4	5	6
$P(\omega)$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$
Gewinn $X(\omega)$	-1	-1	-1	-1	5	-1

Die Zufallsgrösse X nimmt also nur zwei Werte an, $x_1 = -1$ und $x_2 = 5$. Mit welchen Wahrscheinlichkeiten werden diese Werte angenommen? Erwartungswert?

2. Nun gewinnt man bei jeder Augenzahl 1 CHF.

$\omega = \text{Augenzahl}$	1	2	3	4	5	6
$P(\omega)$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$
Gewinn $X(\omega)$	1	1	1	1	1	1

Das ist natürlich ein langweiliges Spiel. Ohne Rechnung erkennen wir, dass der erwartete Gewinn, (d.h. $\mu = E(X)$) 1 CHF beträgt.

3. Nun gewinnt man 6 CHF bei der Augenzahl 5 und sonst 0 CHF.

$\omega = \text{Augenzahl}$	1	2	3	4	5	6
$P(\omega)$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$
Gewinn $X(\omega)$	0	0	0	0	6	0

Wie gross ist der Erwartungswert?

Die letzten beiden Beispiele haben also denselben Erwartungswert, doch das 3. Beispiel verspricht deutlich mehr Spannung als das 2. Beispiel. Dies wird durch die sogenannte Varianz der Zufallsgrösse beschrieben.

4.2 Varianz und Standardabweichung

Die Varianz $\sigma^2 = \text{Var}(X)$ einer Zufallsgrösse X misst die mittlere quadratische Abweichung vom Erwartungswert $\mu = E(X)$.

Definition Die *Varianz* einer Zufallsgrösse X ist definiert durch

$$\sigma^2 = \text{Var}(X) = E((X - \mu)^2) = \sum_{k=1}^n p_k (x_k - \mu)^2.$$

Die *Standardabweichung* oder *Streuung* σ ist definiert als die positive Quadratwurzel der Varianz, das heisst

$$\sigma = \sqrt{\text{Var}(X)} = \sqrt{E((X - \mu)^2)} = \sqrt{\sum_{k=1}^n p_k (x_k - \mu)^2}.$$

Betrachten wir nun nochmals das 2. und das 3. Beispiel von vorher. Im 2. Beispiel gibt es keine Streuung. Wir haben nur einen Wert $x_1 = 1$ und somit ist $x_1 - \mu = 0$, das heisst $\sigma^2 = \text{Var}(X) = 0$. Im 3. Beispiel sieht es anders aus:

Wie für die Varianz der beschreibenden Statistik können wir die Formel für die Varianz umformen:

$$\begin{aligned} \text{Var}(X) &= \sum_{k=1}^n p_k (x_k - \mu)^2 = \sum_{k=1}^n p_k (x_k^2 - 2x_k\mu + \mu^2) \\ &= \underbrace{\sum_{k=1}^n p_k x_k^2}_{=E(X^2)} - 2\mu \underbrace{\sum_{k=1}^n p_k x_k}_{=\mu} + \mu^2 \underbrace{\sum_{k=1}^n p_k}_{=1} = E(X^2) - \mu^2 = E(X^2) - (E(X))^2. \end{aligned}$$

Satz 4.1 Es gilt

$$\sigma^2 = \text{Var}(X) = E(X^2) - (E(X))^2.$$

Beispiel

4. Wieder werfen wir einen Würfel. Der Gewinn $X(\omega)$ entspricht nun genau der gewürfelten Augenzahl. Wie oft müssen wir würfeln, um (durchschnittlich) einen Gewinn von 1000 CHF einstreichen zu können?

ω	1	2	3	4	5	6
$P(\omega)$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$	$\frac{1}{6}$
$X(\omega) = x_k$	1	2	3	4	5	6

Wir berechnen zunächst den Erwartungswert und die Streuung:

Bei jedem Wurf können wir also mit einem Gewinn von 3.50 CHF rechnen. Für einen Gewinn von 1000 CHF müssen wir demnach (durchschnittlich)

würfeln.

In all den bisherigen Beispielen waren die Wahrscheinlichkeiten der Ergebnisse ω jeweils gleich gross. Das muss nicht so sein.

Beispiel

5. Wir werfen zwei Würfel gleichzeitig. Als Zufallsgrösse wählen wir die halbe Augensumme (d.h. der Durchschnitt der beiden geworfenen Augenzahlen).

ω	2	3	4	5	6	7	8	9	10	11	12
$P(\omega) = p_k$	$\frac{1}{36}$	$\frac{2}{36}$	$\frac{3}{36}$	$\frac{4}{36}$	$\frac{5}{36}$	$\frac{6}{36}$	$\frac{5}{36}$	$\frac{4}{36}$	$\frac{3}{36}$	$\frac{2}{36}$	$\frac{1}{36}$
$X(\omega) = x_k$	1	1,5	2	2,5	3	3,5	4	4,5	5	5,5	6

Für den Erwartungswert erhalten wir

$$\mu = E(X) = \frac{1}{36} \cdot 1 + \frac{2}{36} \cdot 1,5 + \frac{3}{36} \cdot 2 + \frac{4}{36} \cdot 2,5 + \dots + \frac{1}{36} \cdot 6 = \frac{126}{36} = 3,5$$

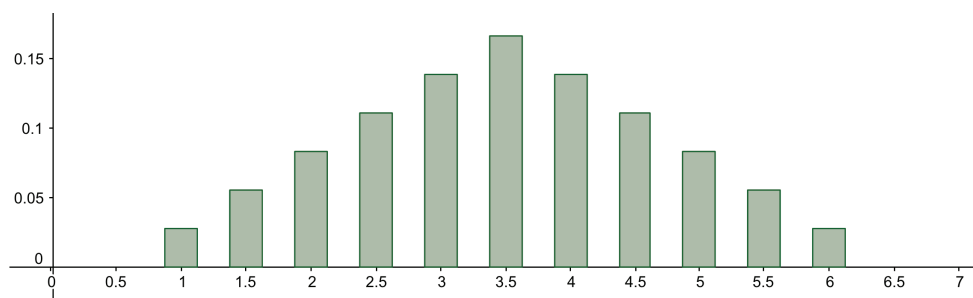
genau wie im 4. Beispiel. Die Varianz und die Streuung sind nun allerdings kleiner als im 4. Beispiel. Es gilt

$$\sigma^2 = \text{Var}(X) = E(X^2) - (E(X))^2 = \frac{1}{36} \cdot 1^2 + \frac{2}{36} \cdot 1,5^2 + \frac{3}{36} \cdot 2^2 + \dots + \frac{1}{36} \cdot 6^2 - 3,5^2 \approx 1,46$$

und damit ist $\sigma \approx 1,21$.

Eine diskrete Zufallsgrösse kann man auch graphisch darstellen. In einem Stabdiagramm errichtet man über jedem Wert x_k einen Stab der Länge p_k .

Für das letzte Beispiel sieht das so aus:



Definition Sei X eine diskrete Zufallsgrösse. Man nennt die Menge

$$\{ (x_1, p_1), (x_2, p_2), (x_3, p_3), \dots \}$$

die *Wahrscheinlichkeitsverteilung* oder *Verteilung* von X .

4.3 Kombination von Zufallsgrössen

Zwei Zufallsgrössen X und Y können addiert und multipliziert werden. Wie hängen der Erwartungswert und die Varianz der neuen Zufallsgrösse von X und Y ab?

Beispiel

Wieder würfeln wir. Es gilt also $P(\omega) = \frac{1}{6}$ für jedes Ergebnis ω .

ω	1	2	3	4	5	6
Zufallsgrösse X	1	2	2	3	5	5
Zufallsgrösse Y	0	1	1	0	0	10
$X + Y$						
$X \cdot Y$						

Nun vergleichen wir die Erwartungswerte der verschiedenen Zufallsgrössen. Zunächst gilt $E(X) = 3$ und $E(Y) = 2$. Weiter finden wir

Bei der Addition der zwei Zufallsgrössen haben sich also deren Erwartungswerte ebenfalls addiert. Dies gilt allgemein, und zwar gilt noch ein wenig mehr.

Satz 4.2 Für zwei Zufallsgrössen X, Y und reelle Zahlen a, b, c gilt

$$E(aX + bY + c) = aE(X) + bE(Y) + c.$$

Der Beweis erfolgt durch Nachrechnen:

$$\begin{aligned} E(aX + bY + c) &= p_1(ax_1 + by_1 + c) + \dots + p_n(ax_n + by_n + c) \\ &= a(p_1x_1 + \dots + p_nx_n) + b(p_1y_1 + \dots + p_ny_n) + c \underbrace{(p_1 + \dots + p_n)}_{=1} \\ &= aE(X) + bE(Y) + c. \end{aligned}$$

Mit der Multiplikation von zwei Zufallsgrössen scheint es nicht so einfach zu gehen. Schauen wir uns nochmals ein Beispiel an.

Beispiel

ω	1	2	3	4	5	6
Zufallsgrösse X	3	3	5	5	5	3
Zufallsgrösse Y	1	3	2	3	1	2
$X \cdot Y$	3	9	10	15	5	6

Es gilt $E(X) = 4$ und $E(Y) = 2$. Für $E(X \cdot Y)$ erhalten wir

In diesem Beispiel sind die Werte von Y gleichmässig über die Werte von X verteilt und umgekehrt, das heisst, die Werte von Y sind unabhängig von den Werten von X . Man nennt die Zufallsgrössen stochastisch unabhängig. Die präzise Definition hat mit unabhängigen Ereignissen zu tun.

Definition Zwei Zufallsgrössen X und Y heissen *stochastisch unabhängig*, falls für alle x_k und y_ℓ die Ereignisse $(X = x_k)$ und $(Y = y_\ell)$ unabhängig sind, also falls

$$P((X = x_k) \text{ und } (Y = y_\ell)) = P(X = x_k) \cdot P(Y = y_\ell).$$

Wollen wir überprüfen, dass im zweiten Beispiel die Zufallsgrössen X und Y stochastisch unabhängig sind, dann müssen wir sechs Gleichungen nachweisen:

$$\begin{aligned}P((X = 3) \text{ und } (Y = 1)) &= P(X = 3) \cdot P(Y = 1) \\P((X = 3) \text{ und } (Y = 2)) &= P(X = 3) \cdot P(Y = 2) \\P((X = 3) \text{ und } (Y = 3)) &= P(X = 3) \cdot P(Y = 3)\end{aligned}$$

und dann nochmals die drei Gleichungen, wobei wir $X = 3$ durch $X = 5$ ersetzen. Wir überprüfen hier nur die erste Gleichung:

Im Gegensatz dazu sind X und Y vom ersten Beispiel nicht stochastisch unabhängig. Um dies nachzuweisen, genügt es, ein Pärchen (x_k, y_ℓ) zu finden, welches die Gleichung in der Definition nicht erfüllt.

Satz 4.3 Sind die Zufallsgrössen X und Y stochastisch unabhängig, dann gilt

$$E(X \cdot Y) = E(X) \cdot E(Y) .$$

Wegen der Formel $Var(X) = E(X^2) - (E(X))^2$ können auch Aussagen über die Varianz von Kombinationen von Zufallsgrössen gemacht werden. Im folgenden Satz sind nun alle Regeln zu Erwartungswert und Varianz zusammengestellt.

Satz 4.4 Seien X, Y Zufallsgrössen und a, b, c reelle Zahlen. Dann gilt:

- (1) $E(aX + bY + c) = aE(X) + bE(Y) + c$
- (2) $Var(aX + c) = a^2 Var(X)$

Falls X und Y stochastisch unabhängig sind, gilt weiter:

- (3) $E(X \cdot Y) = E(X) \cdot E(Y)$
- (4) $Var(X + Y) = Var(X) + Var(Y)$

4.4 Schätzen von Erwartungswert und Varianz

Ein quantitatives Merkmal X einer Grundgesamtheit kann als Zufallsgrösse aufgefasst werden. Interessiert man sich für den Erwartungswert und die Varianz von X , dann können diese beiden Grössen nur dann berechnet werden, wenn die Anzahl Elemente N der Grundgesamtheit nicht zu gross ist. Andernfalls, wenn N sehr gross oder unendlich ist, muss man sich mit einer Schätzung von Erwartungswert und Varianz begnügen. Wie dies zu verstehen ist, wird hier anhand eines Beispiels gezeigt.

Betrachten wir als Grundgesamtheit zum Beispiel die Menge aller Studierenden der Vorlesung Mathematik II für Naturwissenschaften. Das Merkmal, für das wir uns interessieren, sei das Alter. Es geht hier also um die Zufallsgrösse $X = (\text{Alter eines*r zufällig ausgewählten Studierenden der Grundgesamtheit})$. Interessieren wir uns für das durchschnittliche Alter der Studierenden, dann entspricht dies dem Erwartungswert

$$\mu = E(X) = \frac{1}{N}(x_1 + \dots + x_N) ,$$

wobei x_i das Alter des*r i -ten Studierenden ist.

Das Überprüfen des Alters von jedem*r Studierenden ist nun allerdings zu aufwendig. Deshalb entnehmen wir eine zufällige Stichprobe vom Umfang n (n klein gegenüber der Anzahl N der Studierenden) und versuchen damit, das unbekannte Durchschnittsalter der Grundgesamtheit zu schätzen. Eine solche *Zufallsstichprobe* vom Umfang n ist eine Folge von unabhängigen, identisch verteilten Zufallsgrössen $(X_1, X_2, \dots, X_n)$, wobei X_i die Merkmalsausprägung (hier also die vorkommenden Alter) des i -ten Elementes in der Stichprobe bezeichnet. Identisch verteilt bedeutet insbesondere, dass die Erwartungswerte und die Varianzen der X_i übereinstimmen, das heisst, $E(X_i) = \mu$ und $Var(X_i) = \sigma^2$ für alle i . Wenn N klein ist, sind die $X_1, X_2, \dots, X_n$ nur dann unabhängig und identisch verteilt, wenn die Studierenden mit Zurücklegen ausgewählt werden. Wir gehen hier jedoch von einem sehr grossen N aus, so dass wir von fast unabhängigen und identisch verteilten Zufallsgrössen $X_1, X_2, \dots, X_n$ ausgehen können, auch wenn wir Studierende ohne Zurücklegen auswählen.

Wird eine Stichprobe gezogen, so nehmen $X_1, \dots, X_n$ die konkreten Werte $x_1, \dots, x_n$ an.

Als *Schätzfunktion* $\hat{\mu}$ für das unbekannte Durchschnittsalter μ wählen wir das arithmetische Mittel $\bar{X}$ der Stichprobe,

$$\hat{\mu} = \bar{X} = \frac{1}{n}(X_1 + \dots + X_n).$$

Erhalten wir beispielsweise die konkrete Stichprobe (20, 22, 19, 20, 24), dann ist das arithmetische Mittel davon $\bar{x} = 21$. Dieser Wert hängt jedoch von der gewählten Stichprobe ab. Daher dürfen wir nicht davon ausgehen, dass er die gesuchte Zahl μ genau trifft. Wir erwarten jedoch von einer guten Schätzfunktion, dass die Schätzwerte wenigstens im Mittel richtig sind. Und tatsächlich gilt (mit Satz 4.4)

Für die Varianz erhalten wir

Mit wachsender Stichprobengröße n wird die Streuung also immer kleiner.

Für die Varianz $\sigma^2 = \text{Var}(X)$ der Grundgesamtheit wählen wir als Schätzfunktion die empirische Varianz s^2 der Stichprobe,

$$\hat{\sigma}^2 = s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \hat{\mu})^2 = \frac{1}{n-1} \left(\sum_{i=1}^n X_i^2 - n\hat{\mu}^2 \right).$$

Auch hier erwarten wir, dass wenigstens der Erwartungswert von $\hat{\sigma}^2$ mit der Varianz σ^2 übereinstimmt. Wir rechnen dies nach. Wegen Satz 4.1 gilt

$$\begin{aligned} E(X_i^2) &= \text{Var}(X_i) + (E(X_i))^2 = \sigma^2 + \mu^2 \\ E(\hat{\mu}^2) &= \text{Var}(\hat{\mu}) + (E(\hat{\mu}))^2 = \frac{\sigma^2}{n} + \mu^2. \end{aligned}$$

Damit folgt

$$\begin{aligned} E(\hat{\sigma}^2) &= \frac{1}{n-1} \left(\sum_{i=1}^n E(X_i^2) - nE(\hat{\mu}^2) \right) = \frac{1}{n-1} \left(\sum_{i=1}^n (\sigma^2 + \mu^2) - n \left(\frac{\sigma^2}{n} + \mu^2 \right) \right) \\ &= \frac{1}{n-1} (n\sigma^2 + n\mu^2 - \sigma^2 - n\mu^2) \\ &= \sigma^2. \end{aligned}$$

Genau aus diesem Grund haben wir in Kapitel 1 in der Definition der empirischen Varianz durch $n-1$ dividiert und nicht durch die naheliegendere Zahl n !

Würden wir die empirische Varianz mit dem Faktor $\frac{1}{n}$ definieren, nämlich als

$$\tilde{s}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \hat{\mu})^2 = \frac{n-1}{n} s^2,$$

dann würden wir damit die Varianz σ^2 systematisch unterschätzen, denn

$$E(\tilde{s}^2) = \frac{n-1}{n} E(s^2) = \sigma^2 - \frac{\sigma^2}{n}.$$

5 Binomial- und Poissonverteilung

In diesem Kapitel untersuchen wir zwei wichtige diskrete Verteilungen (d.h. Verteilungen von diskreten Zufallsgrößen): die Binomial- und die Poissonverteilung.

5.1 Die Binomialverteilung

Für die Binomialverteilung brauchen wir die Binomialkoeffizienten, die aus der Schule bekannt sein sollten. Wir frischen hier das Wichtigste darüber kurz auf.

Binomialkoeffizienten

Sei $n \geq 0$ in $\mathbb{Z}$.

Satz 5.1 *Es gibt $n!$ verschiedene Möglichkeiten, n Elemente anzuordnen.*

Jede Anordnung heisst *Permutation* der n Elemente. Es gibt also $n!$ Permutationen von n Elementen. Dabei gilt

$$n! = n \cdot (n-1) \cdot \dots \cdot 2 \cdot 1 \quad \text{für } n \geq 1 \quad \text{und} \quad 0! = 1.$$

Satz 5.2 *Es gibt*

$$n \cdot (n-1) \cdot \dots \cdot (n-k+1) = \frac{n!}{(n-k)!}$$

Möglichkeiten, aus n Elementen k auszuwählen und diese anzuordnen.

Beispiel

Wie gross ist die Wahrscheinlichkeit, dass unter 23 Personen (mindestens) zwei am gleichen Tag Geburtstag haben? Diese Frage ist als *Geburtstagsparadoxon* bekannt.

Wieviele verschiedene Möglichkeiten gibt es, aus n Elementen k auszuwählen? Wir wählen also wieder aus n Elementen k aus, aber die Anordnung dieser k ausgewählten Elemente spielt keine Rolle. Offensichtlich gibt es nun weniger Möglichkeiten. Wir müssen durch die Anzahl der Anordnungsmöglichkeiten, nämlich $k!$, dividieren.

Satz 5.3 *Es gibt*

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

Möglichkeiten, aus n Elementen k auszuwählen.

Der Ausdruck

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

heisst *Binomialkoeffizient*.

Wenn Sie auf Ihrem Taschenrechner keine Taste zur Berechnung von Binomialkoeffizienten haben, sollten Sie den linken Ausdruck von Satz 5.3 zur Berechnung benutzen.

Beispiele

Bernoulli-Experimente

Definition Ein Zufallsexperiment mit genau zwei möglichen Ausgängen heisst *Bernoulli-Experiment*.

Die beiden Ausgänge können oft als “Erfolg” (E) und “Misserfolg” (M) interpretiert werden.

Beispiel

Beim Wurf eines Würfels wollen wir nur wissen, ob die Augenzahl 2 geworfen wird oder nicht. Es gilt also $P(\text{Erfolg}) = \frac{1}{6}$.

Definition Eine *Bernoulli-Kette* ist eine Folge von gleichen Bernoulli-Experimenten. Wird ein Bernoulli-Experiment n -mal hintereinander ausgeführt, so spricht man von einer Bernoulli-Kette der *Länge* n .

Beispiel

Wir werfen einen Würfel viermal hintereinander. “Erfolg” sei wieder der Wurf der Augenzahl 2. Bei jedem einzelnen Wurf gilt also $P(\text{Erfolg}) = \frac{1}{6}$. Bei vier Würfeln können zwischen 0 und 4 Erfolge eintreten. Wie gross sind die Wahrscheinlichkeiten dafür?

Schauen wir uns die Wahrscheinlichkeit für genau 2 Erfolge genauer an.

Analog finden wir für genau k Erfolge die Wahrscheinlichkeiten

$$P_4(k) = P(k\text{-mal Erfolg}) = \binom{4}{k} \left(\frac{1}{6}\right)^k \left(\frac{5}{6}\right)^{4-k}.$$

Binomialverteilung

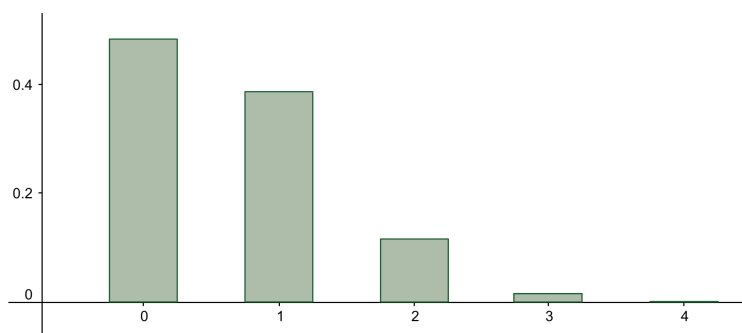
Definiert man im vorhergehenden Beispiel die Zufallsgrösse

$$X = (\text{Anzahl der Erfolge}),$$

so nimmt X die Werte $x_k = k = 0, 1, 2, 3$ oder 4 an und für die zugehörigen Wahrscheinlichkeiten gilt

$$p_k = P(X = k) = P_4(k).$$

Diese Wahrscheinlichkeitsverteilung ist ein Beispiel einer Binomialverteilung. Graphisch sieht sie so aus:



Definition Gegeben sei eine Bernoulli-Kette der Länge n , wobei Erfolg im einzelnen Experiment mit der Wahrscheinlichkeit p eintritt. Sei X die Anzahl Erfolge in den n Experimenten. Dann ist die Wahrscheinlichkeit von k Erfolgen gleich

$$P(X = k) = P_n(k) = \binom{n}{k} p^k (1-p)^{n-k}.$$

Man nennt die Zufallsgrösse X *binomialverteilt* und ihre Wahrscheinlichkeitsverteilung *Binomialverteilung* mit den Parametern n, p .

Weiter ist die Wahrscheinlichkeit, in n gleichen Bernoulli-Experimenten *höchstens* ℓ Erfolge zu haben, gleich

$$P_n(k \leq \ell) = P_n(0) + P_n(1) + \cdots + P_n(\ell) = \sum_{k=0}^{\ell} P_n(k).$$

Für die Berechnung der Wahrscheinlichkeiten $P_n(k)$ und $P_n(k \leq \ell)$ können die Tabellen in den Formelsammlungen oder die Tabellen von Hans Walser benutzt werden.

Wegen $P_n(0) + P_n(1) + \cdots + P_n(n) = 1$ (eine bestimmte Anzahl von Erfolgen tritt ja mit Sicherheit ein), gilt

$$P_n(k \geq \ell) = 1 - P(k \leq \ell - 1).$$

In den Tabellen sind die Binomialverteilungen nur für Wahrscheinlichkeiten $p \leq 0,5$ aufgeführt. Ist die Wahrscheinlichkeit eines Erfolgs gleich $p > 0,5$, so muss mit der Wahrscheinlichkeit des Misserfolgs $q = 1 - p < 0,5$ gerechnet werden.

Beispiele

1. Ein Würfel wird 10-mal geworfen. Erfolg sei das Werfen der Augenzahl 2.

- $P(2\text{-mal Erfolg}) =$

- $P(\text{höchstens 2-mal Erfolg}) =$

- $P(\text{mindestens 3-mal Erfolg}) =$

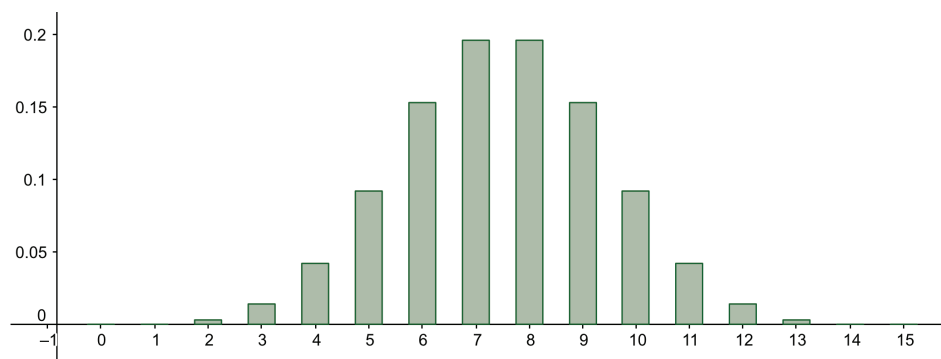
- $P(4 \leq k \leq 8) =$

- $P(7\text{-mal Misserfolg}) =$

2. Eine Münze wird 15-mal geworfen, also ist $n = 15$ und $p = 1 - p = \frac{1}{2}$.

- $P(9\text{-mal Kopf}) =$

Wegen $p = 1 - p$ ist bei diesem Beispiel die Binomialverteilung symmetrisch um die Werte $k = 7$ und 8 :



Erwartungswert und Varianz

Mit welcher Anzahl von Erfolgen können wir durchschnittlich in unserer Bernoulli-Kette rechnen? Wie gross ist die Varianz?

Um diese Fragen zu beantworten, schreiben wir die (binomialverteilte) Zufallsgrösse X als Summe $X = X_1 + \dots + X_n$ von unabhängigen (und identisch verteilten) Zufallsgrössen X_i , wobei X_i gleich 1 ist, falls der Erfolg im i -ten Experiment eingetreten ist, und 0 sonst. Für den Erwartungswert und die Varianz von X_i gilt damit

$$\begin{aligned} E(X_i) &= p \cdot 1 + (1 - p) \cdot 0 = p \\ \text{Var}(X_i) &= E(X_i^2) - (E(X_i))^2 = p - p^2 = p(1 - p). \end{aligned}$$

Mit Satz 4.4 folgt

$$\begin{aligned} E(X) &= E(X_1 + \dots + X_n) = E(X_1) + \dots + E(X_n) = np \\ \text{Var}(X) &= \text{Var}(X_1 + \dots + X_n) = \text{Var}(X_1) + \dots + \text{Var}(X_n) = np(1 - p). \end{aligned}$$

Satz 5.4 Für eine binomialverteilte Zufallsgrösse X gilt

$$\begin{aligned} E(X) &= np \\ \text{Var}(X) &= np(1 - p). \end{aligned}$$

Beispiele

1. Im ersten Beispiel von vorher (10-maliger Wurf eines Würfels) erhalten wir

Durchschnittlich können wir also mit 1,67 Erfolgen bei 10 Würfeln rechnen.

2. Im zweiten Beispiel von vorher (15-maliger Wurf einer Münze) erhalten wir

$$\begin{aligned} E(X) &= 15 \cdot \frac{1}{2} = 7,5 \\ \text{Var}(X) &= 15 \cdot \frac{1}{2} \cdot \frac{1}{2} = 3,75 \quad \implies \quad \sigma = \sqrt{\text{Var}(X)} \approx 1,94. \end{aligned}$$

In diesem Beispiel ist die Binomialverteilung also symmetrisch um den Erwartungswert.

5.2 Die Poissonverteilung

In den Jahren 2014 – 2017 gab es im Kanton Basel-Stadt durchschnittlich 10 Verkehrsunfälle pro Jahr wegen Bedienung des Telefons während der Fahrt. Mit welcher Wahrscheinlichkeit wird es im Jahr 2021 genau 6 Verkehrsunfälle mit derselben Ursache geben?

Pro Monat erhält eine Person durchschnittlich 10 Werbeanrufe. Mit welcher Wahrscheinlichkeit erhält diese Person im nächsten Monat 6 Werbeanrufe?

Die gesuchte Wahrscheinlichkeit ist für beide Fragen dieselbe. In beiden Situationen kennen wir die durchschnittliche Anzahl von “Erfolgen” pro Zeiteinheit. Wir haben jedoch keine Kenntnis über die Anzahl der Experimente (Anzahl Autofahrten, bzw. Anzahl Telefonanrufe). Wir können aber davon ausgehen, dass n gross ist. Wir kennen auch die Wahrscheinlichkeit p des Erfolgs im einzelnen Experiment nicht. Doch wir nehmen an, dass p klein ist. Man nennt solche Situationen “seltene Ereignisse”.

Die bekannte durchschnittliche Anzahl von Erfolgen bezeichnet man mit λ . Die Wahrscheinlichkeit $P(k)$, dass in einer bestimmten Zeiteinheit (oder Längeneinheit, Flächeneinheit, usw.) genau k Erfolge eintreten, ist gegeben durch

$$P(k) = \frac{\lambda^k}{k!} e^{-\lambda}.$$

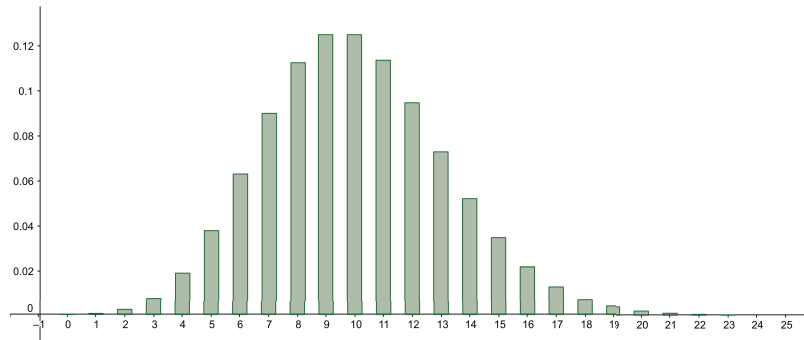
Für die beiden Beispiele finden wir also die Wahrscheinlichkeit

Definition Eine Zufallsgrösse X , die jeden der Werte $k = 0, 1, 2, \dots$ mit den Wahrscheinlichkeiten

$$P(X = k) = P(k) = \frac{\lambda^k}{k!} e^{-\lambda}$$

annehmen kann, heisst *poissonverteilt* mit dem Parameter λ . Die zugehörige Verteilung heisst *Poissonverteilung*.

Für die beiden Beispiele sieht die Verteilung so aus:



Nicht überraschend ist hier $P(k)$ am grössten für $k = \lambda = 10$, die durchschnittliche Anzahl von Erfolgen (es gilt $P(10) = 0,12511$). Wir werden unten gleich nachweisen, dass λ der Erwartungswert ist.

Es fällt weiter auf, dass $P(9)$ genau so gross wie $P(10)$ ist. Allgemein gilt $P(\lambda - 1) = P(\lambda)$, falls λ eine ganze Zahl ist, denn

Erwartungswert und Varianz

Für den Erwartungswert berechnen wir

$$\mu = E(X) = \sum_{k=0}^{\infty} P(k) \cdot k = \sum_{k=0}^{\infty} \frac{\lambda^k}{k!} e^{-\lambda} \cdot k = e^{-\lambda} \sum_{k=0}^{\infty} \frac{k \lambda^k}{k!}.$$

Da der erste Summand ($k = 0$) null ist, folgt

$$\mu = e^{-\lambda} \sum_{k=1}^{\infty} \frac{k \lambda^k}{k!} = e^{-\lambda} \sum_{k=1}^{\infty} \frac{\lambda^k}{(k-1)!} = \lambda e^{-\lambda} \sum_{k=1}^{\infty} \frac{\lambda^{k-1}}{(k-1)!}.$$

Die letzte Summe ist nichts anderes als $1 + \lambda + \frac{\lambda^2}{2!} + \frac{\lambda^3}{3!} + \dots = e^{\lambda}$, also erhalten wir

$$\mu = \lambda e^{-\lambda} e^{\lambda} = \lambda.$$

Die Varianz kann ähnlich berechnet werden (vgl. Übungsblatt 4).

Satz 5.5 *Für eine poissonverteilte Zufallsgrösse X gilt*

$$\begin{aligned} E(X) &= \lambda \\ \text{Var}(X) &= \lambda. \end{aligned}$$

Näherung für die Binomialverteilung

Ist bei einer Binomialverteilung die Anzahl n der Bernoulli-Experimente gross und gleichzeitig die Wahrscheinlichkeit p des Erfolgs im Einzelexperiment sehr klein, dann kann die Poissonverteilung mit dem Parameter $\lambda = np$ als Näherung für die Binomialverteilung benutzt werden. Tatsächlich ist diese Näherung normalerweise bereits für $n > 10$ und $p < 0,05$ ausreichend genau.

Beispiel

Eine Maschine stellt Artikel her. Aus Erfahrung weiss man, dass darunter 4 % defekte Artikel sind. Die Artikel werden in Kisten zu je 100 Stück verpackt. Wie gross ist die Wahrscheinlichkeit, dass in einer zufällig ausgewählten Kiste genau 5 defekte Artikel sind?

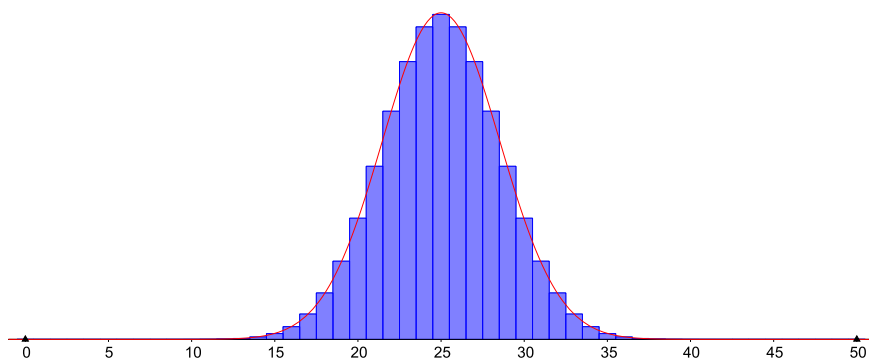
1. Exakte Berechnung mit der Binomialverteilung:

2. Näherung mit der Poissonverteilung:

6 Die Normalverteilung

Im letzten Kapitel haben wir die Binomial- und die Poissonverteilung untersucht. Dies sind Wahrscheinlichkeitsverteilungen von diskreten Zufallsgrößen. Nehmen wir nun an, die Zufallsgröße X ordne jeder Tablette einer Packung Aspirin ihr Gewicht zu. Dann kann X (nicht abzählbar) unendlich viele verschiedene Werte annehmen, zum Beispiel jeden reellen Wert zwischen 20 mg und 30 mg. Eine solche Zufallsgröße heisst stetig. Die wichtigste Wahrscheinlichkeitsverteilung einer stetigen Zufallsgröße ist die Normalverteilung.

Weiter können wir die Normalverteilung als Näherung für die Binomialverteilung benutzen (wenn n gross genug ist). Werfen wir zum Beispiel eine Münze 50-mal und Erfolg sei der Wurf von Zahl. Dann ist die Zufallsgröße $X = (\text{Anzahl Erfolge})$ binomialverteilt mit $n = 50$ und $p = 1 - p = \frac{1}{2}$. Die Binomialverteilung (blau) sieht so aus:



Eingezeichnet in rot ist der Graph der Funktion

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2},$$

wobei $\mu = np$ und $\sigma = \sqrt{npq}$ mit $q = 1 - p$. Der Graph von f heisst (*Gaußsche*) *Glockenkurve*. Sind die Wahrscheinlichkeiten $P(X \leq k)$ einer (stetigen) Zufallsgröße X gegeben durch

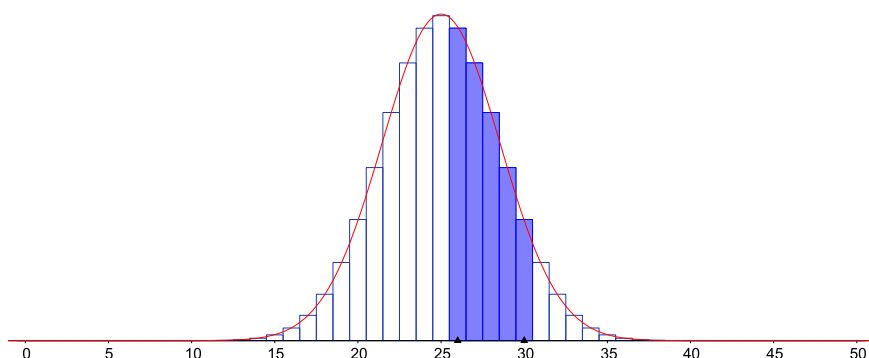
$$P(X \leq k) = \int_{-\infty}^k f(x) dx,$$

dann nennt man X normalverteilt.

Wie gross ist nun die Wahrscheinlichkeit, mit 50 Würfeln zwischen 26 und 30 Erfolge zu erzielen? Die Binomialverteilung liefert

$$P_{50}(26 \leq k \leq 30) = \sum_{k=26}^{30} P_{50}(k) = \sum_{k=26}^{30} \binom{50}{k} \left(\frac{1}{2}\right)^k \left(\frac{1}{2}\right)^{50-k},$$

doch die Wahrscheinlichkeiten $P_{50}(k)$ sind in der Tabelle nicht zu finden. Eine exakte Berechnung wäre mit einem CAS möglich, aber tatsächlich reicht eine Näherung mit Hilfe der Funktion f von oben. Die folgende Abbildung zeigt den passenden Ausschnitt aus dem Balkendiagramm.



Die blauen Rechtecke haben die Breite 1 und die Höhe $P_{50}(k)$. Die gesuchte Wahrscheinlichkeit ist also gleich dem Flächeninhalt der fünf blauen Rechtecke. Diesen Flächeninhalt können wir nun mit Hilfe des Integrals über $f(x)$ approximieren.

Allerdings haben wir nun ein neues Problem, denn die Funktion f ist nicht elementar integrierbar (d.h. ihre Stammfunktion ist nicht aus elementaren Funktionen zusammengesetzt). Wir könnten ein CAS zu Hilfe nehmen, welches Integrale über f näherungsweise berechnet. Praktischer (und in den meisten Fällen auch ausreichend genau) ist jedoch die Verwendung von Tabellen. Wie das funktioniert, untersuchen wir im nächsten Abschnitt. Danach werden wir bereit sein, Wahrscheinlichkeiten von Binomialverteilungen zu approximieren und Wahrscheinlichkeiten von Normalverteilungen zu berechnen.

6.1 Eigenschaften der Glockenkurve

Wie im Beispiel oben ersichtlich, hat die Glockenkurve ein globales Maximum und zwei Wendepunkte.

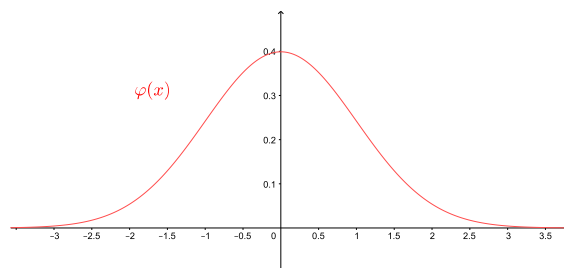
Satz 6.1 *Die Funktion $f(x)$ hat eine (lokale und globale) Maximalstelle in $x = \mu$ und zwei Wendestellen in $x = \mu \pm \sigma$.*

Im speziellen Fall $\mu = 0$ und $\sigma = 1$ wird $f(x)$ zur Funktion

$$\varphi(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2},$$

deren Graphen man *Standardglockenkurve* nennt. Die Funktion $\varphi(x)$ hat die folgenden Eigenschaften:

- In $(0, \frac{1}{\sqrt{2\pi}})$ hat $\varphi(x)$ ein Maximum.
- In $x = \pm 1$ hat $\varphi(x)$ zwei Wendestellen.
- Die Standardglockenkurve ist symmetrisch zur y -Achse, denn $\varphi(-x) = \varphi(x)$.
- Es gilt $\lim_{x \rightarrow \pm\infty} \varphi(x) = 0$.



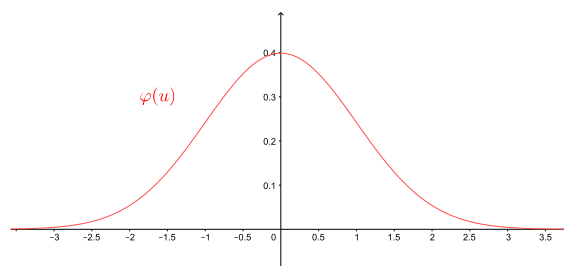
In der Tabelle (Seite 11 der Tabellen von H. Walser oder in jeder Formelsammlung) sind die Werte der Stammfunktion

$$\Phi(u) = \int_{-\infty}^u \varphi(x) dx = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^u e^{-\frac{1}{2}x^2} dx$$

zu finden. Graphisch gesehen gilt:

$\Phi(u)$ = Flächeninhalt links von u zwischen x -Achse und Graph von φ

Beispiel



Die Werte aller Integrale über der Funktion $\varphi(x)$ genügen, da wir jedes Integral über $f(x)$ (durch Substitution) in ein Integral über $\varphi(x)$ umformen können.

Satz 6.2 *Es gilt*

$$\int_a^b f(t) dt = \int_{\frac{a-\mu}{\sigma}}^{\frac{b-\mu}{\sigma}} \varphi(x) dx = \Phi\left(\frac{b-\mu}{\sigma}\right) - \Phi\left(\frac{a-\mu}{\sigma}\right).$$

Beweis durch Substitution:

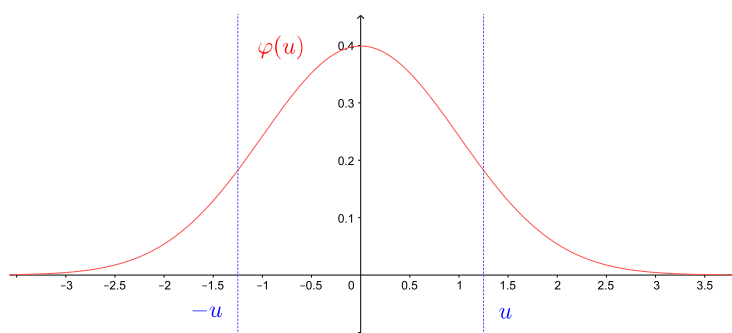
Eine weitere wichtige Eigenschaft der Standardglockenkurve ist, dass der Flächeninhalt der gesamten Fläche unter der Kurve gleich 1 ist.

Satz 6.3

$$\int_{-\infty}^{\infty} \varphi(x) dx = 1.$$

Es folgen sofort zwei weitere Eigenschaften:

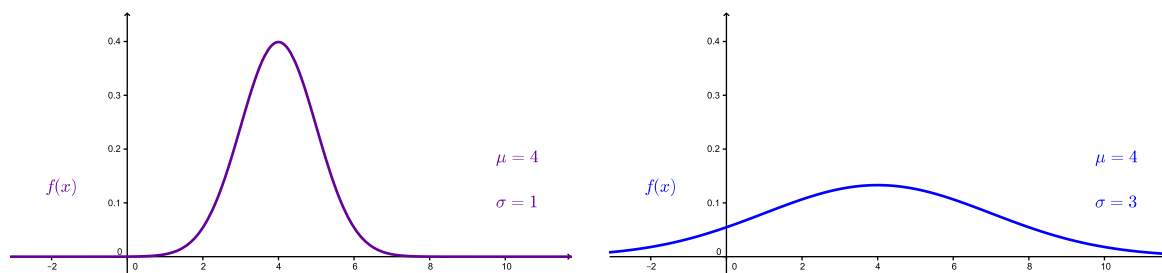
$$\begin{aligned}\Phi(0) &= \frac{1}{2} \\ \Phi(-u) &= 1 - \Phi(u)\end{aligned}$$



Wegen Satz 6.2 ist der Flächeninhalt nicht nur unter der Standardglockenkurve sondern unter jeder beliebigen Glockenkurve gleich 1,

$$\int_{-\infty}^{\infty} f(x) dx = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} dx = 1.$$

Abhängig von der Grösse von σ ist die Glockenkurve hoch und schmal oder tief und breit.



6.2 Approximation der Binomialverteilung

Im Beispiel auf den Seiten 54–55 haben wir gesehen, dass die Wahrscheinlichkeiten $P_{50}(k)$ der dort betrachteten Binomialverteilung durch die Werte der Funktion f approximiert werden können. Allgemein gilt der folgende Satz.

Satz 6.4 (Lokaler Grenzwertsatz von de Moivre und Laplace)

Die Wahrscheinlichkeit $P_n(k)$ einer Binomialverteilung (mit der Erfolgswahrscheinlichkeit p im Einzelerperiment) kann approximiert werden durch

$$P_n(k) \approx f(k) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{k-\mu}{\sigma}\right)^2},$$

wobei $\mu = np$ und $\sigma = \sqrt{npq}$ mit $q = 1 - p$.

Diese Näherung ist (in den meisten Fällen) ausreichend genau, falls $\sigma^2 = npq > 9$.

In demselben Beispiel haben wir gesehen, dass die Wahrscheinlichkeit $P_{50}(26 \leq k \leq 30)$ durch ein Integral über f approximiert werden kann. Schauen wir die blauen Rechtecke auf Seite 55 genau an, dann sehen wir, dass wir als Integrationsgrenzen nicht 26 und 30 wählen müssen, sondern 25,5 und 30,5. Die Breite des ersten blauen Rechtecks liegt auf der x -Achse zwischen 25,5 und 26,5. Addieren wir zu 25,5 die fünf Rechtecksbreiten (je der Länge 1), dann endet die Breite des letzten blauen Rechtecks bei $25,5 + 5 = 30,5$. Wir erhalten damit die Näherung

$$P_{50}(26 \leq k \leq 30) \approx \int_{25,5}^{30,5} f(t) dt.$$

Mit Hilfe von Satz 6.2 können wir nun das Integral auf der rechten Seite problemlos berechnen.

Satz 6.5 Mit denselben Bezeichnungen wie in Satz 6.4 gilt die Näherung

$$P_n(a \leq k \leq b) \approx \int_{a-\frac{1}{2}}^{b+\frac{1}{2}} f(t) dt = \Phi\left(\frac{b+\frac{1}{2}-\mu}{\sigma}\right) - \Phi\left(\frac{a-\frac{1}{2}-\mu}{\sigma}\right).$$

Weiter gilt

$$P_n(k \leq b) \approx \Phi\left(\frac{b+\frac{1}{2}-\mu}{\sigma}\right).$$

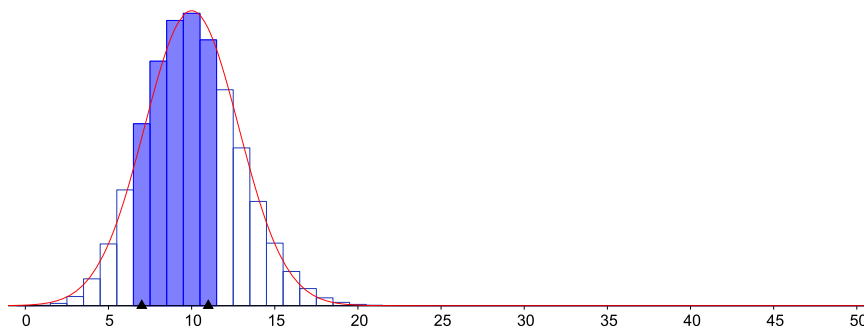
Beispiel

Wie gross ist die Wahrscheinlichkeit $P_{50}(26 \leq k \leq 30)$ vom Beispiel vorher?

Die Näherung von Satz 6.5 ist auch dann gut, wenn die Binomialverteilung nicht symmetrisch um den Erwartungswert ist.

Beispiele

1. Sei $n = 50$, $p = 0,2$. Dann ist $\mu = 10$ und $\sigma = 2\sqrt{2} \approx 2,83$. Mit Hilfe von Geogebra erhält man $P_{50}(7 \leq k \leq 11) = 0,6073$. Die Näherung von Satz 6.5 liefert $P_{50}(7 \leq k \leq 11) \approx 0,5945$.



2. In Mitteleuropa besitzen 45 % der Menschen die Blutgruppe A. Wie gross ist die Wahrscheinlichkeit, unter 100 zufälligen Blutspendern höchstens 40 mit dieser Blutgruppe vorzufinden?

6.3 Normalverteilte Zufallsgrössen

Zu Beginn dieses Kapitels haben wir ein Beispiel einer sogenannten stetigen Zufallsgrösse gesehen. Im Gegensatz zu einer diskreten Zufallsgrösse nimmt eine stetige Zufallsgrösse (nicht abzählbar) unendlich viele reelle Werte an, das heisst, die Werte eines ganzen Intervalles.

Genauer heisst eine Zufallsgrösse X *stetig*, wenn es eine integrierbare Funktion $\delta(x)$ gibt, so dass die Wahrscheinlichkeit $P(X \leq x)$ gegeben ist durch

$$P(X \leq x) = \int_{-\infty}^x \delta(t) dt .$$

Die Funktion $\delta(x)$ heisst *Dichtefunktion* von X und erfüllt die Eigenschaften

$$\int_{-\infty}^{\infty} \delta(x) dx = 1 \quad \text{und} \quad \delta(x) \geq 0 \quad \text{für alle } x \in \mathbb{R} .$$

Die Funktion $F(x) = P(X \leq x)$ nennt man *Verteilungsfunktion* von X . Die Wahrscheinlichkeit $P(X \leq x)$ einer stetigen Zufallsgrösse X entspricht also dem Flächeninhalt der Fläche zwischen dem Graphen von δ und der x -Achse zwischen $-\infty$ und x . Es gilt deshalb, dass

$$P(X \leq x) = P(X < x) \quad \text{und} \quad P(X \geq x) = P(X > x).$$

Insbesondere folgt (da $P(X = x) = P(X \leq x) - P(X < x)$)

$$P(X = x) = 0.$$

Die Funktion $\delta(x) = f(x)$ von den Abschnitten vorher ist die wichtigste Dichtefunktion.

Definition Eine Zufallsgrösse X heisst *normalverteilt* mit den Parametern μ und σ , wenn sie die Dichtefunktion

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

besitzt. Die zugehörige Wahrscheinlichkeitsverteilung heisst *Normalverteilung* oder auch *Gauß-Verteilung*. Die Parameter μ und σ sind der Erwartungswert, bzw. die Standardabweichung der Verteilung.

Wir haben in Abschnitt 6.1 gesehen, dass der Spezialfall $\mu = 0$ und $\sigma = 1$ eine wichtige Rolle spielt.

Definition Eine Zufallsgrösse Z heisst *standardnormalverteilt*, wenn sie normalverteilt mit den Parametern $\mu = 0$ und $\sigma = 1$ ist. Ihre Dichtefunktion ist damit

$$\varphi(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2}.$$

Insbesondere gilt

$$P(Z \leq x) = \int_{-\infty}^x \varphi(t) dt = \Phi(x).$$

Mit Hilfe von Satz 6.2 können auch die Wahrscheinlichkeiten einer beliebigen normalverteilten Zufallsgrösse (d.h. mit beliebigen Parametern μ und σ) berechnet werden.

Satz 6.6 Sei X eine normalverteilte Zufallsgrösse mit den Parametern μ und σ . Dann gilt

$$P(X \leq x) = \int_{-\infty}^x f(t) dt = \Phi\left(\frac{x-\mu}{\sigma}\right).$$

Damit folgt

$$P(a \leq X \leq b) = \Phi\left(\frac{b-\mu}{\sigma}\right) - \Phi\left(\frac{a-\mu}{\sigma}\right).$$

Man kann eine Zufallsgrösse X wie in Satz 6.6 auch direkt *standardisieren*; standardnormalverteilt ist die Zufallsgrösse

$$Z = \frac{X - \mu}{\sigma}.$$

Beispiele

1. Gegeben sind normalverteilte Messwerte (d.h. die Zufallsgrösse $X = (\text{Messwert})$ ist normalverteilt) mit dem Erwartungswert $\mu = 4$ und der Standardabweichung $\sigma = 2$. Wie gross ist die Wahrscheinlichkeit, dass ein Messwert (a) höchstens 6 ist (b) mindestens 2 ist und (c) zwischen 3,8 und 7 liegt?

(a)

(b)

(c)

$$\begin{aligned} P(3,8 \leq X \leq 7) &= \Phi\left(\frac{7-4}{2}\right) - \Phi\left(\frac{3,8-4}{2}\right) = \Phi(1,5) - \Phi(-0,1) \\ &= \Phi(1,5) - (1 - \Phi(0,1)) = 0,473 \end{aligned}$$

2. Das Gewicht von gewissen automatisch gepressten Tabletten ist erfahrungsgemäss normalverteilt mit $\mu = 25$ mg und $\sigma = 0,7$ mg.

(a) Mit welcher Wahrscheinlichkeit ist das Gewicht einer einzelnen Tablette zwischen 23,8 mg und 26,2 mg (d.h. im Bereich $\mu \pm 1,2$ mg)?

(b) Mit welcher Wahrscheinlichkeit ist das Gewicht von allen 30 Tabletten einer Packung zwischen 23,8 mg und 26,2 mg?

3. Schokoladentafeln werden abgefüllt. Das Abfüllgewicht ist erfahrungsgemäss normalverteilt mit $\mu = 100$ Gramm und $\sigma = 5$ Gramm. Man bestimme den Toleranzbereich $\mu \pm c\sigma$ so, dass 90 % aller Abfüllgewichte in diesen Bereich fallen.

Zwischen 91,8 und 108,2 Gramm liegen also 90 % aller Abfüllgewichte.

Wahrscheinlichkeiten unabhängig von den Werten von μ und σ

Im vorhergehenden Beispiel haben wir festgestellt, dass $c = 1,645$ unabhängig von μ und σ ist. Man kann nun analog für beliebige μ und σ zu einer vorgegebenen Wahrscheinlichkeit den zugehörigen, um μ symmetrischen Bereich angeben:

Wahrscheinlichkeit in %	50 %	90 %	95 %	99 %
Bereich	$\mu \pm 0,675 \sigma$	$\mu \pm 1,645 \sigma$	$\mu \pm 1,96 \sigma$	$\mu \pm 2,576 \sigma$

Diese Tabelle liest sich so: Der um μ symmetrische Bereich, in den eine normalverteilte Zufallsgrösse mit Erwartungswert μ und Varianz σ^2 beispielsweise mit einer Wahrscheinlichkeit von 95 % fällt, ist $\mu \pm 1,96 \sigma$.

Wir können auch umgekehrt fragen: Mit welcher Wahrscheinlichkeit liegt eine normalverteilte Zufallsgrösse beispielsweise im Bereich $\mu \pm \sigma$?

Analog finden wir die folgenden Werte.

Bereich	$\mu \pm \sigma$	$\mu \pm 2\sigma$	$\mu \pm 3\sigma$	$\mu \pm 4\sigma$
Wahrscheinlichkeit in %	68,26 %	95,45 %	99,73 %	≈ 100 %

Diese Tabelle liest sich nun so: Die Wahrscheinlichkeit, dass eine normalverteilte Zufallsgrösse mit Erwartungswert μ und Varianz σ^2 beispielsweise im Bereich $\mu \pm 2\sigma$ liegt, beträgt 95,45 %.

6.4 Der zentrale Grenzwertsatz

Sind X und Y zwei unabhängige und normalverteilte Zufallsgrössen, dann ist auch die Summe $X + Y$ eine normalverteilte Zufallsgrösse. Der zentrale Grenzwertsatz verallgemeinert diese Aussage.

Satz 6.7 (Zentraler Grenzwertsatz) *Seien $X_1, \dots, X_n$ unabhängige und identisch verteilte Zufallsgrössen (sie brauchen nicht normalverteilt zu sein). Ihr Erwartungswert sei jeweils μ und die Varianz σ^2 . Dann hat die Summe $S_n = X_1 + \dots + X_n$ den Erwartungswert $n\mu$ und die Varianz $n\sigma^2$.*

Für die zugehörige standardisierte Zufallsgrösse

$$Z_n = \frac{S_n - n\mu}{\sqrt{n}\sigma} = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}}$$

gilt

$$\lim_{n \rightarrow \infty} P(Z_n \leq x) = \Phi(x).$$

In Worten bedeutet dies (grob): Ist ein Merkmal (d.h. Zufallsgrösse) eine Summe von vielen (kleinen) zufälligen, unabhängigen Einflüssen, so können die Wahrscheinlichkeiten dieses Merkmals näherungsweise durch eine Normalverteilung beschrieben werden. Ein solches Merkmal ist zum Beispiel der Messfehler bei einer Messung, die Füllmenge von automatisch abgefüllten Flaschen oder der Intelligenzquotient eines Menschen. Es gibt zahlreiche weitere Beispiele. Die Normalverteilung ist deshalb eine äusserst wichtige Verteilung der Statistik.

Der zentrale Grenzwertsatz erklärt schliesslich auch die gute Näherung der Normalverteilung an eine Binomialverteilung für grosse n (Satz 6.4).

Beispiel

Wir werfen eine Münze n -mal. Wir definieren die Zufallsgrössen X_i durch $X_i = 1$, falls beim i -ten Wurf Zahl eintritt und $X_i = 0$ sonst (d.h. bei Kopf). Die X_i sind damit unabhängig und identisch verteilt mit $\mu = E(X_i) = \frac{1}{2}$ und $\sigma^2 = Var(X_i) = \frac{1}{4}$. Dann ist

$$S_n = X_1 + \dots + X_n$$

die Anzahl Zahl bei n Würfeln und entspricht für $n = 50$ genau der Zufallsgrösse X des Beispiels auf Seite 54. Wir haben dort schon bemerkt, dass die Verteilung von $X = S_{50}$ durch eine Normalverteilung angenähert werden kann. Gemäss zentralem Grenzwertsatz können die Wahrscheinlichkeiten $P(Z_n \leq x)$ der zugehörigen standardisierten Zufallsgrösse

$$Z_n = \frac{S_n - n\mu}{\sqrt{n}\sigma} = \frac{S_n - \frac{n}{2}}{\sqrt{n}\frac{1}{2}}$$

für sehr grosse n näherungsweise durch $\Phi(x)$ berechnet werden.

7 Statistische Testverfahren

In diesem Kapitel geht es darum, eine Annahme (Hypothese) über eine Grundgesamtheit aufgrund einer Stichprobe entweder beizubehalten oder zu verwerfen.

7.1 Testen von Hypothesen

Wie kann man beispielsweise testen, ob ein neues Medikament wirklich wirkt oder ob die kranken Personen nicht einfach von selbst wieder gesund werden? Oder eine Lady behauptet, sie könne am Geschmack des Tees erkennen, ob zuerst die Milch oder zuerst der Tee in die Tasse gegossen wurde. Kann sie das wirklich oder blufft (bzw. rät) sie nur?

Beispiel eines einseitigen Tests

Betrachten wir das Beispiel mit dem Medikament genauer. Wir gehen von einer Krankheit aus, bei welcher 70 % der kranken Personen ohne Medikament von selbst wieder gesund werden. Ein neues Medikament gegen diese Krankheit wurde hergestellt und wird nun an $n = 10$ Personen getestet.

Wir gehen von einer sogenannten *Nullhypothese* H_0 aus.

Nullhypothese H_0 : Das Medikament nützt nichts.

Wir nehmen weiter an, dass das Medikament nicht schadet, also im besten Fall nützt oder sonst keine Wirkung hat. Dies bedeutet, dass der Test *einseitig* ist.

Die 10 Testpersonen sind also krank und nehmen das Medikament ein. Wieviele dieser Testpersonen müssen gesund werden, damit wir mit gewisser Sicherheit sagen können, dass das Medikament wirklich nützt und wir H_0 verwerfen können?

Vor der Durchführung des Experiments wählen wir eine *kritische Zahl* m von Genesenden und studieren das Ereignis $A = (m \text{ oder mehr Testpersonen werden von selbst gesund})$. Wie gross ist die Wahrscheinlichkeit $P(A)$? Hier haben wir eine Binomialverteilung mit $n = 10$ und $p = P(\text{eine Testperson wird von selbst gesund}) = 0,7$.

Für $m = 9$ zum Beispiel erhalten wir $P(A) = P_{10}(k \geq 9) = 0,1493$, was etwa 14,9 % entspricht. Wenn also 9 oder 10 Testpersonen gesund werden und wir deshalb die Nullhypothese H_0 verwerfen, ist die *Irrtumswahrscheinlichkeit* (also die Wahrscheinlichkeit, dass wir fälschlicherweise die Nullhypothese verwerfen) gleich 14,9 %. Das ist zuviel.

Wir erhöhen deshalb die kritische Zahl auf $m = 10$. Wenn alle 10 Testpersonen gesund werden, beträgt nun die Irrtumswahrscheinlichkeit beim Verwerfen von H_0 nur noch $P(A) = P_{10}(10) = 0,7^{10} = 0,0285 = 2,85 \%$. In allen anderen Fällen (d.h. wenn 9 oder weniger Personen gesund werden), müssen wir allerdings H_0 beibehalten.

Um mehr "Spielraum" zu haben, erhöhen wir die Anzahl der Testpersonen auf $n = 20$. Die Wahrscheinlichkeit des Ereignisses A ist nun gegeben durch

$$P(A) = P_{20}(k \geq m) = \sum_{k=m}^{20} \binom{20}{k} 0,7^k 0,3^{20-k}.$$

Für $m = 18$ zum Beispiel erhalten wir $P(A) = 0,0355 = 3,55\%$. Wir können also die Nullhypothese H_0 mit einer Irrtumswahrscheinlichkeit von etwa $3,6\%$ verwerfen, falls 18 oder mehr Personen gesund werden. In allen anderen Fällen müssen wir H_0 beibehalten, wobei auch das ein Fehler sein kann.

Fehlerarten

Bei einem Testverfahren kann man sich auf zwei verschiedene Arten irren.

Fehler erster Art. Die Nullhypothese H_0 ist richtig, das heisst, das Medikament ist tatsächlich wirkungslos. Doch wegen eines zufällig guten Ergebnisses verwerfen wir die Nullhypothese. Dies wird als *Fehler erster Art* bezeichnet.

Im Beispiel vorher mit $n = 20$ Testpersonen und $m = 18$ tritt ein Fehler erster Art mit einer Wahrscheinlichkeit von $\alpha = 3,6\%$ auf.

Fehler zweiter Art. Die Nullhypothese H_0 ist falsch, das heisst, das Medikament wirkt. Doch wegen eines zufällig schlechten Ergebnisses behalten wir die Nullhypothese bei. Die Wahrscheinlichkeit eines *Fehlers zweiter Art* bezeichnet man mit β .

Es gibt also vier Möglichkeiten, wie die Realität und die Testentscheidung zusammentreffen können:

		Realität	
		H_0 ist richtig	H_0 ist falsch
Testent- scheidung	H_0 beibehalten	ok	Fehler 2. Art, β -Fehler
	H_0 verwerfen	Fehler 1. Art, α -Fehler	ok

Die Wahrscheinlichkeit für einen Fehler erster Art wird zu Beginn des Tests durch Vorgabe von α nach oben beschränkt. In den Naturwissenschaften üblich ist eine Toleranz von bis zu $\alpha = 5\%$. Man spricht vom *Signifikanzniveau* α . Dieser Fehler ist also kontrollierbar. Gleichzeitig sollte jedoch die Wahrscheinlichkeit eines Fehlers zweiter Art nicht zu gross sein. Dieser Fehler kann allerdings nicht vorgegeben werden.

Beim Testen wählt man also den Fehler mit dem grösseren Risiko zum Fehler erster Art. Man wählt dementsprechend die Nullhypothese H_0 so, dass das irrtümliche Festhalten an H_0 nicht so schlimm ist, bzw. weniger schlimm als das irrtümliche Verwerfen von H_0 und Annehmen der *Alternativhypothese* H_1 ist.

Beim Testen eines neuen Medikaments ist es also besser, dieses als unwirksam anzunehmen (obwohl es wirkt) und weiterhin das bisher übliche Medikament zu verwenden, anstatt das neue Medikament gegen die Krankheit einzusetzen, obwohl es nichts nützt.

Beispiel eines zweiseitigen Tests

Wir wollen testen, ob eine Münze gefälscht ist. Wir gehen von der Nullhypothese H_0 aus, dass dem nicht so ist.

$$H_0: p(\text{Kopf}) = \frac{1}{2}$$

Die Alternativhypothese H_1 ist in diesem Fall, dass $p(\text{Kopf}) \neq \frac{1}{2}$.

$$H_1: p(\text{Kopf}) > \frac{1}{2} \text{ oder } p(\text{Kopf}) < \frac{1}{2}$$

Wir müssen daher *zweiseitig* testen.

Wir werfen die Münze zum Beispiel $n = 10$ Mal. Als Verwerfungsbereich der Anzahl Köpfe für H_0 wählen wir die Menge $\{0, 1, 2\} \cup \{8, 9, 10\} = \{0, 1, 2, 8, 9, 10\}$. Ist also die Anzahl der Köpfe bei 10 Würfeln sehr klein (nämlich 0, 1 oder 2) oder sehr gross (nämlich 8, 9 oder 10), dann verwerfen wir die Nullhypothese H_0 .

Wie gross ist damit der Fehler erster Art?

Wir erhalten also $\alpha = 10,9\%$. Dieser Fehler ist zu gross.

Wir werfen nochmals $n = 20$ Mal. Als Verwerfungsbereich für H_0 wählen wir nun die Menge $\{0, 1, 2, 3, 4, 16, 17, 18, 19, 20\}$. Für den Fehler erster Art erhalten wir nun

Das heisst, $\alpha = 1,2\%$. Tritt also bei den 20 Würfeln eine Anzahl Köpfe des Verwerfungsbereichs auf, dann verwerfen wir die Nullhypothese und nehmen an, dass die Münze gefälscht ist. Dabei irren wir uns mit einer Wahrscheinlichkeit von $\alpha = 1,2\%$.

Nun geben wir das Signifikanzniveau 5% vor. Wir werfen die Münze wieder 20-mal. Wie müssen wir den Verwerfungsbereich wählen, damit $\alpha \leq 5\%$? Dabei soll der Verwerfungsbereich so gross wie möglich sein. Wir suchen also die grösste natürliche Zahl x , so dass

Hier lohnt es sich, in der Tabelle der summierten Binomialverteilung nachzusehen. Für alle $x \leq 5$ gilt $P_{20}(k \leq x) \leq 0,025$. Mit $x = 5$ erhalten wir den grössten Verwerfungsbereich,

$$\{0, 1, 2, 3, 4, 5, 15, 16, 17, 18, 19, 20\}.$$

Wie gross ist nun α tatsächlich? Wir finden

$$\alpha = 2 P_{20}(k \leq 5) = 2 \cdot 0,021 = 0,042 \quad \implies \quad \alpha = 4,2\%$$

Werfen wir die Münze 100-mal, dann verwenden wir die Normalverteilung als Näherung für die Binomialverteilung. Wir wollen auch hier den grösstmöglichen Verwerfungsbereich bestimmen, so dass $\alpha \leq 5\%$. Wir haben also $n = 100$ und $p = \frac{1}{2}$ (die Nullhypothese) wie bisher. Die Parameter für die Normalverteilung sind damit

Wegen $\sigma^2 = 25 > 9$ können wir die Normalverteilung als Näherung für die Binomialverteilung nutzen.

Es muss nun gelten:

Näherung mit Normalverteilung X :

Mit der Tabelle von Seite 62 folgt

Damit $\alpha \leq 0,05$, müssen wir x auf 39 abrunden (den Verwerfungsbereich verkleinern bedeutet α verkleinern). Der grösstmögliche Verwerfungsbereich mit $\alpha \leq 5\%$ ist also

Wie gross ist nun α tatsächlich?

Weiteres Beispiel

Ein Hersteller von Überraschungseiern behauptet, dass in mindestens 14 % der Eier Figuren von Disney-Filmen stecken. Wir wollen dies testen und nehmen eine Stichprobe von $n = 1000$ Eiern, in welchen wir 130 Disney-Figuren finden. Genügt dieses Ergebnis, um die Behauptung des Herstellers auf dem Signifikanzniveau von 5 % zu widerlegen?

Sei p der wahre aber unbekannte Anteil der Eier mit Disney-Figuren. Der Hersteller behauptet, dass $p \geq 0,14$. Dies ist unsere Nullhypothese H_0 .

$$H_0: p \geq 0,14$$

Die Alternativhypothese H_1 ist, dass es in weniger als 14 % der Eier Disney-Figuren gibt.

$$H_1: p < 0,14$$

Wir haben also einen einseitigen Test. Der Fehler erster Art α soll höchstens 5 % betragen.

Die Zufallsgrösse X sei die Anzahl der Disney-Figuren in der Stichprobe. Da der Umfang $n = 1000$ der Stichprobe sehr gross ist, dürfen wir von einer Binomialverteilung ausgehen (die 1000 Ziehungen sind praktisch unabhängig), die wir durch eine Normalverteilung annähern.

Für den Verwerfungsbereich $\{0, 1, 2, \dots, x\}$ suchen wir also die grösste Zahl x , so dass

Die approximierende Normalverteilung X hat die Parameter

Es gilt also

und wir müssen x bestimmen, so dass

Mit der Tabelle von Seite 62 folgt

Als Verwerfungsbereich erhalten wir also $\{0, 1, \dots, 121\}$ und da wir 130 Disney-Figuren gefunden haben, können wir die Behauptung des Herstellers nicht mit der gewünschten Sicherheit verwerfen.

Alternativ könnten wir auch vom Verwerfungsbereich $\{0, 1, \dots, 130\}$, dem kleinsten Verwerfungsbereich, der 130 enthält, ausgehen und direkt den Fehler α berechnen:

$$\alpha = P_{1000}(k \leq 130) \approx \Phi\left(\frac{130 + \frac{1}{2} - 140}{10,97}\right) = \Phi(-0,87) = 1 - \Phi(0,87) = 0,19215$$

Das heisst, wenn wir aufgrund obiger Stichprobe dem Hersteller widersprechen, irren wir uns mit einer Wahrscheinlichkeit von $\alpha = 19,2\%$, was zuviel ist.

7.2 Der t -Test für Mittelwerte

In diesem Abschnitt interessieren wir uns für den Mittelwert μ eines Merkmals (d.h. den Erwartungswert μ einer Zufallsgrösse) einer (oder zweier) Grundgesamtheit(en). Anhand einer Stichprobe wollen wir Aussagen über den unbekannten Mittelwert μ machen.

Vertrauensintervall

Gegeben ist also ein Merkmal einer Grundgesamtheit, wobei wir annehmen, dass dieses Merkmal normalverteilt ist. Sowohl der Mittelwert μ als auch die Varianz σ^2 dieses Merkmals sind unbekannt. Wir entnehmen dieser Grundgesamtheit eine Stichprobe und berechnen in Abhängigkeit dieser Stichprobe ein sogenanntes 95 %-Vertrauensintervall für μ . Dies bedeutet, dass der unbekannte Mittelwert μ mit einer Wahrscheinlichkeit von 95 % in diesem Vertrauensintervall liegt.

Wir benötigen dazu den sogenannten Standardfehler. In Kapitel 1 hatten wir die Standardabweichung

$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$$

für eine Messreihe $x_1, \dots, x_n$ definiert. Wir bezeichnen diese nun mit $s = SD$ (für standard deviation). Weiter haben wir in Abschnitt 4.4 gesehen, dass die Varianz des Mittelwerts $\bar{X}$ einer Stichprobe gegeben ist durch σ^2/n , wobei σ^2 die Varianz des Merkmals der Grundgesamtheit ist. Da diese jedoch unbekannt ist, schätzen wir sie durch s^2 (wie in Abschnitt 4.4). Dies führt zum *Standardfehler* (standard error) der Stichprobe

$$SE = \frac{SD}{\sqrt{n}} = \frac{s}{\sqrt{n}} = \sqrt{\frac{1}{n(n-1)} \sum_{i=1}^n (x_i - \bar{x})^2}.$$

Der Standardfehler ist also umso kleiner, je grösser der Umfang der Stichprobe ist.

Beispiel

Wie in Abschnitt 4.4 seien alle Studierenden der Vorlesung Mathematik II die Grundgesamtheit und das Merkmal sei das Alter (d.h. die Zufallsgrösse X ordnet jedem*r Studierenden sein*ihr Alter zu). Wir können davon ausgehen, dass dieses Merkmal normalverteilt ist. Das Durchschnittsalter der Studierenden, das heisst der Mittelwert μ , ist unbekannt, ebenso die Varianz σ^2 . Das Ziel ist, ein Intervall anzugeben, in welchem der Mittelwert μ mit 95 %-iger Wahrscheinlichkeit liegt.

Dazu nehmen wir eine Stichprobe. Wir notieren also das Alter von beispielsweise 8 zufällig ausgewählten Studierenden. Wir erhalten (zum Beispiel) die Zahlen

20	22	19	20	21	23	21	24
----	----	----	----	----	----	----	----

Wir berechnen

$$\bar{x} = 21,25, \quad s = SD = 1,669, \quad SE = \frac{s}{\sqrt{8}} = 0,590.$$

Wir wissen (von Abschnitt 4.4), dass der Erwartungswert des Stichprobenmittelwerts $\bar{X}$ gleich dem Mittelwert μ ist. Wäre nun der Stichprobenumfang sehr gross (etwa $n \geq 30$) und

die Varianz σ^2 bekannt, dann wäre (nach dem zentralen Grenzwertsatz) die standardisierte Zufallsgrösse

$$Z = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}}$$

(näherungsweise) standardnormalverteilt. Die Zufallsgrösse Z würde gemäss der Tabelle auf Seite 62 mit einer Wahrscheinlichkeit von 95 % einen Wert zwischen $-1,96$ und $1,96$ annehmen.

Nun haben wir jedoch einen kleinen Stichprobenumfang und die Varianz σ^2 ist unbekannt, so dass wir die obige Bemerkung in zwei Punkten korrigieren müssen. Erstens ersetzen wir (wie schon weiter oben bemerkt) $\sigma/\sqrt{n}$ durch den Standardfehler $SE = s/\sqrt{n}$. Zweitens ist nun die “standardisierte” Zufallsgrösse

$$Z = \frac{\bar{X} - \mu}{SE}$$

nicht standardnormalverteilt, sondern sie folgt der sogenannten *Studentschen t-Verteilung*. Diese hängt vom Stichprobenumfang n , bzw. vom *Freiheitsgrad*

$$\nu = n - 1$$

ab. Ist n gross, dann sieht die t -Verteilung wie die Normalverteilung aus; für kleine n ist die Kurve jedoch flacher und breiter. Die Studentsche t -Verteilung wurde von William Sealy Gosset eingeführt; der Name stammt von seinem Pseudonym “Student”, unter dem er publizierte.

Die Rolle der Zahl $1,96$ oben übernimmt nun der *kritische Schrankenwert* t_{krit} , den wir aus der Tabelle (Seite 13) ablesen können. In unserem Beispiel ist $\nu = 8 - 1 = 7$ und da wir eine Wahrscheinlichkeit von 95 % suchen, ist das Signifikanzniveau $\alpha = 5 \% = 0,05$. Wir finden den Tabellenwert

$$t_{\text{krit}} = 2,365.$$

Damit gilt

$$0,95 = P(-2,365 \leq Z \leq 2,365) = P(-2,365 \cdot SE \leq \bar{X} - \mu \leq 2,365 \cdot SE),$$

also liegt μ mit 95 %-iger Wahrscheinlichkeit im Intervall $[\bar{X} - 2,365 \cdot SE, \bar{X} + 2,365 \cdot SE]$. Setzen wir unseren konkreten Stichprobenmittelwert $\bar{x} = 21,25$ sowie den Standardfehler $SE = 0,590$ ein, erhalten wir das Vertrauensintervall

$$[21,25 - 2,365 \cdot 0,590; 21,25 + 2,365 \cdot 0,590] = [19,854; 22,646].$$

Allgemein ist das 95 %-Vertrauensintervall, das den unbekannten Mittelwert μ mit einer Wahrscheinlichkeit von 95 % überdeckt, gegeben durch

$$[\bar{x} - t_{\text{krit}} \cdot SE, \bar{x} + t_{\text{krit}} \cdot SE].$$

Das Vertrauensintervall ist vom Mittelwert $\bar{x}$ der Stichprobe, und damit von der Stichprobe abhängig. Eine andere Stichprobe ergibt möglicherweise ein anderes Vertrauensintervall. Insbesondere verkleinert ein grösserer Stichprobenumfang die Länge des Intervalles wesentlich.

Anstelle der Berechnung eines Vertrauensintervalles könnte man auch testen, ob eine bestimmte Zahl μ_0 als Mittelwert μ wahrscheinlich ist. Zum Beispiel testen wir wie folgt:

$$\begin{array}{ll} \text{Nullhypothese:} & \mu = \mu_0 = 23 \\ \text{Alternativhypothese:} & \mu \neq \mu_0 = 23 \\ \text{Signifikanzniveau:} & \alpha = 5\% \end{array}$$

Dies ist also ein zweiseitiger Test.

Als Testgrösse verwenden wir

$$t = \frac{|\bar{x} - \mu_0|}{SE}.$$

Damit liegt μ_0 im Vertrauensintervall, genau dann wenn $t \leq t_{\text{krit}}$. Es gilt also:

$$t > t_{\text{krit}} \implies \text{Nullhypothese verwerfen}$$

Mit unseren Messwerten erhalten wir

Unsere Stichprobe steht also auf dem 5%-Niveau im Widerspruch zur Nullhypothese, das heisst, wir können davon ausgehen, dass $\mu_0 = 23$ nicht der Mittelwert μ ist.

Allgemeines Vorgehen

Gesucht: Mittelwert μ eines normalverteilten Merkmals einer Grundgesamtheit.

Gegeben: Stichprobe vom Umfang n mit Mittelwert $\bar{x}$ und Standardfehler SE .

Vertrauensintervall zum Niveau $1 - \alpha$:

$$[\bar{x} - t_{\alpha,\nu} \cdot SE, \bar{x} + t_{\alpha,\nu} \cdot SE],$$

wobei $t_{\alpha,\nu} = t_{\text{krit}}$ der kritische Schrankenwert (gemäss Tabelle der Studentschen t -Verteilung) für das Signifikanzniveau α und den Freiheitsgrad $\nu = n - 1$ ist.

Oder mit Testgrösse:

$$t = \frac{|\bar{x} - \mu_0|}{SE}$$

Entscheid: $t > t_{\alpha,\nu} \implies \mu \neq \mu_0$

Vergleich der Mittelwerte zweier Normalverteilungen

Gegeben sind zwei Grundgesamtheiten mit je einem normalverteilten Merkmal, dessen Mittelwert μ_x , bzw. μ_y ist. Wir wollen testen, ob die beiden Mittelwerte μ_x und μ_y gleich sind. Die Varianzen brauchen nicht bekannt zu sein, sie werden aber als gleich vorausgesetzt.

Für den Test brauchen wir je eine Stichprobe aus den beiden Grundgesamtheiten. Die beiden folgenden Fälle sind praktisch besonders wichtig:

1. Die beiden Stichproben sind gleich gross. Je ein Wert der einen und ein Wert der anderen Stichprobe gehören zusammen, da sie von demselben Merkmalsträger stammen (zum Beispiel das Körpergewicht vor und nach einer Diät oder Messwerte von demselben Objekt, gemessen mit zwei verschiedenen Messgeräten). Man spricht von *gepaarten Stichproben*.
2. Die beiden Stichproben sind unabhängig und nicht notwendigerweise gleich gross.

1. Gepaarte Stichproben

12 Personen machen eine Diät. Verringert diese Diät das Körpergewicht auch wirklich? Bei den Proband*innen wird deshalb das Körpergewicht (in kg) vor und nach der Diät gemessen:

Proband i	Gewicht vorher x_i	Gewicht nachher y_i	Differenz $d_i = x_i - y_i$
1	84,5	83	1,5
2	72,5	72,5	0
3	69	64,5	4,5
4	88,5	89,5	-1
5	104,5	94	10,5
6	63	57,5	5,5
7	93,5	95,5	-2
8	67	60	7
9	76,5	75	1,5
10	58,5	54,5	4
11	79,5	73,5	6
12	92	83,5	8,5
	$\bar{x} = 79,083$	$\bar{y} = 75,250$	$\bar{d} = 3,833$
	$s_x = 13,879$	$s_y = 14,133$	$s_d = 3,898$

Wir wollen nun testen, ob $\mu_x = \mu_y$. Dabei können wir auf den vorhergehenden Test für einen einzelnen Mittelwert zurückgreifen, indem wir wie folgt testen (zweiseitig):

Nullhypothese: $\mu_d = \mu_x - \mu_y = 0$

Alternativhypothese: $\mu_d \neq 0$

Signifikanzniveau: $\alpha = 5\%$

Wir berechnen also die Testgrösse

Da dieses t grösser als der kritische Schrankenwert $t_{\text{krit}} = t_{\alpha, \nu} = 2,201$ (gemäss Tabelle der t -Verteilung, Freiheitsgrad $\nu = 11$) ist, kann die Nullhypothese auf dem 5%-Niveau verworfen werden. Wir können davon ausgehen, dass die Diät tatsächlich wirkt und sich das Körpergewicht damit allgemein verringert.

2. Unabhängige Stichproben

Ein neues Düngemittel (N) für Sojabohnen wurde entwickelt und nun soll getestet werden, ob dieses das Wachstum der Sojabohnen besser als das bisher verwendete Düngemittel (B) fördert. Dabei wird davon ausgegangen, dass das neue Düngemittel N nicht schlechter als das Düngemittel B wirkt. Es werden 20 gleichartige Sojapflanzen zufällig ausgewählt, 12 davon mit dem Düngemittel N und die restlichen 8 mit dem Düngemittel B gedüngt. Nach einer bestimmten Zeit wird die Höhe der Pflanzen gemessen. Die durchschnittliche Höhe der $n_x = 12$ Pflanzen, die mit Mittel N gedüngt wurden, beträgt $\bar{x} = 35,6$ cm mit einer Standardabweichung von $s_x = 1,8$ cm und die durchschnittliche Höhe der $n_y = 8$ Pflanzen, die mit Mittel B gedüngt wurden, beträgt $\bar{y} = 33,2$ cm mit $s_y = 1,9$ cm.

Da wir nun keine Stichprobenpaare (x_i, y_i) mehr haben, müssen wir die vorhergehende Testmethode leicht anpassen. Wir testen wie folgt:

Nullhypothese: Die Düngemittel N und B wirken gleich gut ($\mu_x = \mu_y$).
 Alternativhypothese: Düngemittel N wirkt besser als Düngemittel B ($\mu_x > \mu_y$).
 Signifikanzniveau: $\alpha = 1\%$

Dies ist also ein einseitiger Test. (Man könnte auch zweiseitig testen, die Alternativhypothese wäre in diesem Fall $\mu_x \neq \mu_y$.)

Wir verwenden die Testgrösse

$$t = \frac{|\bar{x} - \bar{y}|}{SE_{\bar{x} - \bar{y}}}, \quad \text{wobei} \quad SE_{\bar{x} - \bar{y}} = \sqrt{\frac{n_x + n_y}{n_x n_y} \sqrt{\frac{s_x^2(n_x - 1) + s_y^2(n_y - 1)}{n_x + n_y - 2}}}$$

der Standardfehler der Differenz $\bar{x} - \bar{y}$ ist. Für die Testgrösse erhalten wir also

$$t = |\bar{x} - \bar{y}| \sqrt{\frac{n_x n_y}{n_x + n_y} \sqrt{\frac{n_x + n_y - 2}{s_x^2(n_x - 1) + s_y^2(n_y - 1)}}.$$

In unserem Beispiel erhalten wir

Nun benutzen wir wieder die Tabelle der t -Verteilung. Der Freiheitsgrad ist hier

$$\nu = n_x + n_y - 2,$$

also $\nu = 18$. Die Tabelle (für $\alpha = 0,01$, einseitiger Test) gibt uns den kritischen Schrankenwert $t_{\text{krit}} = 2,552$. Dieser ist kleiner als unser berechneter Wert $t = 2,858$. Wir können also die Nullhypothese bei einer zugelassenen Wahrscheinlichkeit von 1% für den Fehler erster Art verwerfen. Das Düngemittel N wirkt besser als das bisher verwendete Düngemittel B.

Allgemeines Vorgehen

Test: Gilt $\mu_x = \mu_y$ für die Mittelwerte μ_x, μ_y von zwei normalverteilten Merkmalen?

Gegeben: Je eine Stichprobe vom Umfang n_x, n_y mit Mittelwerten $\bar{x}, \bar{y}$ und Standardabweichungen s_x, s_y

α Signifikanzniveau, $\nu = n_x + n_y - 2$ Freiheitsgrad, $t_{\alpha, \nu}$ (gemäss Tabelle der t -Verteilung)

Testgrösse:

$$t = |\bar{x} - \bar{y}| \sqrt{\frac{n_x n_y}{n_x + n_y} \sqrt{\frac{n_x + n_y - 2}{s_x^2(n_x - 1) + s_y^2(n_y - 1)}}}$$

Entscheid: $t > t_{\alpha, \nu} \implies \mu_x \neq \mu_y$

Nicht normalverteilte Merkmale

Sind die Merkmale der Grundgesamtheiten nicht normalverteilt, so kann der t -Test als Näherung trotzdem verwendet werden (der Näherungsfehler ist umso kleiner, je grösser die Stichproben sind). Ansonsten kann der sogenannte *Wilcoxon-Mann-Whitney-Test*, der keine Annahmen über die Verteilungen der Merkmale der Grundgesamtheiten macht, verwendet werden. Auf diesen Test gehen wir hier aber nicht ein.

7.3 Der Varianzenquotiententest

Wenn wir den t -Test für zwei Stichproben anwenden wollen, müssen wir in den beiden Grundgesamtheiten dieselbe Varianz der Merkmale voraussetzen, das heisst $\sigma_x^2 = \sigma_y^2$. Wie können wir diese Bedingung überprüfen?

Beispiel

Wir vergleichen zwei verschiedene Pipettier-Methoden, um 50 ml abzumessen.

Die beiden Stichproben ergeben die folgenden Messwerte:

	automatische Pipette
1	48,82
2	50,88
3	51,22
4	49,75
5	50,19
6	50,01
7	49,98
8	48,29
9	49,82
10	51,02
	Mittelwert: 49,998
	Varianz: 0,8612

	manuelle Pipette
1	50,11
2	50,11
3	49,92
4	50,63
5	49,91
6	50,26
7	50,05
8	50,09
	Mittelwert: 50,135
	Varianz: 0,0526

Die Varianz bei der Stichprobe der automatischen Pipette ist grösser. Können wir daraus schliessen, dass allgemein die Varianz der Messwerte bei der automatischen Pipette grösser ist oder dass die Varianzen allgemein unterschiedlich sind bei den beiden Pipettier-Methoden?

Wir testen (zweiseitig) wie folgt:

Nullhypothese: $\sigma_x^2 = \sigma_y^2$

Alternativhypothese: $\sigma_x^2 \neq \sigma_y^2$

Signifikanzniveau: $\alpha = 5\%$

Wir verwenden hier die Testgrösse

$$F = \frac{s_x^2}{s_y^2} \quad \text{mit } s_x^2 \geq s_y^2.$$

Sind die Varianzen gleich, dann ist F nahe bei 1.

Damit in unserem Beispiel die Bedingung $s_x^2 \geq s_y^2$ erfüllt ist, müssen wir x für die automatische und y für die manuelle Pipettierung wählen. Für unsere Testgrösse erhalten wir also

Die Testgrösse F folgt der sogenannten F -Verteilung (nach Ronald Aylmer Fisher). Wir benutzen also die Tabelle der F -Verteilung. Dazu brauchen wir noch den Freiheitsgrad vom Zähler $\nu_x = 10 - 1 = 9$ und vom Nenner $\nu_y = 8 - 1 = 7$. Die Tabelle gibt den kritischen Schrankenwert

$$F_{\text{krit}} = 3,68.$$

Wie beim t -Test gilt nun allgemein

$$F > F_{\text{krit}} \implies \text{Nullhypothese verwerfen}$$

Da unsere berechnete Testgrösse $F = 16,37$ grösser als F_{krit} ist, können wir unsere Nullhypothese auf dem 5 %-Niveau verwerfen. Die Varianzen bei den beiden Pipettier-Methoden sind also nicht gleich.

Allgemeines Vorgehen

Test: Gilt $\sigma_x^2 = \sigma_y^2$ für die Varianzen σ_x^2, σ_y^2 von zwei normalverteilten Merkmalen?

Gegeben: Je eine Stichprobe mit den Varianzen s_x^2, s_y^2 .

α Signifikanzniveau, ν_x, ν_y Freiheitsgrade, F_{α, ν_x, ν_y} (gemäss Tabelle der F -Verteilung)

Testgrösse:

$$F = \frac{s_x^2}{s_y^2} \quad \text{mit } s_x^2 \geq s_y^2$$

$$\text{Entscheid: } F > F_{\alpha, \nu_x, \nu_y} \implies \sigma_x^2 \neq \sigma_y^2$$

7.4 Korrelationsanalyse

In Kapitel 2 haben wir die Korrelationskoeffizienten von Pearson und von Spearman kennengelernt. Hier wollen wir nun aus dem Korrelationskoeffizienten der Messwertpaare einer Stichprobe Aussagen über den Korrelationskoeffizienten der Messwertpaare der Grundgesamtheit machen.

Beispiel

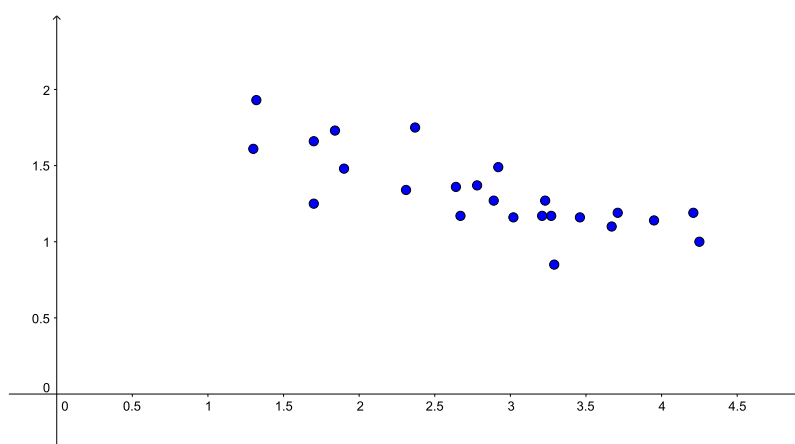
Blätter von Bäumen, aus denen ein bestimmter Wirkstoff gewonnen werden kann, sollen geerntet werden. Wir überlegen uns, ob der Wirkstoffgehalt in einem Blatt davon abhängt, wie hoch das Blatt am Baum hängt. Wenn nicht, könnte man einfach die leicht zugänglichen Blätter in niedriger Höhe ernten, ohne Leitern verwenden zu müssen.

Wir pflücken daher als Stichprobe 24 Blätter in unterschiedlicher Höhe und notieren ihren Wirkstoffgehalt. Wir erhalten die folgenden Messwertpaare (in m, bzw. mg/100 g):

Nr. i	Höhe x	Wirkstoffgehalt y
1	1,70	1,66
2	2,31	1,34
3	2,89	1,27
4	1,30	1,61
5	3,21	1,17
6	1,84	1,73
7	3,27	1,17
8	4,21	1,19
9	1,32	1,93
10	3,67	1,10
11	2,78	1,37
12	3,71	1,19

Nr. i	Höhe x	Wirkstoffgehalt y
13	3,23	1,27
14	3,29	0,85
15	3,46	1,16
16	3,95	1,14
17	1,70	1,25
18	2,92	1,49
19	2,67	1,17
20	3,02	1,16
21	2,37	1,75
22	2,64	1,36
23	4,25	1,00
24	1,90	1,48

Streudiagramm:



Der Korrelationskoeffizient von Pearson beträgt

$$r_{xy} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}} \approx -0,7767$$

Das Streudiagramm und der Korrelationskoeffizient weisen auf eine Korrelation der Wertepaare der Stichprobe hin. Aber gibt es auch eine Korrelation der Wertepaare der Grundgesamtheit (d.h. zwischen der Wuchshöhe und dem Wirkstoffgehalt von beliebigen Blättern an den Bäumen)? Wir bezeichnen mit ϱ den Korrelationskoeffizienten der Wertepaare der Grundgesamtheit.

Um den folgenden Test anwenden zu können, müssen beide Merkmale (d.h. die Zufallsgrößen x und y) der Grundgesamtheit normalverteilt sein. Man nennt eine solche Verteilung *bivariate Normalverteilung*. Im Gegensatz zum t -Test reagiert dieser Test empfindlich auf Abweichungen von der Normalverteilung.

Nun testen wir (zweiseitig) wie folgt:

$$\begin{array}{ll} \text{Nullhypothese:} & \varrho = 0 \\ \text{Alternativhypothese:} & \varrho \neq 0 \\ \text{Signifikanzniveau:} & \alpha = 5\% \end{array}$$

Die Testgröße ist der Betrag des Korrelationskoeffizienten der Messwertpaare der Stichprobe, das heisst $|r_{xy}| \approx 0,7767$. Den *kritischen Schrankenwert* r_{krit} entnehmen wir der Tabelle (Seite 20). Wir finden

$$r_{\text{krit}} = 0,404.$$

Es gilt allgemein

$$|r_{xy}| > r_{\text{krit}} \implies \text{Nullhypothese verwerfen}$$

Wir können in unserem Beispiel also die Nullhypothese auf dem 5 %-Niveau verwerfen. Es gibt eine Korrelation zwischen der Wuchshöhe und dem Wirkstoffgehalt eines Blattes. Allerdings

ist die Korrelation entgegengesetzt (da $r_{xy} < 0$), das heisst, je höher das Blatt sich befindet, desto geringer ist der Wirkstoffgehalt. Bei der Ernte bleiben wir also schön am Boden und pflücken die unteren Blätter.

Wenn wir nicht davon ausgehen können, dass die Merkmale (Zufallsgrössen) x und y der Grundgesamtheit bivariat normalverteilt sind, können wir für die Testgrösse den Rangkorrelationskoeffizienten von Spearman benutzen.

Zum Beispiel wollen wir testen, ob die an die Spieler des FC Basel vergebenen Noten nach einem Fussballspiel in den beiden Zeitungen bz Basel (bz) und Basler Zeitung (BaZ) korrelieren. Als Stichprobe untersuchen wir die Noten des Champions League Spiels Manchester City gegen den FCB vom 7. März 2018:

Spieler	Note bz	Note BaZ	Rang r_{bz}	Rang r_{BaZ}	$d = r_{bz} - r_{BaZ}$	d^2
T. Vaclík	5,5	5,4	2,5	2,5	0	0
M. Suchy	5	4,7	7,5	8	-0,5	0,25
F. Frei	5	5	7,5	4	3,5	12,25
L. Lacroix	5	4,4	7,5	10	-2,5	6,25
M. Lang	5	5,4	7,5	2,5	5	25
G. Serey Dié	5,5	4,9	2,5	5,5	-3	9
L. Zuffi	5	4,9	7,5	5,5	2	4
B. Riveros	5	4,7	7,5	8	-0,5	0,25
K. Bua	5	4,7	7,5	8	-0,5	0,25
D. Oberlin	4,5	3,6	12	11,5	0,5	0,25
M. Elyounoussi	6	5,6	1	1	0	0
V. Stocker	5	3,6	7,5	11,5	-4	16
					Summe	73,5

Für den Rangkorrelationskoeffizienten von Spearman erhalten wir also

$$r_S = 1 - \frac{6}{n(n^2 - 1)} \sum_{i=1}^n d_i^2 = 0,743.$$

Der *kritische Schrankenwert* $r_{S,krit}$ der Tabelle (Seite 21) beträgt

$$r_{S,krit} = 0,591.$$

Wir können also bei diesem Test die Nullhypothese, dass keine Korrelation besteht, verwerfen. Die Benotungen der FCB-Spieler in den beiden Zeitungen korrelieren.

Allgemeines Vorgehen

Test: Gilt $\varrho = 0$ für den Korrelationskoeffizienten ϱ der Wertepaare der Grundgesamtheit?

Gegeben: Eine Stichprobe von Wertepaaren.

Testgrössen:

$$\begin{array}{ll} |r_{xy}| & \text{falls Wertepaare der Grundgesamtheit bivariat normalverteilt} \\ |r_S| & \text{sonst} \end{array}$$

$$\text{Entscheid: } |r_{xy}| > r_{krit} \quad \text{bzw.} \quad |r_S| > r_{S,krit} \quad \implies \quad \varrho \neq 0$$

7.5 Der χ^2 -Test

Es gibt verschiedene Varianten und Anwendungsmöglichkeiten des χ^2 -Tests. Wir behandeln hier die Variante, mit der man testen kann, ob zwei Zufallsgrößen stochastisch unabhängig sind.

Beispiel

Zwei Behandlungen für eine bestimmte Krankheit wurden klinisch untersucht. Die Behandlung 1 erhielten 124 Patienten, die Behandlung 2 erhielten 109 Patienten. Die Resultate können in einer *Vierfelder-Tafel* übersichtlich dargestellt werden:

	Behandlung 1	Behandlung 2	Total
wirksam	102	78	180
unwirksam	22	31	53
Total	124	109	233

Wir testen wie folgt:

Nullhypothese: Die Behandlungen haben dieselbe Wirkungswahrscheinlichkeit

Signifikanzniveau: $\alpha = 5\%$

Nun nehmen wir unsere Vierfelder-Tafel, wobei nur die Randhäufigkeiten gegeben sind:

	Behandlung 1	Behandlung 2	Total
wirksam	x	z	180
unwirksam	y	w	53
Total	124	109	233

Wie gross sind die Häufigkeiten x, y, z, w , wenn wir von der Nullhypothese ausgehen?

Wir haben die folgenden Wahrscheinlichkeiten:

Die Nullhypothese besagt, dass die Ereignisse $A = (\text{Behandlung 1})$ und $B = (\text{wirksam})$ stochastisch unabhängig sind. Unter dieser Bedingung erhalten wir die folgenden Werte für x, y, z, w (wir runden diese Zahlen auf ganze Zahlen, was im Allgemeinen jedoch nicht nötig ist):

Wir sehen, dass wir nur eine der vier Zahlen x, y, z, w mit Hilfe von Wahrscheinlichkeiten berechnen müssen. Die anderen Zahlen ergeben sich direkt mit den vorgegebenen Randhäufigkeiten. Dies bedeutet, dass wir nur einen Freiheitsgrad haben.

Wenn wir die Werte für x, y, z, w mit den tatsächlich gemessenen Häufigkeiten in der ersten Tabelle vergleichen, stellen wir Abweichungen fest. Diese Abweichungen stellen wir wieder in einer Tabelle dar, und zwar berechnen wir in jedem Feld die folgende Grösse:

$$\frac{(\text{gemessene Häufigkeit} - \text{erwartete Häufigkeit})^2}{\text{erwartete Häufigkeit}}$$

Damit erhalten wir die folgende Tabelle:

	Behandlung 1	Behandlung 2
wirksam	$\frac{(102-96)^2}{96} = \frac{36}{96}$	$\frac{(78-84)^2}{84} = \frac{36}{84}$
unwirksam	$\frac{(22-28)^2}{28} = \frac{36}{28}$	$\frac{(31-25)^2}{25} = \frac{36}{25}$

Es ist kein Zufall, dass alle Zähler gleich sind. Wegen den gegebenen Randhäufigkeiten (bzw. wegen des Freiheitsgrads 1) bewirkt eine Veränderung in einem Feld eine gleich grosse Veränderung in den drei anderen Feldern (in der gleichen Zeile und in der gleichen Spalte mit umgekehrtem Vorzeichen). Quadriert ergibt dies in allen vier Feldern dieselbe Zahl.

Die Testgrösse χ^2 ist nun die Summe der berechneten Zahlen in dieser Tabelle:

$$\chi^2 = \frac{36}{96} + \frac{36}{84} + \frac{36}{28} + \frac{36}{25} = 3,529$$

Je grösser χ^2 ist, desto unwahrscheinlicher ist die Nullhypothese. Die Testgrösse χ^2 folgt der sogenannten χ^2 -Verteilung mit einem Freiheitsgrad. Den *kritischen Schrankenwert* χ_{krit}^2 entnehmen wir der entsprechenden χ^2 -Tabelle (für den Freiheitsgrad 1). Wir finden

$$\chi_{\text{krit}}^2 = 3,84.$$

Wie bei allen anderen Tests gilt allgemein

$$\chi^2 > \chi_{\text{krit}}^2 \implies \text{Nullhypothese verwerfen}$$

Für das Signifikanzniveau $\alpha = 5\%$ müssen wir in unserem Beispiel also die Nullhypothese beibehalten. Wir müssen davon ausgehen, dass die Wirkungswahrscheinlichkeit der beiden Behandlungen gleich gross ist.

Allgemeines Vorgehen

Test: Sind zwei Ereignisse (bzw. Merkmale) A und B stochastisch unabhängig?

Gegeben: Vierfelder-Tafel mit den beobachteten Häufigkeiten:

	A	$\bar{A}$	Total
B	a	b	$a + b$
$\bar{B}$	c	d	$c + d$
Total	$a + c$	$b + d$	$n = a + b + c + d$

Die Wahrscheinlichkeiten $P(A)$ und $P(B)$ berechnet man mit Hilfe der Häufigkeiten,

$$P(A) = \frac{a+c}{n} \quad \text{und} \quad P(B) = \frac{a+b}{n}.$$

Geht man nun vor wie im Beispiel, erhält man die Testgrösse

$$\chi^2 = \frac{n(ad-bc)^2}{(a+b)(a+c)(b+d)(c+d)}.$$

Vergleich mit χ^2_{krit} aus der Tabelle für den Freiheitsgrad 1.

Entscheid: $\chi^2 > \chi^2_{\text{krit}} \implies A$ und B sind nicht stochastisch unabhängig

Dieser Test kann nur für “grosse” Stichproben verwendet werden, das heisst unter den Bedingungen

$$n = a + b + c + d \geq 30, \quad a + b \geq 10, \quad a + c \geq 10, \quad b + d \geq 10, \quad c + d \geq 10.$$

Für kleinere Stichproben kann der sogenannte exakte Fisher-Test für Vierfelder-Tafeln verwendet werden.

Mehr als vier Felder

Eine Verallgemeinerung des χ^2 -Tests betrachten wir an einem Beispiel.

Beispiel

Hängt eine gute Note in der Mathematik I Prüfung mit dem Studiengang zusammen? In der Tabelle sind die Anzahl Mathematik I Prüfungsteilnehmer*innen vom HS21 aufgelistet, getrennt nach Studiengang und Resultat:

	Chemie	Bio	Geo	Pharma	Total
Note ≥ 5	20	9	4	44	77
Note < 5	12	17	16	68	113
Total	32	26	20	112	190

Wir testen wie folgt:

Nullhypothese: Gute Note und Studiengang sind voneinander unabhängig

Signifikanzniveau: $\alpha = 1\%$

Wie bei der Vierfelder-Tafel berechnen wir die Häufigkeiten unter Annahme der Nullhypothese. Wir erhalten die folgenden Häufigkeiten:

	Chemie	Bio	Geo	Pharma	Total
Note ≥ 5	12,97	10,54	8,10	45,39	77
Note < 5	19,03	15,46	11,90	66,61	113
Total	32	26	20	112	190

Die folgende Tabelle zeigt die Differenzen zwischen den tatsächlichen und den unter der Annahme der Nullhypothese erwarteten Häufigkeiten:

	Chemie	Bio	Geo	Pharma
Note ≥ 5	7,03	-1,54	-4,1	-1,39
Note < 5	-7,03	1,54	4,1	1,39

Diese Differenzen müssen wir nun quadrieren und durch die erwarteten Häufigkeiten dividieren. Die Summe dieser Zahlen ergibt

$$\chi^2 = 10,35.$$

Der Freiheitsgrad ist hier 3. Der kritische Schrankenwert aus der Tabelle (für $\alpha = 1\%$) ist

$$\chi_{\text{krit}}^2 = 11,34$$

und damit grösser als unsere berechnete Testgrösse. Wir müssen die Nullhypothese auf dem 1 %-Niveau beibehalten. Es gibt keinen hochsignifikanten Zusammenhang zwischen guter Note und Studiengang. Auf dem 5 %-Niveau könnten wir allerdings die Nullhypothese verwerfen, da $\chi_{\text{krit},5\%}^2 = 7,81$.

7.6 Vertrauensintervall für eine Wahrscheinlichkeit

Den Begriff des Vertrauensintervalles haben wir schon beim t -Test angetroffen. Dort ging es um ein Vertrauensintervall für den Mittelwert eines normalverteilten Merkmals einer Grundgesamtheit. Dieses Vertrauensintervall hing von der konkreten Stichprobe ab.

Hier geht es nun um ein Vertrauensintervall für eine Wahrscheinlichkeit. Auch hier ist das Vertrauensintervall abhängig von der konkreten Stichprobe.

Beispiel

Von 60 zufällig in einer Plantage ausgewählten Sträuchern sind 18 krank, das heisst, die relative Häufigkeit für einen kranken Strauch in der Stichprobe beträgt

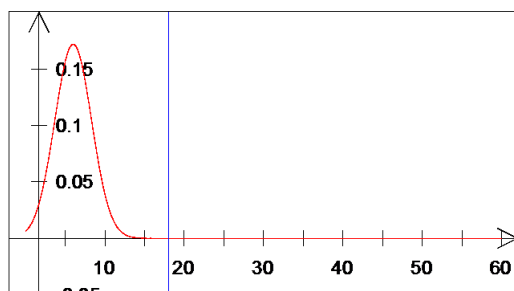
$$h_{\text{krank}} = \frac{18}{60} = 0,3.$$

Wie gross ist nun der Anteil der kranken Sträucher in der Grundgesamtheit? Das heisst, wie gross ist die Wahrscheinlichkeit p_{krank} , dass ein zufällig ausgewählter Strauch krank ist?

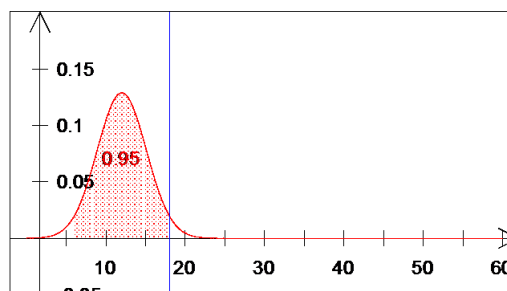
Gesucht ist ein *zweiseitiges Vertrauensintervall* für p_{krank} zum Niveau $1 - \alpha = 95\%$.

Wir fassen die Anzahl 18 der kranken Sträucher in der Stichprobe als Realisation einer binomial verteilten Zufallsgrösse auf. Dabei ist $n = 60$ und wir suchen die (unbekannte) Einzelwahrscheinlichkeit p . Für die praktische Rechnung verwenden wir die Näherung mit der Normalverteilung.

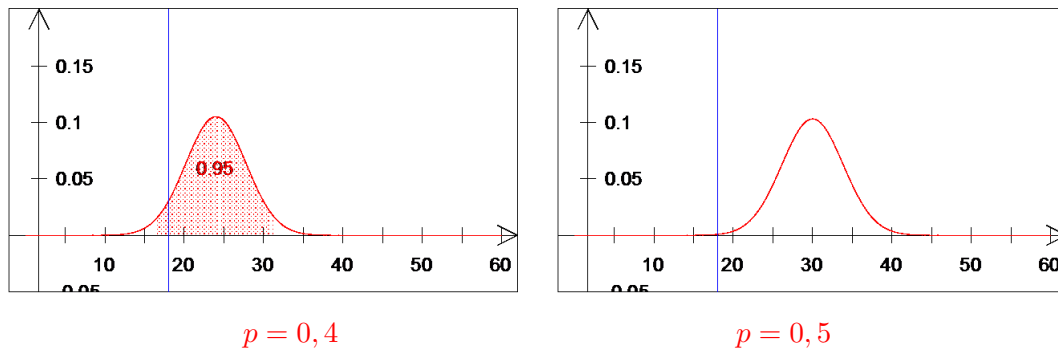
Wir können verschiedene Werte für p ausprobieren, um eine Idee zu bekommen.



$$p = 0,1$$



$$p = 0,2$$



Wir sehen, dass die untere Grenze des Vertrauensintervalles ungefähr bei 0,2 liegt und die obere Grenze ungefähr bei 0,4.

Aus der Graphik für $p = 0,2$ erkennen wir, dass für die untere Grenze

$$\mu + 1,96\sigma = 18$$

gelten muss. Der Faktor 1,96 kommt aus der Tabelle von Seite 62 im Skript (wir suchen ein 95 %-Vertrauensintervall). Für den Erwartungswert μ und die Varianz σ^2 der Binomialverteilung gilt

Wir erhalten damit die Gleichung

Dies ist eine quadratische Gleichung für p . Die Lösungen sind 0,199 und 0,425. Die untere Grenze ist also 0,199 (die Zahl 0,425 ist keine Lösung der ersten, unquadrierten Gleichung).

Für die obere Grenze des Vertrauensintervalles muss

$$\mu - 1,96\sigma = 18$$

gelten. Auch diese Gleichung führt zu einer quadratischen Gleichung für p , und zwar zu derselben Gleichung wie oben. Für die obere Grenze erhalten wir daher 0,425.

Das Vertrauensintervall für p_{krank} zum Niveau $1 - \alpha = 95\%$ ist also gegeben durch

$$[0,199; 0,425] .$$

Allgemeines Vorgehen

Gegeben ist die Anzahl Elemente einer Stichprobe (vom Umfang n) mit einer bestimmten Eigenschaft. Diese Anzahl (im Beispiel die Zahl 18) ist gleich nh für die relative Häufigkeit h . Der Anteil der Elemente der Grundgesamtheit mit dieser Eigenschaft sei p . Die Grenzen für ein 95 %-Vertrauensintervall für p erhalten wir als Lösungen der quadratischen Gleichung

$$1,96^2 np(1-p) = (nh - np)^2 .$$

Kürzen mit n ergibt die Gleichung

$$1,96^2 p(1-p) = n(h-p)^2.$$

Für Vertrauensintervalle mit anderen Prozentzahlen müssen wir die Zahl 1,96 entsprechend ändern. Zum Beispiel müssen wir sie für ein 99 %-Vertrauensintervall durch 2,576 ersetzen (wie der Tabelle auf Seite 62 zu entnehmen ist).

Für grosse Stichproben können die Grenzen des 95 %-Vertrauensintervalles mit der Formel

$$h \pm 1,96 \sqrt{\frac{h(1-h)}{n}}$$

berechnet werden.

In unserem Beispiel erhalten wir mit dieser Formel 0,184 für die untere und 0,416 für die obere Grenze.

8 Lineare Abbildungen

In diesem Kapitel untersuchen wir lineare Abbildungen von $\mathbb{R}^n$ nach $\mathbb{R}^m$ wie zum Beispiel Spiegelungen, Drehungen, Streckungen und Orthogonalprojektionen in $\mathbb{R}^2$ und $\mathbb{R}^3$. Man nennt eine Abbildung linear, wenn man sie durch eine Matrix beschreiben kann. Die Komposition (d.h. Verknüpfung) von zwei linearen Abbildungen kann dadurch einfach berechnet werden. Weiter können an der Matrix einer linearen Abbildung die wichtigsten Eigenschaften der Abbildung abgelesen werden.

8.1 Definition und Beispiele

Im letzten Semester haben wir reelle Funktionen (d.h. Funktionen von $\mathbb{R}$ nach $\mathbb{R}$) betrachtet. Nun kann man nicht nur Zahlen aus $\mathbb{R}$, sondern auch Punkten in $\mathbb{R}^2$, $\mathbb{R}^3$ oder allgemein $\mathbb{R}^n$ eine reelle Zahl zuordnen. Zum Beispiel

Eine Funktion $f : \mathbb{R}^n \longrightarrow \mathbb{R}$ nennt man *reellwertige Funktion von n reellen Variablen*.

Seien nun $f_1, \dots, f_m$ reellwertige Funktionen von n reellen Variablen, das heisst

$$\begin{aligned}w_1 &= f_1(x_1, \dots, x_n) \\w_2 &= f_2(x_1, \dots, x_n) \\&\vdots \\w_m &= f_m(x_1, \dots, x_n) .\end{aligned}$$

Durch diese Gleichungen wird jedem Punkt $(x_1, \dots, x_n)$ in $\mathbb{R}^n$ genau ein Punkt $(w_1, \dots, w_m)$ in $\mathbb{R}^m$ zugeordnet. Wir erhalten damit eine Abbildung

$$T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$$

durch

$$T(x_1, \dots, x_n) = (w_1, \dots, w_m) .$$

Der Buchstabe T steht für *Transformation*, denn Abbildungen von $\mathbb{R}^n$ nach $\mathbb{R}^m$ werden auch Transformationen genannt.

Beispiel

Die Gleichungen

$$\begin{aligned}w_1 &= x_1 + x_2 \\w_2 &= 3x_1x_2 \\w_3 &= x_1^2 - x_2^2\end{aligned}$$

definieren eine Abbildung $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^3$ durch

$$T(x_1, x_2) = (x_1 + x_2, 3x_1x_2, x_1^2 - x_2^2) .$$

Fasst man n -Tupel nicht als Punkte P in $\mathbb{R}^n$ sondern als Ortsvektoren $\vec{x} = \overrightarrow{OP}$ in $\mathbb{R}^n$ auf, dann erhalten wir eine Abbildung, welche Vektoren in $\mathbb{R}^n$ auf Vektoren in $\mathbb{R}^m$ abbildet. Tatsächlich werden wir Elemente aus $\mathbb{R}^n$ und $\mathbb{R}^m$ stets als Vektoren betrachten, da wir ja mit der Struktur dieser Räume als Vektorräume vertraut sind und insbesondere lineare Abbildungen den $\mathbb{R}^n$ auf einen Vektorraum in $\mathbb{R}^m$ abbilden. Die Funktionsvorschrift $T(x_1, \dots, x_n) = (w_1, \dots, w_m)$ bedeutet in diesem Fall also, dass der Vektor $\vec{x}$ in $\mathbb{R}^n$ mit den Komponenten $x_1, \dots, x_n$ auf den Vektor $\vec{w}$ in $\mathbb{R}^m$ mit den Komponenten $w_1, \dots, w_m$ abgebildet wird. Oft beschreiben wir die Abbildung direkt mit Vektoren: $T(\vec{x}) = \vec{w}$.

Definition Eine Abbildung $T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ definiert durch $T(x_1, \dots, x_n) = (w_1, \dots, w_m)$ heisst *linear*, wenn

$$\begin{aligned} w_1 &= a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n \\ w_2 &= a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n \\ &\vdots \\ w_m &= a_{m1}x_1 + a_{m2}x_2 + \dots + a_{mn}x_n \end{aligned}$$

für reelle Zahlen a_{ij} ($1 \leq i \leq m, 1 \leq j \leq n$).

Das heisst, alle Variablen $x_1, \dots, x_n$ kommen in den Komponentenfunktionen $w_1, \dots, w_m$ linear (d.h. zur ersten Potenz oder gar nicht) vor. Dieses Gleichungssystem kann man als Matrixmultiplikation schreiben

$$\begin{pmatrix} w_1 \\ \vdots \\ w_m \end{pmatrix} = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & & \vdots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix},$$

das heisst

$$T(\vec{x}) = \vec{w} = [T]\vec{x},$$

wobei $[T] = (a_{ij})$ die sogenannte *Darstellungsmatrix* der linearen Abbildung T ist. Die Einträge der Darstellungsmatrix hängen von der Wahl der Basen von $\mathbb{R}^n$ und $\mathbb{R}^m$ ab. Wir wählen vorerst stets die Standardbasen; die Darstellungsmatrix nennt man in diesem Fall auch *Standardmatrix* von T .

Beispiele

$$1. \quad T : \mathbb{R}^3 \longrightarrow \mathbb{R}^3, \quad T(x_1, x_2, x_3) = (x_1 + x_2 - x_3, 2x_1 - 3x_3, 5x_2)$$

2. $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^3, \quad T(x_1, x_2) = (x_1 + 2x_2, x_1, 0)$

3. $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2, \quad T(x_1, x_2) = (x_1^2 + x_2, 3x_1)$

4. $T : \mathbb{R}^3 \longrightarrow \mathbb{R}^2, \quad T(x_1, x_2, x_3) = (x_1 + x_2 + 4x_3, 1)$

5. $T : \mathbb{R} \longrightarrow \mathbb{R}$ ist linear genau dann, wenn $T(x) = ax$ für eine reelle Zahl a .
In diesem Fall ist $[T] = (a)$.

Eine Abbildung $T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ ist also linear, wenn es eine $m \times n$ -Matrix A gibt (nämlich $A = [T]$), so dass $T(\vec{x}) = A\vec{x}$. Insbesondere gilt

$$T(\vec{0}) = A\vec{0} = \vec{0}$$

in $\mathbb{R}^m$. Bei einer linearen Abbildung gilt stets, dass $\vec{0}$ in $\mathbb{R}^n$ auf $\vec{0}$ in $\mathbb{R}^m$ abgebildet wird.

Lineare Abbildungen werden also durch Matrizen beschrieben. Umgekehrt beschreibt jede Matrix eine lineare Abbildung. Ist A eine $m \times n$ -Matrix, so definiert diese eine lineare Abbildung $T : \mathbb{R}^n \longrightarrow \mathbb{R}^m$ durch $T(\vec{x}) = A\vec{x}$.

Beispiel

Die Matrix

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

definiert eine lineare Abbildung $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ durch

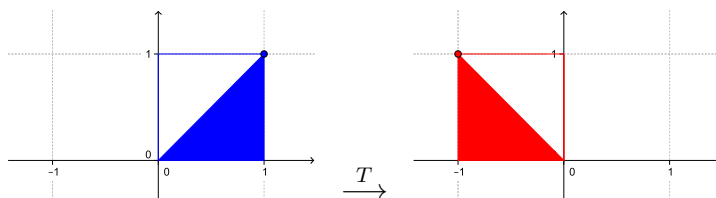
$$T(\vec{x}) = A\vec{x} = \begin{pmatrix} 1 & 2 \\ 3 & 0 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} x_1 + 2x_2 \\ 3x_1 \end{pmatrix}.$$

Was ist also zum Beispiel das Bild des Vektors $\vec{x} = \begin{pmatrix} 7 \\ 5 \end{pmatrix}$ unter T ?

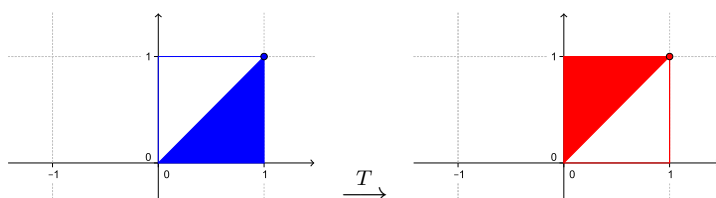
Beispiele von linearen Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$

Die wichtigsten linearen Abbildungen von $\mathbb{R}^2$ nach $\mathbb{R}^2$ kennen Sie aus der Schule: Spiegelungen, Projektionen, Drehungen und Streckungen.

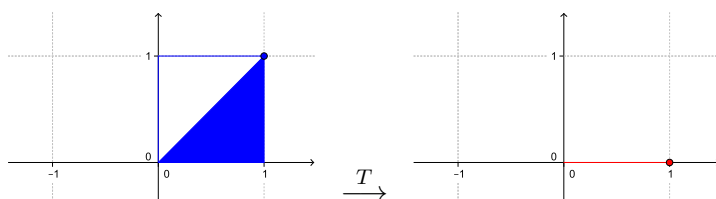
1. Spiegelung an der y -Achse



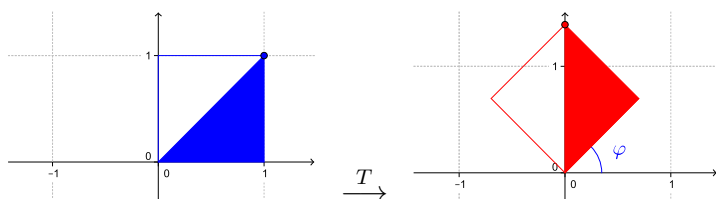
2. Spiegelung an der Geraden $y = x$



3. Orthogonalprojektion auf die x -Achse



4. Drehung um den Winkel φ um den Ursprung



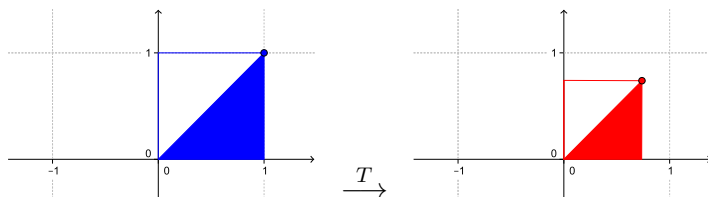
Es gilt (Herleitung später)

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

Für den Drehwinkel $\varphi = 90^\circ$, bzw. $\varphi = 45^\circ$ erhalten wir damit die Darstellungsmatrizen

$$[T] = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix} , \text{ bzw. } [T] = \begin{pmatrix} \frac{\sqrt{2}}{2} & -\frac{\sqrt{2}}{2} \\ \frac{\sqrt{2}}{2} & \frac{\sqrt{2}}{2} \end{pmatrix} .$$

5. Streckung um den Faktor k mit dem Ursprung als Streckzentrum



Spezialfälle:

- $k = 1$: $T = \text{Id}$ Identität, $T(x, y) = (x, y)$, $[T] = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$
- $k = 0$: $T = 0$ Nullabbildung, $T(x, y) = (0, 0)$, $[T] = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$

Komposition von linearen Abbildungen

Seien $T_1 : \mathbb{R}^n \rightarrow \mathbb{R}^k$ und $T_2 : \mathbb{R}^k \rightarrow \mathbb{R}^m$ zwei lineare Abbildungen. Wir verknüpfen die beiden Abbildungen wie folgt:

$$\begin{array}{ccccc} \mathbb{R}^n & \xrightarrow{T_1} & \mathbb{R}^k & \xrightarrow{T_2} & \mathbb{R}^m \\ \vec{x} & \mapsto & T_1(\vec{x}) & \mapsto & T_2(T_1(\vec{x})) \end{array}$$

Die Verknüpfung von T_1 und T_2 ist demnach eine Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$, die *Komposition* von T_2 mit T_1 genannt und mit $T_2 \circ T_1$ bezeichnet wird (vgl. letztes Semester, Kapitel 1); das heisst

$$(T_2 \circ T_1)(\vec{x}) = T_2(T_1(\vec{x})) \quad \text{für alle } \vec{x} \text{ in } \mathbb{R}^n.$$

Satz 8.1 Sind T_1 und T_2 linear, dann ist auch die Komposition $T_2 \circ T_1$ linear. Für die Darstellungsmatrix $[T_2 \circ T_1]$ der Komposition gilt

$$[T_2 \circ T_1] = [T_2][T_1].$$

Die Reihenfolge der beiden Matrizen auf der rechten Seite der Gleichung ist dabei wichtig, denn im Allgemeinen gilt ja $[T_2][T_1] \neq [T_1][T_2]$.

Die folgende Zeile beweist den Satz 8.1. Für $\vec{x}$ in $\mathbb{R}^n$ gilt

$$(T_2 \circ T_1)(\vec{x}) = T_2(T_1(\vec{x})) = T_2([T_1]\vec{x}) = [T_2][T_1]\vec{x}.$$

Beispiel

Sei $T_1 : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Spiegelung an der Geraden $y = x$ und $T_2 : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Orthogonalprojektion auf die y -Achse. Man bestimme die Komposition $T_2 \circ T_1$.

Gilt im Beispiel $T_1 \circ T_2 = T_2 \circ T_1$? Wir berechnen

$$[T_1 \circ T_2] = [T_1][T_2] = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \neq [T_2][T_1] = [T_2 \circ T_1]$$

Es folgt, dass $T_1 \circ T_2 \neq T_2 \circ T_1$.

8.2 Eigenschaften

Eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$ hat die Eigenschaft, dass sie mit den Vektorraumoperationen in $\mathbb{R}^n$ und $\mathbb{R}^m$ verträglich ist; das heisst, es spielt keine Rolle, ob zwei Vektoren in $\mathbb{R}^n$ zuerst addiert und danach abgebildet werden oder ob sie zuerst abgebildet und danach ihre Bilder in $\mathbb{R}^m$ addiert werden (entsprechend für die skalare Multiplikation).

Satz 8.2 *Eine Abbildung T von $\mathbb{R}^n$ nach $\mathbb{R}^m$ ist genau dann linear, wenn für alle $\vec{x}, \vec{y}$ in $\mathbb{R}^n$ und k in $\mathbb{R}$ die folgenden zwei Linearitätsbedingungen gelten:*

- (1) $T(\vec{x} + \vec{y}) = T(\vec{x}) + T(\vec{y})$
- (2) $T(k\vec{x}) = kT(\vec{x})$

Insbesondere ist das Bild des Vektorraums $\mathbb{R}^n$ unter einer linearen Abbildung stets wieder ein Vektorraum in $\mathbb{R}^m$.

Warum gilt Satz 8.2? Nun, ist T eine lineare Abbildung, dann gilt $T(\vec{x}) = A\vec{x}$ für die Darstellungsmatrix $A = [T]$ von T , und die Linearitätsbedingungen (1) und (2) folgen aufgrund der Distributivgesetze für Matrizen.

Umgekehrt, ist T eine Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$, welche die Linearitätsbedingungen (1) und (2) aus Satz 8.2 erfüllt, dann gilt $T(\vec{x}) = A\vec{x}$ für diejenige Matrix A , deren Spalten die Bilder (unter T) der Basisvektoren $\vec{e}_1, \dots, \vec{e}_n$ von $\mathbb{R}^n$ sind, das heisst für

$$A = (T(\vec{e}_1) \ \cdots \ T(\vec{e}_n)).$$

Die Abbildung T lässt sich also als Matrixmultiplikation schreiben und ist deshalb linear.

Damit gilt insbesondere die folgende wichtige Tatsache.

Satz 8.3 *Sei T eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^m$. Dann stehen in den Spalten der Darstellungsmatrix $[T]$ die Bilder der Basisvektoren $\vec{e}_1, \dots, \vec{e}_n$ von $\mathbb{R}^n$.*

Beispiel

Sei $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Spiegelung an der y -Achse.

Der Satz 8.3 ist vor allem dann praktisch, wenn eine lineare Abbildung geometrisch beschrieben ist und die Abbildungsvorschrift $T(\vec{x}) = \vec{w}$ nicht bekannt ist. Bestimmt man die Bilder der Basisvektoren, dann ist die Abbildung eindeutig beschrieben.

Beispiele

1. Sei $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ die Drehung um den Winkel φ um den Ursprung.

2. Sei $T : \mathbb{R}^3 \longrightarrow \mathbb{R}^3$ die Drehung um den Winkel φ um die x -Achse. Dies bedeutet, dass der Basisvektor $\vec{e}_1$ bei der Drehung unverändert bleibt und die Basisvektoren $\vec{e}_2$ und $\vec{e}_3$ eine Drehung in der yz -Ebene beschreiben, analog zur Drehung in der Ebene im Beispiel vorher. Damit erhalten wir

$$T(\vec{e}_1) = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, T(\vec{e}_2) = \begin{pmatrix} 0 \\ \cos \varphi \\ \sin \varphi \end{pmatrix}, T(\vec{e}_3) = \begin{pmatrix} 0 \\ -\sin \varphi \\ \cos \varphi \end{pmatrix} \implies [T] = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \varphi & -\sin \varphi \\ 0 & \sin \varphi & \cos \varphi \end{pmatrix}$$

8.3 Basiswechsel

Tatsächlich gilt Satz 8.3 nicht nur für die Standardbasis $\vec{e}_1, \dots, \vec{e}_n$ von $\mathbb{R}^n$, sondern für jede beliebige Basis $\mathcal{B} = \{\vec{u}_1, \dots, \vec{u}_n\}$ von $\mathbb{R}^n$. Ist $T : \mathbb{R}^n \longrightarrow \mathbb{R}^n$ eine lineare Abbildung (wir beschränken uns auf den Fall $m = n$), dann erhält man die Darstellungsmatrix $[T]_{\mathcal{B}}$ von T bezüglich der Basis $\mathcal{B}$, indem man die Koordinaten der Bilder der Basisvektoren $\vec{u}_1, \dots, \vec{u}_n$ in die Spalten der Matrix schreibt.

Beispiel

Sei $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ die Spiegelung an der Geraden g mit dem Richtungsvektor $\vec{u}_1 = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix}$.

Bezüglich der Standardbasis wäre die Darstellungsmatrix $[T]$ von T gegeben durch

$$[T] = \begin{pmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ \frac{\sqrt{3}}{2} & -\frac{1}{2} \end{pmatrix}.$$

Damit haben wir zwei verschiedene Möglichkeiten, das Bild eines Vektors, zum Beispiel $\vec{u}_1$, zu berechnen. In der Standardbasis rechnen wir

$$T(\vec{u}_1) = [T] \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ \frac{\sqrt{3}}{2} & -\frac{1}{2} \end{pmatrix} \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} = \vec{u}_1$$

und in der Basis $\mathcal{B} = \{\vec{u}_1, \vec{u}_2\}$ erhalten wir

$$T(\vec{u}_1) = [T]_{\mathcal{B}} \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}_{\mathcal{B}} \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}} = \vec{u}_1.$$

Im Gegensatz zur Matrix $[T]$ ist die Darstellungsmatrix $[T]_{\mathcal{B}}$ sofort erkennbar als Matrix einer Spiegelung. Zudem ist das Rechnen mit Diagonalmatrizen einfacher. Aus diesen Gründen ist ein Basiswechsel manchmal sinnvoll.

Tatsächlich gibt es eine (algebraische) Beziehung zwischen den Matrizen $[T]_{\mathcal{B}}$ und $[T]$. Wir schreiben die Basisvektoren der Basis $\mathcal{B}$ als Spalten in eine 2×2 -Matrix. Eine übliche Bezeichnung dieser Matrix ist P^{-1} (tatsächlich ist diese Matrix invertierbar, da die Spalten linear unabhängig sind). Wir setzen also

$$P^{-1} = (\vec{u}_1 \ \vec{u}_2) = \begin{pmatrix} \sqrt{3} & -1 \\ 1 & \sqrt{3} \end{pmatrix}.$$

Damit ist

$$P = (P^{-1})^{-1} = \frac{1}{4} \begin{pmatrix} \sqrt{3} & 1 \\ -1 & \sqrt{3} \end{pmatrix}$$

und es gilt

$$[T]_{\mathcal{B}} = P [T] P^{-1}.$$

Die Matrix P^{-1} beschreibt dabei den Basiswechsel von der Basis $\mathcal{B}$ zur Standardbasis und die Matrix P beschreibt den Basiswechsel von der Standardbasis zur Basis $\mathcal{B}$. Das heisst,

$$P^{-1} \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}} = \begin{pmatrix} \sqrt{3} & -1 \\ 1 & \sqrt{3} \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}} = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} \quad \text{und} \quad P \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} = \frac{1}{4} \begin{pmatrix} \sqrt{3} & 1 \\ -1 & \sqrt{3} \end{pmatrix} \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}_{\mathcal{B}}.$$

Satz 8.4 Sei $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung und $\mathcal{B}$ eine beliebige Basis von $\mathbb{R}^n$. Dann gilt

$$[T]_{\mathcal{B}} = P [T] P^{-1}$$

wobei in den Spalten von P^{-1} die Basisvektoren von $\mathcal{B}$ stehen.

Die Kunst ist nun, eine Basis $\mathcal{B}$ zu finden, so dass die Matrix $[T]_{\mathcal{B}}$ diagonal ist. Dies ist tatsächlich “nur” ein Handwerk, welches wir im nächsten Kapitel erlernen. Man muss die sogenannten Eigenwerte und Eigenvektoren der Matrix $[T]$ berechnen.

Lineare Abbildungen zwischen allgemeinen Vektorräumen (für Interessierte)

In der Literatur werden lineare Abbildungen oft direkt durch die Linearitätsbedingungen von Satz 8.2 definiert. Dies hat den Vorteil, dass man sich nicht auf Abbildungen von $\mathbb{R}^n$ nach $\mathbb{R}^m$ beschränken muss. Man geht von zwei (reellen) Vektorräumen V und W aus und nennt eine Abbildung $T : V \rightarrow W$ *linear*, wenn gilt

$$(1) \quad T(\mathbf{u} + \mathbf{v}) = T(\mathbf{u}) + T(\mathbf{v})$$

$$(2) \quad T(k \mathbf{v}) = k T(\mathbf{v})$$

für alle $\mathbf{u}, \mathbf{v} \in V$ und $k \in \mathbb{R}$.

Beispiele

1. Sei $V = \{ax^2 + bx + c \mid a, b, c \in \mathbb{R}\}$ die Menge aller Polynome vom Grad ≤ 2 . Wir haben im letzten Semester (Kapitel 9, Seite 148) gesehen, dass V ein Vektorraum ist. Sei nun $T : V \rightarrow V$ definiert durch die Ableitung

$$T(p) = p',$$

das heisst, $T(ax^2 + bx + c) = 2ax + b$ (man leitet das Polynom $p = p(x)$ nach x ab). Dies ist eine lineare Abbildung, denn für Polynome $p_1, p_2 \in V$ und $k \in \mathbb{R}$ gilt mit den Ableitungsregeln

$$\begin{aligned} T(p_1 + p_2) &= (p_1 + p_2)' = p_1' + p_2' = T(p_1) + T(p_2) \\ T(kp_1) &= (kp_1)' = kp_1' = kT(p_1) \end{aligned}$$

2. Sei V die Menge aller 2×2 -Matrizen und $W = \mathbb{R}$. Wir betrachten $T : V \rightarrow W$ definiert durch

$$T(A) = \det(A)$$

für $A \in V$. Dies ist keine lineare Abbildung, da im Allgemeinen gilt

$$T(A + B) = \det(A + B) \neq \det(A) + \det(B) = T(A) + T(B).$$

Wählt man je eine Basis für V und W , dann kann man die lineare Abbildung $T : V \rightarrow W$ durch eine (reelle) $m \times n$ -Matrix beschreiben, wobei $m = \dim(W)$ und $n = \dim(V)$.

Betrachten wir den Spezialfall $W = V$. Wir wählen also eine Basis $\mathcal{B}$ von V . Nach Satz 8.3 (leicht angepasst), stehen dann in den Spalten der Darstellungsmatrix $[T]_{\mathcal{B}}$ die Koordinaten der Bilder der Basisvektoren.

Beispiel

Betrachten wir das 1. Beispiel oben mit $V = \{ax^2 + bx + c \mid a, b, c \in \mathbb{R}\}$ und $T(p) = p'$. Eine Basis für V ist $\mathcal{B} = \{x^2, x, 1\}$. Nun bestimmen wir die (Koordinaten der) Bilder der Basisvektoren:

$$\left. \begin{aligned} T(x^2) &= 2x &= 0 \cdot x^2 + 2 \cdot x + 0 \cdot 1 \\ T(x) &= 1 &= 0 \cdot x^2 + 0 \cdot x + 1 \cdot 1 \\ T(1) &= 0 &= 0 \cdot x^2 + 0 \cdot x + 0 \cdot 1 \end{aligned} \right\} \implies [T]_{\mathcal{B}} = \begin{pmatrix} 0 & 0 & 0 \\ 2 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}$$

Die Matrix $[T]_{\mathcal{B}}$ beschreibt nun die Abbildung T im folgenden Sinn.

Sei beispielsweise $p(x) = 3x^2 - 2x + 5$. Der *Koordinatenvektor* von p bzgl. der Basis $\mathcal{B}$ ist $\vec{v} = \begin{pmatrix} 3 \\ -2 \\ 5 \end{pmatrix}$. Dann ist der Koordinatenvektor von $T(p)$ gegeben durch

$$[T]_{\mathcal{B}} \vec{v} = \begin{pmatrix} 0 & 0 & 0 \\ 2 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} 3 \\ -2 \\ 5 \end{pmatrix} = \begin{pmatrix} 0 \\ 6 \\ -2 \end{pmatrix}.$$

Wir erhalten also

$$T(p) = 0 \cdot x^2 + 6 \cdot x + (-2) \cdot 1 = 6x - 2 \quad (= p'(x)).$$

8.4 Bedeutung der Determinante einer Darstellungsmatrix

An der Determinante der Darstellungsmatrix einer linearen Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ (bzw. $T : V \rightarrow V$) können gewisse Eigenschaften von T abgelesen werden.

Umkehrbare lineare Abbildungen

Eine lineare Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ ist umkehrbar, wenn es eine lineare Abbildung $T^{-1} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ gibt, so dass

$$T \circ T^{-1} = T^{-1} \circ T = \text{Id}$$

die Identität ist (vgl. letztes Semester, Kapitel 1). Für die Darstellungsmatrizen (bzgl. der Standardbasis oder einer anderen Basis) folgt $[T][T^{-1}] = [T^{-1}][T] = E$ (wobei E die Einheitsmatrix bezeichnet). Das heisst, $[T]$ ist invertierbar und

$$[T^{-1}] = [T]^{-1}.$$

In Worten: Die Darstellungsmatrix der Umkehrabbildung T^{-1} ist die Inverse $[T]^{-1}$ der Darstellungsmatrix $[T]$ von T .

Ist umgekehrt $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ eine lineare Abbildung und $[T]$ invertierbar, dann ist T umkehrbar und T^{-1} ist definiert durch $T^{-1}(\vec{x}) = [T]^{-1}\vec{x}$ für alle $\vec{x}$ in $\mathbb{R}^n$.

Es gilt also

$$T \text{ umkehrbar} \iff [T] \text{ invertierbar}.$$

Mit Satz 8.5 (b), Seite 144, vom letzten Semester erhalten wir den folgenden Satz.

Satz 8.5 Sei $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ linear. Dann gilt

$$T \text{ umkehrbar} \iff \det([T]) \neq 0.$$

Wegen Satz 8.4 ist die Standardmatrix $[T]$ invertierbar, genau dann wenn die Darstellungsmatrix $[T]_{\mathcal{B}}$ zu jeder anderen Basis $\mathcal{B}$ invertierbar ist. Der Satz 8.5 gilt also auch für $[T]_{\mathcal{B}}$ anstelle von $[T]$.

Beispiele

1. Sei $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ die Orthogonalprojektion auf die x -Achse.

2. Sei $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$ die Drehung um den Ursprung um den Winkel φ .

Volumenänderung

Sei A eine $n \times n$ -Matrix und $T : \mathbb{R}^n \longrightarrow \mathbb{R}^n$ die lineare Abbildung $T(\vec{x}) = A\vec{x}$. Eine "geometrische Figur" mit Volumen V wird durch T abgebildet auf eine Figur mit Volumen V' .

Satz 8.6 Es gilt $V' = |\det(A)| \cdot V$.

Für $n = 2$ muss Volumen V durch Flächeninhalt F ersetzt werden.

Betrachten wir einen Spezialfall in $\mathbb{R}^2$, nämlich das von $\vec{e}_1$ und $\vec{e}_2$ aufgespannte Quadrat mit Flächeninhalt $F = 1$. Durch T wird dieses Quadrat auf das Parallelogramm, aufgespannt von $T(\vec{e}_1)$ und $T(\vec{e}_2)$, abgebildet. Mit Satz 8.6, Seite 145, vom letzten Semester folgt

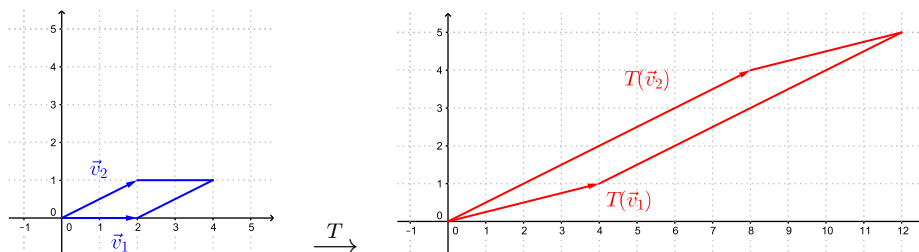
$$F' = |\det(T(\vec{e}_1) \ T(\vec{e}_2))| = |\det(A)| = |\det(A)| \cdot F,$$

wie in Satz 8.6 behauptet.

Beispiel

Gegeben seien das Parallelogramm aufgespannt von $\vec{v}_1 = \begin{pmatrix} 2 \\ 0 \end{pmatrix}$, $\vec{v}_2 = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$ und $T : \mathbb{R}^2 \longrightarrow \mathbb{R}^2$, $T(\vec{x}) = A\vec{x}$ mit

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



Spiegelungen und Drehungen im $\mathbb{R}^2$ und $\mathbb{R}^3$

Mit einer Spiegelung im $\mathbb{R}^2$ ist hier eine Spiegelung an einer Geraden durch den Ursprung gemeint und mit einer Spiegelung im $\mathbb{R}^3$ eine Spiegelung an einer Ebene durch den Ursprung. Mit einer Drehung im $\mathbb{R}^2$ ist eine Drehung um den Ursprung gemeint und mit einer Drehung im $\mathbb{R}^3$ eine Drehung um eine Drehachse durch den Ursprung. Solche Spiegelungen und Drehungen sind lineare Abbildungen.

Diese Spiegelungen und Drehungen sind lngen- und winkeltreu, das heisst, Strecken werden auf Strecken gleicher Lnge abgebildet und die Winkel bleiben erhalten. Ist also A die Standardmatrix einer solchen Spiegelung oder Drehung, dann gilt

$$\|A\vec{v}\| = \|\vec{v}\|$$

fur einen beliebigen Vektor $\vec{v}$ in $\mathbb{R}^2$ oder $\mathbb{R}^3$.

Die Lnge eines Vektors kann man mit dem Skalarprodukt ausdrcken:

Es gilt also

$$A^T A = E.$$

Satz 8.7 *Die Darstellungsmatrix einer Drehung oder Spiegelung im $\mathbb{R}^2$ oder $\mathbb{R}^3$ bzgl. einer Orthonormalbasis ist orthogonal.*

Ist A die Darstellungsmatrix einer Drehung oder Spiegelung T bzgl. einer beliebigen Basis von $\mathbb{R}^2$ oder $\mathbb{R}^3$, dann gilt

$$|\det(A)| = 1$$

denn der Flcheninhalt (bzw. das Volumen) einer geometrischen Figur bleibt durch eine Spiegelung oder Drehung unverndert.

Genauer gilt das Folgende:

$$\begin{aligned} T \text{ Drehung} &\implies \det(A) = 1 \\ T \text{ Spiegelung} &\implies \det(A) = -1 \end{aligned}$$

Dies folgt aus der Tatsache, dass eine Drehung orientierungserhaltend und eine Spiegelung orientierungsumkehrend ist.

Es gilt namlich allgemein fur eine lineare Abbildung T in $\mathbb{R}^2$ oder $\mathbb{R}^3$, dass

$$\begin{aligned} \det([T]) > 0 &\iff \text{die Orientierung einer geometrischen Figur bleibt erhalten} \\ \det([T]) < 0 &\iff \text{die Orientierung einer geometrischen Figur wird umgekehrt} \end{aligned}$$

9 Eigenwerte und Eigenvektoren

Wir haben im vorhergehenden Kapitel gesehen, dass eine lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$ durch verschiedene Darstellungsmatrizen beschrieben werden kann (je nach Wahl der Basis für $\mathbb{R}^n$). Im Idealfall kann die Matrix diagonal gewählt werden. Dazu muss eine Basis aus sogenannten Eigenvektoren existieren.

9.1 Bestimmung von Eigenwerten und Eigenvektoren

Definition Sei A eine (reelle) $n \times n$ -Matrix. Ist λ eine reelle Zahl und $\vec{u} \neq \vec{0}$ in $\mathbb{R}^n$ mit

$$A\vec{u} = \lambda\vec{u},$$

dann heisst λ *Eigenwert* von A und $\vec{u}$ ein zu λ dazugehöriger *Eigenvektor* von A .

Wir beschränken uns im Moment auf reelle Eigenwerte und Eigenvektoren, werden aber in einem späteren Abschnitt auf komplexe Matrizen, Eigenwerte und Eigenvektoren eingehen.

Beispiel

Der Vektor $\vec{u} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ ist ein Eigenvektor zum Eigenwert $\lambda = 2$ der Matrix

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

Um die Eigenwerte und Eigenvektoren einer Matrix A zu bestimmen, schreibt man die Gleichung $A\vec{u} = \lambda\vec{u}$ als $\lambda\vec{u} - A\vec{u} = \vec{0}$ oder

$$(\lambda E - A)\vec{u} = \vec{0}. \quad (\text{EV})$$

Dies ist ein homogenes lineares Gleichungssystem. Für welche λ hat dieses System Lösungen $\vec{u} \neq \vec{0}$? Ist die Matrix $\lambda E - A$ invertierbar, dann ist $\vec{u} = (\lambda E - A)^{-1}\vec{0} = \vec{0}$ die einzige Lösung. Es gibt also Lösungen $\vec{u} \neq \vec{0}$ genau dann, wenn $\lambda E - A$ nicht invertierbar ist, das heisst genau dann wenn

$$\det(\lambda E - A) = 0. \quad (\text{EW})$$

Diese Gleichung heisst *charakteristische Gleichung* von A . Ihre Lösungen sind die Eigenwerte von A . Der Ausdruck

$$p(\lambda) = \det(\lambda E - A)$$

ist ein Polynom in λ , das sogenannte *charakteristische Polynom*. Es hat den Grad n und da ein Polynom vom Grad n höchstens n verschiedene Nullstellen hat, hat eine $n \times n$ -Matrix höchstens n verschiedene Eigenwerte.

Um zu jedem gefundenen Eigenwert λ die zugehörigen Eigenvektoren zu bestimmen, muss anschliessend das homogene lineare System (EV) gelöst werden.

Beispiel

Gesucht sind die Eigenwerte und Eigenvektoren der Matrix

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

Wir sehen im Beispiel und auch allgemein anhand der Gleichung (EV), dass die Eigenvektoren zu einem festen Eigenwert zusammen mit dem Nullvektor einen reellen Vektorraum, den sogenannten *Eigenraum*, bilden. Das heisst, für einen fixen Eigenwert sind Vielfache $\neq \vec{0}$ von einem Eigenvektor und Linearkombinationen $\neq \vec{0}$ von Eigenvektoren auch Eigenvektoren zum selben Eigenwert. Es genügt also, eine Basis von jedem Eigenraum anzugeben.

Vorgehen zur Bestimmung von Eigenwerten und Eigenvektoren einer Matrix A

- (1) Bestimme alle Lösungen λ der charakteristischen Gleichung

$$\det(\lambda E - A) = 0. \quad (\text{EW})$$

Dies sind die Eigenwerte von A .

- (2) Bestimme für jeden Eigenwert λ die Lösungen $\vec{u} \neq \vec{0}$ des linearen Gleichungssystems

$$(\lambda E - A)\vec{u} = \vec{0}. \quad (\text{EV})$$

Diese sind die Eigenvektoren von A zum Eigenwert λ und bilden, zusammen mit $\vec{0}$, den Eigenraum von λ , so dass die Angabe einer Basis dieses Eigenraums genügt.

Beispiel

Gesucht sind die Eigenwerte und Eigenvektoren der Matrix

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

Geometrische Interpretation

Eine $n \times n$ -Matrix A kann aufgefasst werden als Standardmatrix $A = [T]$ einer linearen Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$. Die Gleichung $A\vec{u} = \lambda\vec{u}$ für ein λ in $\mathbb{R}$ und $\vec{u}$ in $\mathbb{R}^n$ ist dadurch gleichbedeutend mit

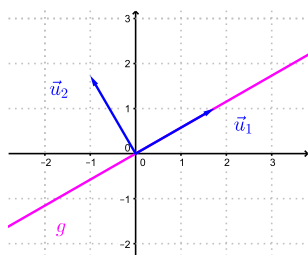
$$T(\vec{u}) = \lambda\vec{u}.$$

Das heisst, die Eigenvektoren von A sind genau diejenigen Vektoren, welche durch T lediglich gestreckt oder gestaucht (mit einer Richtungsumkehrung, falls $\lambda < 0$) werden.

Beispiele

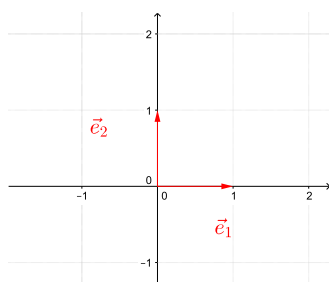
1. Sei $A = \begin{pmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ \frac{\sqrt{3}}{2} & -\frac{1}{2} \end{pmatrix}$.

Wir haben im letzten Kapitel (Seite 91) gesehen, dass $A = [T]$ die Standardmatrix der Spiegelung $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ an der Geraden g mit dem Richtungsvektor $\vec{u}_1 = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix}$ ist.



2. Sei $A = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$.

Wir wissen (vgl. Seite 87), dass diese Matrix die Standardmatrix $A = [T]$ der Orthogonalprojektion $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ auf die x -Achse ist.



Ist also A die Darstellungsmatrix einer bekannten linearen Abbildung, so können die Eigenwerte und Eigenvektoren ohne Rechnung angegeben werden.

Umgekehrt, ist die Darstellungsmatrix einer linearen Abbildung gegeben, so kann der Typ dieser Abbildung durch Berechnen der Eigenwerte und -vektoren bestimmt werden.

Eine lineare Abbildung $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ hat durch die Gleichung $T(\vec{u}) = \lambda\vec{u}$ eindeutig bestimmte Eigenwerte und -vektoren. Es gibt jedoch verschiedene Darstellungsmatrizen, die T beschreiben (abhängig von der Wahl der Basis von $\mathbb{R}^n$). Alle diese Darstellungsmatrizen

haben identische Eigenwerte und -vektoren, nämlich genau diejenigen, die durch $T(\vec{u}) = \lambda \vec{u}$ definiert sind. Wir sehen im folgenden Abschnitt, dass die Darstellungsmatrix von T diagonal gewählt werden kann, falls es eine Basis von $\mathbb{R}^n$ gibt, die aus Eigenvektoren von T besteht.

9.2 Diagonalisierung von Matrizen

Definition Eine $n \times n$ -Matrix A heisst *diagonalisierbar*, wenn es eine invertierbare Matrix P gibt, so dass

$$P A P^{-1} = D$$

eine Diagonalmatrix ist.

Überraschenderweise hängt die Diagonalisierbarkeit einer Matrix mit der Existenz von Eigenvektoren ab.

Satz 9.1 Eine $n \times n$ -Matrix ist genau dann diagonalisierbar, wenn sie n linear unabhängige Eigenvektoren hat.

Hat nämlich eine $n \times n$ -Matrix A n linear unabhängige Eigenvektoren $\vec{u}_1, \dots, \vec{u}_n$ und ist P^{-1} die (invertierbare) Matrix mit den Vektoren $\vec{u}_1, \dots, \vec{u}_n$ als Spalten, dann gilt

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

wobei λ_i der Eigenwert zum Eigenvektor $\vec{u}_i$ ist für $i = 1, \dots, n$. Die Diagonalelemente der Diagonalmatrix D sind also gerade die Eigenwerte der Matrix A !

Beispiel

Ist

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

diagonalisierbar? Wenn ja, bestimme D und P^{-1} .

Auf Seite 97 haben wir gesehen, dass $\vec{u}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ ein Eigenvektor zum Eigenwert $\lambda_1 = 2$ und $\vec{u}_2 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$ ein Eigenvektor zum Eigenwert $\lambda_2 = 3$ ist.

Möchte man nur feststellen, ob eine Matrix diagonalisierbar ist, ohne die diagonalisierende Matrix P auszurechnen, so ist der folgende Satz oft anwendbar.

Satz 9.2 *Hat eine $n \times n$ -Matrix n verschiedene Eigenwerte, so ist sie diagonalisierbar.*

Es gilt nämlich, dass Eigenvektoren zu verschiedenen Eigenwerten linear unabhängig sind. Insbesondere bilden die Basen der verschiedenen Eigenräume eine maximale Menge linear unabhängiger Eigenvektoren der Matrix.

Vorgehen zur Diagonalisierung einer $n \times n$ -Matrix A

- (1) Bestimme alle Eigenwerte von A .
 Hat A n paarweise verschiedene Eigenwerte, dann ist A diagonalisierbar.
 Wenn nicht, dann erkennt man erst in Schritt (2), ob A diagonalisierbar ist oder nicht.
- (2) Bestimme zu jedem Eigenwert eine Basis des Eigenraums (bzw. eine maximale Anzahl linear unabhängiger Eigenvektoren).
 Hat A insgesamt n linear unabhängige Eigenvektoren $\vec{u}_1, \dots, \vec{u}_n$, dann ist A diagonalisierbar.
- (3) Schreibe die Eigenvektoren $\vec{u}_1, \dots, \vec{u}_n$ als Spalten in eine Matrix P^{-1} .
- (4) Das Matrixprodukt $D = PAP^{-1}$ ist eine Diagonalmatrix mit den Diagonalelementen $\lambda_1, \dots, \lambda_n$, wobei λ_i der Eigenwert zum Eigenvektor $\vec{u}_i$ ist.

In Schritt (4) kann man die Diagonalmatrix D direkt hinschreiben, ohne Berechnung von PAP^{-1} . Dass $PAP^{-1} = D$ diagonal ist, begründen wir allgemein auf Seite 103.

Beispiele

1. Ist die Matrix

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

diagonalisierbar? Wenn ja, bestimme D und P^{-1} .

Wir haben auf Seite 98 gesehen, dass A die zwei Eigenwerte $\lambda_1 = 1$ und $\lambda_2 = 2$ mit den folgenden drei Eigenvektoren hat:

$$\vec{u}_1 = \begin{pmatrix} 2 \\ -1 \\ -1 \end{pmatrix}, \vec{u}_2 = \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}, \vec{u}_3 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$$

Genauer ist $\vec{u}_1$ eine Basis des Eigenraums zum Eigenwert 1 und $\vec{u}_2, \vec{u}_3$ bilden eine Basis des Eigenraums zum Eigenwert 2. Wegen der Bemerkung nach Satz 9.2 folgt nun direkt, dass die drei Vektoren $\vec{u}_1, \vec{u}_2, \vec{u}_3$ linear unabhängig sind. Also ist die 3×3 -Matrix A nach Satz 9.1 diagonalisierbar. Es gilt

2. Ist die Matrix

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

diagonalisierbar?

Die Eigenwerte sind also

$$\lambda_1 = 2 \quad \text{und} \quad \lambda_2 = -1.$$

- $\lambda_1 = 2$:

Zu $\lambda_1 = 2$ gibt es genau 1 linear unabhängigen Eigenvektor (oder direkt: der Rang von $\lambda_1 E - A$ ist 2, also ist $\dim(\text{Eigenraum}) = 3 - 2 = 1$).

- $\lambda_2 = -1$: Die Matrix $\lambda_2 E - A = \begin{pmatrix} 0 & 0 & 0 \\ 3 & -3 & 0 \\ 0 & -1 & -3 \end{pmatrix}$ hat den Rang 2.

Die Dimension des Eigenraumes ist daher $3 - 2 = 1$. Eine Basis dieses Eigenraumes besteht also aus einem Eigenvektor.

Insgesamt hat die Matrix A also nur 2 linear unabhängige Eigenvektoren. Nach Satz 9.1 ist A deshalb nicht diagonalisierbar.

Das Untersuchen des Eigenwerts $\lambda_2 = -1$ im letzten Beispiel kann man sich dank des folgenden Satzes sparen.

Satz 9.3 *Ist der Eigenwert λ eine k -fache Nullstelle des charakteristischen Polynoms, dann gibt es zu λ höchstens k linear unabhängige Eigenvektoren.*

Beispiel

Das charakteristische Polynom der Matrix

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

ist

$$\det(\lambda E - A) = \lambda^3 - 1 = (\lambda - 1)(\lambda^2 + \lambda + 1).$$

Der einzige reelle Eigenwert ist $\lambda = 1$ und dazu gibt es nur einen linear unabhängigen Eigenvektor. Diese Matrix ist also nicht diagonalisierbar.

Geometrische Interpretation

Wenn wir die gegebene $n \times n$ -Matrix A als Standardmatrix $A = [T]$ einer linearen Abbildung $T: \mathbb{R}^n \rightarrow \mathbb{R}^n$ auffassen, sind Satz 9.1 und die anschliessende Bemerkung einfach einzusehen.

Die n linear unabhängigen Eigenvektoren $\vec{u}_1, \dots, \vec{u}_n$ bilden eine Basis $\mathcal{B}$ von $\mathbb{R}^n$. Die Matrix P^{-1} bedeutet der Basiswechsel von $\mathcal{B}$ zur Standardbasis und Satz 8.4 sagt nun, dass $D = [T]_{\mathcal{B}}$ die Darstellungsmatrix von T bezüglich der Basis $\mathcal{B}$ ist. In den Spalten von D stehen also die Koordinaten der Bilder von $\vec{u}_1, \dots, \vec{u}_n$ unter T :

$$T(\vec{u}_1) = \lambda_1 \vec{u}_1 = \begin{pmatrix} \lambda_1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}_{\mathcal{B}}, \quad T(\vec{u}_2) = \lambda_2 \vec{u}_2 = \begin{pmatrix} 0 \\ \lambda_2 \\ 0 \\ \vdots \\ 0 \end{pmatrix}_{\mathcal{B}}, \quad \dots, \quad T(\vec{u}_n) = \lambda_n \vec{u}_n = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ \lambda_n \end{pmatrix}_{\mathcal{B}}$$

Wir erhalten genau die Diagonalmatrix D mit den Einträgen $\lambda_1, \dots, \lambda_n$ auf der Diagonalen.

Beispiel

Wir betrachten nochmals die Matrix

$$A = [T] = \begin{pmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ \frac{\sqrt{3}}{2} & -\frac{1}{2} \end{pmatrix},$$

welche die Standardmatrix der Spiegelung T an der Geraden g mit dem Richtungsvektor $\vec{u}_1 = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix}$ ist. Auf Seite 99 haben wir festgestellt, dass A die Eigenvektoren $\vec{u}_1 = \begin{pmatrix} \sqrt{3} \\ 1 \end{pmatrix}$ zum Eigenwert $\lambda_1 = 1$ und $\vec{u}_2 = \begin{pmatrix} -1 \\ \sqrt{3} \end{pmatrix}$ zum Eigenwert $\lambda_2 = -1$ hat. Die beiden Eigenvektoren sind linear unabhängig, die Matrix A ist also diagonalisierbar. Gemäss der Bemerkung nach Satz 9.1 gilt

Tatsächlich erkennen wir (vgl. unsere direkte Rechnung auf den Seiten 90–91) die Diagonalmatrix

$$D = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} = [T]_{\mathcal{B}}$$

als Darstellungsmatrix von T bezüglich der Basis $\mathcal{B} = \{\vec{u}_1, \vec{u}_2\}$ und P^{-1} beschreibt den Basiswechsel von der Basis $\mathcal{B}$ zur Standardbasis.

Sind allgemein A und B zwei (verschiedene) $n \times n$ -Matrizen, so dass

$$B = P A P^{-1}$$

für eine invertierbare Matrix P , dann beschreiben A und B dieselbe lineare Abbildung von $\mathbb{R}^n$ nach $\mathbb{R}^n$, aber bezüglich verschiedenen Basen von $\mathbb{R}^n$. Man nennt in diesem Fall die Matrizen A und B *ähnlich*.

Matrixpotenzen

In Anwendungen müssen oft hohe Potenzen einer quadratischen Matrix berechnet werden. Ist die Matrix diagonalisierbar, dann kann diese Berechnung wie folgt vereinfacht werden.

Sei A eine diagonalisierbare Matrix und P invertierbar mit

$$P A P^{-1} = D$$

diagonal. Dann ist $A = P^{-1} D P$ und

Damit gilt

$$A^3 = A^2 \cdot A = P^{-1} D^2 P P^{-1} D P = P^{-1} D^3 P$$

und so weiter. Für jede ganze Zahl $k \geq 1$ folgt

$$A^k = P^{-1} D^k P.$$

Die rechte Seite ist für grosse k viel einfacher zu berechnen als die linke Seite, denn für das Potenzieren von Diagonalmatrizen gilt:

$$D = \begin{pmatrix} d_1 & & 0 \\ & \ddots & \\ 0 & & d_n \end{pmatrix} \implies D^k = \begin{pmatrix} d_1^k & & 0 \\ & \ddots & \\ 0 & & d_n^k \end{pmatrix}.$$

Beispiel

Wir berechnen A^{20} für $A = \begin{pmatrix} 1 & 1 \\ -2 & 4 \end{pmatrix}$.

Wir haben schon gesehen, dass $P A P^{-1} = \begin{pmatrix} 2 & 0 \\ 0 & 3 \end{pmatrix}$ mit $P^{-1} = \begin{pmatrix} 1 & 1 \\ 1 & 2 \end{pmatrix}$ und $P = \begin{pmatrix} 2 & -1 \\ -1 & 1 \end{pmatrix}$.

Es folgt

und wir erhalten

$$A^{20} = \begin{pmatrix} 2^{21} - 3^{20} & -2^{20} + 3^{20} \\ 2^{21} - 2 \cdot 3^{20} & -2^{20} + 2 \cdot 3^{20} \end{pmatrix} = \begin{pmatrix} -3484687249 & 3485735825 \\ -6971471650 & 6972520226 \end{pmatrix}.$$

Anwendungsbeispiele

Wir haben im letzten Semester Markov-Ketten betrachtet und sind dabei auf Matrixpotenzen gestossen, die wir für grosse Exponenten noch nicht berechnen konnten. Jetzt können wir dies mit Hilfe einer Diagonalisierung tun.

Zur Erinnerung: Bei einer Markov-Kette ist ein System mit n verschiedenen Zuständen gegeben. Für unsere Beispiele hier beschränken wir uns gleich auf $n = 2$. Es gibt also zwei verschiedene Zustände. Sei p_{ji} die Wahrscheinlichkeit, dass das System vom Zustand i in den Zustand j wechselt und sei $A = (p_{ij})$ die entsprechende 2×2 -Matrix. Weiter sei $\vec{v}_k = (x_k, y_k)^T$

der Zustandsvektor nach k Zeitebenen, das heisst, x_k , bzw. y_k ist die Wahrscheinlichkeit, dass das System nach k Zeitebenen im Zustand 1, bzw. 2 ist. Dann gilt

$$\vec{v}_k = A^k \vec{v}_0.$$

1. Beispiel. Wir betrachten nochmals den Wolf vom letzten Semester, der sich abwechselnd in der Nähe von Basel und in der Nähe von Liestal aufhält. Wir wissen, dass wenn der Wolf an einem Tag in Basel ist, er am folgenden Tag stets in Liestal herumstreicht und wenn er in Liestal ist, er am folgenden Tag mit einer Wahrscheinlichkeit von $\frac{1}{4}$ in Basel ist. Hier gilt also

$$\vec{v}_k = A^k \vec{v}_0 \quad \text{für die Matrix } A = \begin{pmatrix} 0 & \frac{1}{4} \\ 1 & \frac{3}{4} \end{pmatrix}.$$

Wenn wir den Wolf heute in Basel sehen und wissen wollen, mit welcher Wahrscheinlichkeit er heute in einem Jahr wieder in Basel ist, dann müssen wir A^{365} berechnen. Dafür und insbesondere auch für allgemeinere Aussagen zu den Aufenthaltswahrscheinlichkeiten nach sehr vielen Tagen (also für sehr grosse k) lohnt es sich, die Matrix A zu diagonalisieren. Tatsächlich ist die Matrix A diagonalisierbar, da sie zwei verschiedene Eigenwerte, $\lambda_1 = 1$ und $\lambda_2 = -\frac{1}{4}$, hat. Mit den zugehörigen Eigenvektoren $\vec{u}_1 = \begin{pmatrix} 1 \\ 4 \end{pmatrix}$ und $\vec{u}_2 = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ erhalten wir die Diagonalisierung

$$D = PAP^{-1} = \begin{pmatrix} 1 & 0 \\ 0 & -\frac{1}{4} \end{pmatrix} \quad \text{mit } P^{-1} = \begin{pmatrix} 1 & -1 \\ 4 & 1 \end{pmatrix} \quad \text{und } P = \frac{1}{5} \begin{pmatrix} 1 & 1 \\ -4 & 1 \end{pmatrix}.$$

Für eine übersichtlichere Schreibweise setzen wir $a = -\frac{1}{4}$. Wir berechnen

$$A^k = P^{-1} D^k P = P^{-1} \begin{pmatrix} 1 & 0 \\ 0 & a^k \end{pmatrix} P = \frac{1}{5} \begin{pmatrix} 1 + 4a^k & 1 - a^k \\ 4 - 4a^k & 4 + a^k \end{pmatrix}$$

und erhalten aus $\begin{pmatrix} x_k \\ y_k \end{pmatrix} = \vec{v}_k = A^k \vec{v}_0 = A^k \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}$ die Aufenthaltswahrscheinlichkeiten des Wolfes

$$x_k = \frac{1}{5} \left((1 + 4a^k)x_0 + (1 - a^k)y_0 \right), \quad y_k = \frac{1}{5} \left((4 - 4a^k)x_0 + (4 + a^k)y_0 \right)$$

nach k Tagen in Basel, bzw. Liestal zu sein. Für sehr grosse k stabilisieren sich die Wahrscheinlichkeiten x_k, y_k . Da $\lim_{k \rightarrow \infty} a^k = \lim_{k \rightarrow \infty} \left(-\frac{1}{4}\right)^k = 0$ und $x_0 + y_0 = 1$, gilt

$$\lim_{k \rightarrow \infty} x_k = \frac{1}{5}(x_0 + y_0) = \frac{1}{5}, \quad \lim_{k \rightarrow \infty} y_k = \frac{1}{5}(4x_0 + 4y_0) = \frac{4}{5}.$$

Nach sehr vielen Tagen hält sich also der Wolf mit einer Wahrscheinlichkeit von $\frac{1}{5} = 0,2$ in Basel und mit einer Wahrscheinlichkeit von $\frac{4}{5} = 0,8$ in Liestal auf (und zwar unabhängig davon, wo wir ihn heute sehen).

2. Beispiel. Wir betrachten einen Organismus, der sich durch Zellteilung verdoppelt. Es gibt zwei verschiedene Typen des Organismus, sagen wir X und Y . Im Reagenzglas beobachten wir, dass aus 100 Individuen vom Typ X nach einem Tag 180 Individuen vom Typ X und 20 Individuen vom Typ Y entstehen, während aus 100 Individuen vom Typ Y nach einem Tag 120 Individuen vom Typ Y und 80 Individuen vom Typ X entstehen. Wie entwickelt sich das Verhältnis der Individuenzahlen vom Typ X bzw. Y ? Die Antwort auf diese Frage kann analog zum vorherigen Beispiel herausgefunden werden. Es stellt sich heraus, dass es nach genügend langer Zucht viermal so viele Individuen vom Typ X wie vom Typ Y gibt (unabhängig von der Anfangspopulation).

Systeme von linearen Differentialgleichungen

Wir sind nun endlich bereit, die Lücke im Kapitel 6 vom letzten Semester zu schliessen. Wir wollen nun sehen, wie man ein beliebiges System von zwei oder mehr linearen Differentialgleichungen (mit konstanten Koeffizienten) mit Hilfe von Eigenwerten und Eigenvektoren lösen kann. Wir zeigen hier die Lösungsmethode an einem Beispiel eines Systems von zwei linearen Differentialgleichungen. Analog können damit Systeme von mehr als zwei Differentialgleichungen gelöst werden.

Beispiel

Gesucht sind zwei reelle Funktionen $y_1 = y_1(x)$ und $y_2 = y_2(x)$, so dass

$$\begin{aligned}y_1' &= 3y_1 - 2y_2 \\y_2' &= 2y_1 - 2y_2 .\end{aligned}$$

Wir schreiben dieses System mit Matrizen als

$$\begin{pmatrix} y_1' \\ y_2' \end{pmatrix} = A \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} \quad \text{mit } A = \begin{pmatrix} 3 & -2 \\ 2 & -2 \end{pmatrix} .$$

Die Differentialgleichung $y' = ay$ hat die Lösung $y(x) = ce^{ax}$. Wir versuchen deshalb den Ansatz

$$y_1(x) = c_1 e^{\lambda x}, \quad y_2(x) = c_2 e^{\lambda x}, \quad \text{das heisst} \quad \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = e^{\lambda x} \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} .$$

Damit gilt

$$\begin{pmatrix} y_1' \\ y_2' \end{pmatrix} = \lambda e^{\lambda x} \begin{pmatrix} c_1 \\ c_2 \end{pmatrix}, \quad A \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = e^{\lambda x} A \begin{pmatrix} c_1 \\ c_2 \end{pmatrix} .$$

Das System der zwei Differentialgleichungen ist demnach erfüllt, wenn

Gesucht sind also die Eigenwerte und Eigenvektoren der Matrix A !

Die charakteristische Gleichung von A ist

Zu den Eigenwerten $\lambda_1 = -1$ und $\lambda_2 = 2$ brauchen wir je einen Eigenvektor:

Wir erhalten die Lösungen

Die allgemeine Lösung besteht nun aus allen möglichen Linearkombinationen dieser beiden Lösungen (man kann zeigen, dass der Lösungsraum ein Vektorraum der Dimension 2 ist):

$$\begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \alpha e^{-x} \begin{pmatrix} 1 \\ 2 \end{pmatrix} + \beta e^{2x} \begin{pmatrix} 2 \\ 1 \end{pmatrix}$$

für $\alpha, \beta \in \mathbb{R}$, das heisst,

$$\begin{aligned} y_1(x) &= \alpha e^{-x} + 2\beta e^{2x} \\ y_2(x) &= 2\alpha e^{-x} + \beta e^{2x} . \end{aligned}$$

9.3 Symmetrische Matrizen

Wir haben gesehen, dass das Diagonalisieren einer Matrix einem Basiswechsel der zugehörigen linearen Abbildung entspricht. Speziell praktisch sind natürlich Orthonormalbasen, das heisst Basisvektoren, die zueinander orthogonal sind und alle die Länge 1 haben. Schreibt man die Basisvektoren einer Orthonormalbasis in eine Matrix, dann erhält man eine orthogonale Matrix (vgl. Kapitel 9, Seite 162 vom letzten Semester). Es stellt sich also die Frage, ob und wann die diagonalisierende Matrix P orthogonal gewählt werden kann.

Nehmen wir an, dass P orthogonal ist, das heisst, $P^{-1} = P^T$. Aus $D = PAP^{-1}$ erhalten wir für die Matrix A die Gleichung

$$A = P^{-1}DP = P^TDP .$$

Daraus folgt, dass die Matrix A symmetrisch ist, denn

$$A^T = (P^TDP)^T = P^T D^T (P^T)^T = P^TDP = A .$$

Die Matrix P kann also nur dann orthogonal sein, wenn A symmetrisch ist. Tatsächlich ist diese Bedingung an A nicht nur notwendig, sondern auch hinreichend: Jede symmetrische Matrix kann durch eine orthogonale Matrix diagonalisiert werden.

Satz 9.4 *Jede symmetrische Matrix ist diagonalisierbar, und zwar kann die diagonalisierende Matrix P stets orthogonal gewählt werden.*

Wählt man die diagonalisierende Matrix P orthogonal, dann erhält man wegen $P^{-1} = P^T$ die Diagonalisierung

$$PAP^T = D.$$

Beispiel

Gegeben ist die (symmetrische) Matrix

$$A = \begin{pmatrix} 5 & -2 \\ -2 & 8 \end{pmatrix}.$$

Die charakteristische Gleichung ist

$$\det(\lambda E - A) = \det \begin{pmatrix} \lambda - 5 & 2 \\ 2 & \lambda - 8 \end{pmatrix} = \lambda^2 - 13\lambda + 36 = (\lambda - 4)(\lambda - 9) = 0.$$

Wir finden die folgenden Eigenvektoren:

Die beiden Eigenvektoren $\vec{v}_1$ und $\vec{v}_2$ sind orthogonal. Um eine orthogonale Matrix P zu erhalten, müssen wir $\vec{v}_1$ und $\vec{v}_2$ auf die Länge 1 normieren. Eine Orthonormalbasis von Eigenvektoren bilden also die Vektoren

$$\vec{u}_1 = \frac{1}{\sqrt{5}} \begin{pmatrix} 2 \\ 1 \end{pmatrix}, \quad \vec{u}_2 = \frac{1}{\sqrt{5}} \begin{pmatrix} -1 \\ 2 \end{pmatrix}.$$

Wir erhalten die Diagonalisierung

Satz 9.5 *Bei einer symmetrischen Matrix sind die Eigenvektoren zu verschiedenen Eigenwerten orthogonal.*

Innerhalb eines Eigenraumes muss man eine Orthonormalbasis von Eigenvektoren wählen, damit die Matrix P orthogonal ist.

Hauptachsentransformation

Die Diagonalisierung einer symmetrischen Matrix mit einer orthogonalen Matrix nennt man *Hauptachsentransformation*. Dieser Name wird anhand des folgenden Beispiels klar.

Beispiel

Wie sieht die Kurve C in $\mathbb{R}^2$ mit der Gleichung

$$5x^2 - 4xy + 8y^2 - 36 = 0$$

aus? Wir können diese Gleichung schreiben als

$$\vec{x}^T A \vec{x} - 36 = (x, y) \begin{pmatrix} 5 & -2 \\ -2 & 8 \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} - 36 = 0 \quad \text{mit } A = \begin{pmatrix} 5 & -2 \\ -2 & 8 \end{pmatrix}, \quad \vec{x} = \begin{pmatrix} x \\ y \end{pmatrix}.$$

Dies ist die Matrix A vom vorhergehenden Beispiel.

Wir transformieren nun die Koordinaten, das heisst, wir setzen

$$\vec{y} = P \vec{x} \quad \text{mit } P = \frac{1}{\sqrt{5}} \begin{pmatrix} 2 & 1 \\ -1 & 2 \end{pmatrix}, \quad \vec{y} = \begin{pmatrix} x' \\ y' \end{pmatrix}.$$

Dabei ist P der Basiswechsel von der Standardbasis zur Basis $\mathcal{B} = \{\vec{u}_1, \vec{u}_2\}$ und demnach ist $\vec{y}$ der Koordinatenvektor bezüglich der Basis $\mathcal{B}$. Da P die Matrix einer Drehung ist (denn P ist orthogonal und $\det(P) = 1$), drehen wir also das Koordinatensystem.

Mit $\vec{x} = P^{-1} \vec{y} = P^T \vec{y}$ folgt

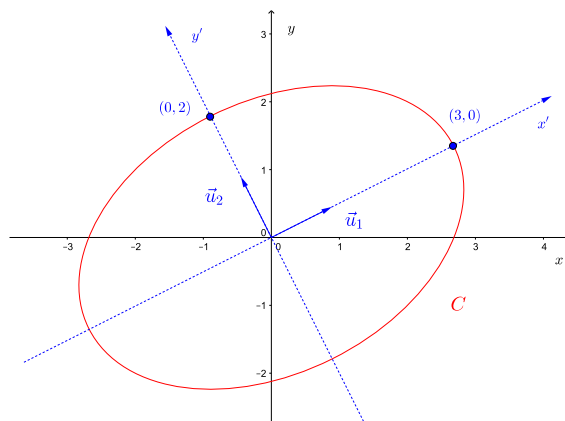
Das heisst, in den neuen Koordinaten (bzgl. der Basis $\mathcal{B}$) hat die Kurve C die Gleichung

$$\vec{y}^T D \vec{y} - 36 = (x', y') \begin{pmatrix} 4 & 0 \\ 0 & 9 \end{pmatrix} \begin{pmatrix} x' \\ y' \end{pmatrix} - 36 = 4x'^2 + 9y'^2 - 36 = 0,$$

das heisst,

$$\frac{x'^2}{9} + \frac{y'^2}{4} = 1.$$

Die Kurve C ist also eine Ellipse, deren Hauptachsen in die Richtung der Eigenvektoren von A zeigen.



Quadratische Formen

Der im letzten Beispiel aufgetretene Ausdruck

$$\vec{x}^T A \vec{x} = 5x^2 - 4xy + 8y^2$$

ist ein Beispiel einer sogenannten quadratischen Form.

Definition Sei A eine symmetrische $n \times n$ -Matrix und $\vec{x} \in \mathbb{R}^n$. Dann heisst

$$\vec{x}^T A \vec{x} = \vec{x} \cdot (A \vec{x})$$

quadratische Form von A .

Da A symmetrisch ist, kann die quadratische Form auch geschrieben werden als

$$\vec{x}^T A \vec{x} = (A \vec{x}) \cdot \vec{x},$$

denn $\vec{x}^T A \vec{x} = \vec{x}^T A^T \vec{x} = (A \vec{x})^T \vec{x} = (A \vec{x}) \cdot \vec{x}$.

Bei der Untersuchung von lokalen Extrema von Funktionen in mehreren Variablen im nächsten Kapitel müssen wir wissen, wann eine quadratische Form nur positive (bzw. negative) Werte annimmt.

Definition Sei A eine symmetrische $n \times n$ -Matrix.

- A heisst *positiv definit*, wenn

$$\vec{x}^T A \vec{x} > 0 \quad \text{für alle } \vec{x} \neq \vec{0} \text{ in } \mathbb{R}^n.$$

- A heisst *negativ definit*, wenn

$$\vec{x}^T A \vec{x} < 0 \quad \text{für alle } \vec{x} \neq \vec{0} \text{ in } \mathbb{R}^n.$$

- A heisst *indefinit*, wenn es $\vec{x} \neq \vec{0}$ mit $\vec{x}^T A \vec{x} > 0$ und $\vec{y} \neq \vec{0}$ mit $\vec{y}^T A \vec{y} < 0$ gibt.

Wir geben im Folgenden zwei Kriterien an, wann eine Matrix positiv (bzw. negativ) definit ist. Wir formulieren diese nur für 2×2 -Matrizen, da wir sie für solche Matrizen im nächsten Kapitel brauchen werden. Beide Kriterien können jedoch auf $n \times n$ -Matrizen verallgemeinert werden.

Beispiel

Sei

$$\vec{x}^T A \vec{x} = 5x^2 - 4xy + 8y^2 \quad \text{mit } A = \begin{pmatrix} 5 & -2 \\ -2 & 8 \end{pmatrix}$$

die quadratische Form vom letzten Beispiel. In den neuen Koordinaten $\begin{pmatrix} x' \\ y' \end{pmatrix} = \vec{y} = P \vec{x}$ (bezüglich der Basis von Eigenvektoren) lautet die quadratische Form

$$\vec{x}^T A \vec{x} = 4x'^2 + 9y'^2.$$

In dieser Darstellung ist klar ersichtlich, dass A positiv definit ist.

Eine solche Koordinatentransformation kann für jede symmetrische 2×2 -Matrix (bzw. $n \times n$ -Matrix) durchgeführt werden. Sei

$$A = \begin{pmatrix} a & b \\ b & c \end{pmatrix}.$$

Dann gilt für $\vec{x} = \begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$, dass

$$\vec{x}^T A \vec{x} = (x, y) A \begin{pmatrix} x \\ y \end{pmatrix} = ax^2 + 2bxy + cy^2.$$

Da A symmetrisch ist, ist A diagonalisierbar mit den Eigenwerten λ_1, λ_2 . Es folgt (wie im Beispiel), dass

$$\vec{x}^T A \vec{x} = \vec{y}^T \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \vec{y} = \lambda_1 x'^2 + \lambda_2 y'^2$$

bezüglich der neuen Koordinaten $\vec{y} = \begin{pmatrix} x' \\ y' \end{pmatrix}$. Die Definitheit von A hängt also von den Eigenwerten von A ab.

Satz 9.6 Seien λ_1, λ_2 die Eigenwerte der symmetrischen 2×2 -Matrix A . Dann gilt:

- A positiv definit $\iff \lambda_1 > 0, \lambda_2 > 0$
- A negativ definit $\iff \lambda_1 < 0, \lambda_2 < 0$
- A indefinit $\iff \lambda_1 > 0, \lambda_2 < 0$ oder $\lambda_1 < 0, \lambda_2 > 0$

Insbesondere ist A indefinit $\iff \det(A) < 0$. Weiter gilt:

- A positiv definit $\iff \det(A) > 0, a > 0$
- A negativ definit $\iff \det(A) > 0, a < 0$

Beispiel

Gegeben sind

$$B = \begin{pmatrix} -1 & 2 \\ 2 & -6 \end{pmatrix} \quad \text{mit} \quad \vec{x}^T B \vec{x} = -x^2 + 4xy - 6y^2,$$

$$C = \begin{pmatrix} 1 & 2 \\ 2 & -6 \end{pmatrix} \quad \text{mit} \quad \vec{x}^T C \vec{x} = x^2 + 4xy - 6y^2,$$

$$D = \begin{pmatrix} 6 & 0 \\ 0 & -1 \end{pmatrix} \quad \text{mit} \quad \vec{x}^T D \vec{x} = 6x^2 - y^2.$$

9.4 Komplexe Matrizen

Anstelle von reellen $n \times n$ -Matrizen können wir auch komplexe $n \times n$ -Matrizen betrachten, das heisst, Matrizen mit Einträgen in $\mathbb{C}$. Eine komplexe $n \times n$ -Matrix beschreibt dann eine lineare Abbildung von $\mathbb{C}^n$ nach $\mathbb{C}^n$.

In $\mathbb{R}^n$ haben wir das Skalarprodukt $\vec{x} \cdot \vec{y} = \vec{x}^T \vec{y}$. Für $\vec{z}, \vec{w}$ in $\mathbb{C}^n$ definieren wir das *Skalarprodukt*

$$\vec{z} \cdot \vec{w} = z_1 \bar{w}_1 + \cdots + z_n \bar{w}_n$$

wobei $\bar{w}_1, \dots, \bar{w}_n$ die zu $w_1, \dots, w_n$ konjugiert komplexen Zahlen sind.

Damit ist $\vec{z} \cdot \vec{z}$ eine nichtnegative reelle Zahl und die *Länge* eines Vektors $\vec{z} \in \mathbb{C}^n$ kann definiert werden durch

$$\|\vec{z}\| = \sqrt{\vec{z} \cdot \vec{z}} = \sqrt{|z_1|^2 + \cdots + |z_n|^2}.$$

Die Begriffe *orthogonale* Vektoren und *Orthonormalbasis* lassen sich ohne Änderung auf den (komplexen) Vektorraum $\mathbb{C}^n$ übertragen.

Beispiel

Sind die Vektoren

$$\vec{z} = \begin{pmatrix} i \\ 1 \end{pmatrix} \quad \text{und} \quad \vec{w} = \begin{pmatrix} 1 \\ i \end{pmatrix}$$

orthogonal?

Weiter ist $\|\vec{z}\| = \sqrt{|i|^2 + |1|^2} = \sqrt{2}$.

Beim Diagonalisieren von reellen Matrizen spielen symmetrische und orthogonale Matrizen eine wichtige Rolle. Analog zur Transponierten definiert man für eine komplexe Matrix A die *konjugiert Transponierte*

$$A^* = \overline{A}^T,$$

wobei $\overline{A}$ durch Konjugieren der einzelnen Elemente von A entsteht und $\overline{A}^T$ die Transponierte von $\overline{A}$ ist.

Beispiel

$$A = \begin{pmatrix} 1 & 2+i \\ 4i & 1-i \end{pmatrix} \quad \Longrightarrow \quad A^* =$$

Definition Sei A eine komplexe $n \times n$ -Matrix.

- A heisst *hermitesch*, falls $A = A^*$.
- A heisst *unitär*, falls $A^{-1} = A^*$.

Hermiteische Matrizen sind also das komplexe Analogon der symmetrischen Matrizen. Unitäre Matrizen übernehmen die Rolle der orthogonalen Matrizen. Dabei ist (wie im Reellen) eine komplexe Matrix unitär genau dann, wenn die Spaltenvektoren eine Orthonormalbasis von $\mathbb{C}^n$ bilden.

Satz 9.7 *Jede hermitesche Matrix ist diagonalisierbar, und zwar kann die diagonalisierende Matrix stets unitär gewählt werden.*

Bei hermiteschen Matrizen sind die Eigenvektoren zu verschiedenen Eigenwerten orthogonal, genau wie bei symmetrischen Matrizen.

Im Reellen sind die symmetrischen Matrizen die einzigen, welche orthogonal diagonalisiert werden können. Im Komplexen gibt es neben den hermiteschen Matrizen noch andere, welche unitär diagonalisiert werden können.

Satz 9.8 *Die Eigenwerte einer hermiteschen Matrix sind reell.*

Beispiel

Wir betrachten die hermitesche Matrix

$$A = \begin{pmatrix} 2 & 1+i \\ 1-i & 3 \end{pmatrix}.$$

Die charakteristische Gleichung lautet:

Die Eigenwerte sind also $\lambda_1 = 1$ und $\lambda_2 = 4$.

Die zugehörigen Eigenvektoren finden wir wie im Reellen (lineares Gleichungssystem lösen). Zu λ_1 und λ_2 finden wir zum Beispiel die Eigenvektoren

$$\vec{v}_1 = \begin{pmatrix} 1+i \\ -1 \end{pmatrix} \quad \text{und} \quad \vec{v}_2 = \begin{pmatrix} 1+i \\ 2 \end{pmatrix}.$$

Da sie zu verschiedenen Eigenwerten gehören, sind sie orthogonal. Um eine unitäre diagonalisierende Matrix P^{-1} zu bekommen, müssen wir die beiden Vektoren noch auf die Länge 1 normieren.

Damit erhalten wir

$$PAP^{-1} = PAP^* = \begin{pmatrix} 1 & 0 \\ 0 & 4 \end{pmatrix}$$

mit der unitären Matrix

$$P^{-1} = P^* = \begin{pmatrix} \frac{1+i}{\sqrt{3}} & \frac{1+i}{\sqrt{6}} \\ \frac{-1}{\sqrt{3}} & \frac{2}{\sqrt{6}} \end{pmatrix}, \quad \text{bzw.} \quad P = \begin{pmatrix} \frac{1-i}{\sqrt{3}} & \frac{-1}{\sqrt{3}} \\ \frac{1-i}{\sqrt{6}} & \frac{2}{\sqrt{6}} \end{pmatrix}.$$

10 Differentialrechnung für Funktionen in mehreren Variablen

Viele Funktionen in den Naturwissenschaften hängen von mehreren Variablen ab. In diesem Kapitel behandeln wir deshalb Methoden zur Untersuchung von Funktionen in mehreren Variablen.

Im letzten Semester haben wir Funktionen $f : \mathbb{R} \rightarrow \mathbb{R}$ betrachtet, in den Kapiteln 8 und 9 in diesem Semester (lineare) Abbildungen $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$. In diesem Kapitel geht es nun hauptsächlich um Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$.

Sei $D \subset \mathbb{R}^2$ eine Teilmenge. Wir wissen schon von Kapitel 8 (Seite 86), dass eine (*reellwertige*) *Funktion* $f : D \rightarrow \mathbb{R}$ *von zwei reellen Variablen* eine Vorschrift ist, die jedem Punkt $(x, y) \in D$ eine reelle Zahl $z = f(x, y)$ zuordnet,

$$\begin{aligned} f : D &\longrightarrow \mathbb{R} \\ (x, y) &\longmapsto z = f(x, y). \end{aligned}$$

Ist D eine Teilmenge von $\mathbb{R}^3$ oder allgemeiner $\mathbb{R}^n$, dann definiert die Zuordnung

$$\begin{aligned} f : D &\longrightarrow \mathbb{R} \\ (x_1, x_2, \dots, x_n) &\longmapsto f(x_1, x_2, \dots, x_n) \end{aligned}$$

eine (*reellwertige*) *Funktion in n Variablen*. Beispiele haben wir schon in Kapitel 8 gesehen.

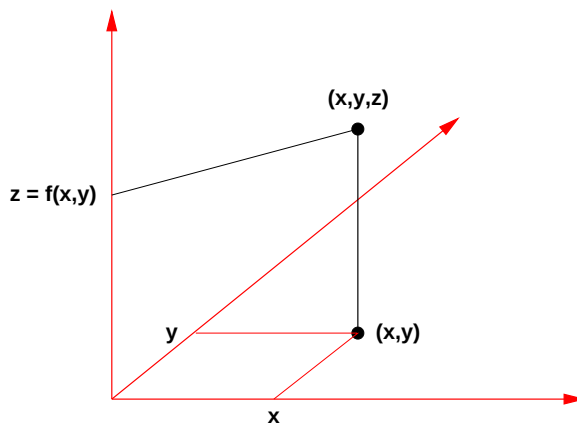
Wir werden in diesem Kapitel vor allem (reellwertige) Funktionen mit Definitionsbereich in $\mathbb{R}^2$ untersuchen. Die meisten Begriffe und Resultate lassen sich problemlos auf den Fall von drei und mehr Variablen übertragen. Im Gegensatz zum allgemeinen Fall hilft uns bei zwei Variablen jedoch die geometrische Anschauung.

10.1 Graphische Darstellung

Sei $D \subset \mathbb{R}^2$ und $f : D \rightarrow \mathbb{R}$. Analog zu reellen Funktionen können wir die Funktion f mit Hilfe ihres Graphen veranschaulichen. Der Graph von f ist definiert durch

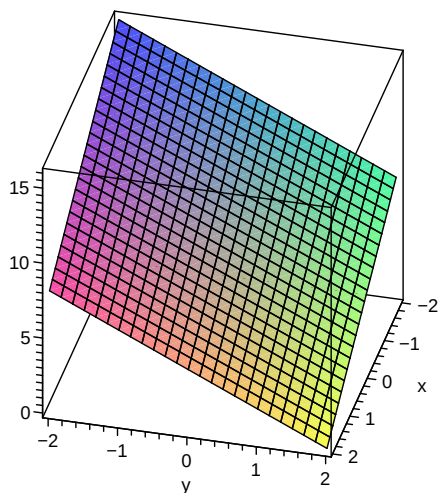
$$\text{Graph}(f) = \{ (x, y, f(x, y)) \mid (x, y) \in D \}.$$

Wir errichten eine Strecke der Länge $z = f(x, y)$ über jedem Punkt $(x, y) \in D$ (bzw. unter $(x, y) \in D$ falls $z < 0$). Die Endpunkte aller dieser Strecken bilden eine Fläche im Raum, welche der Graph von f ist.

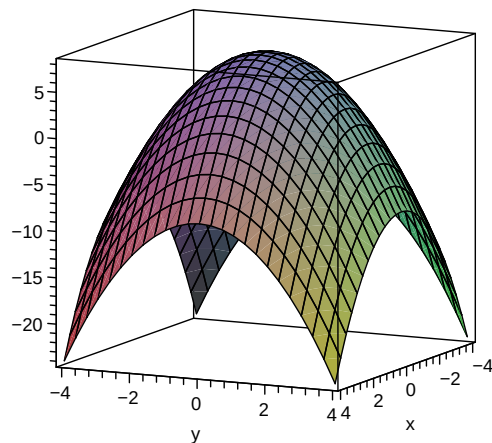


Beispiele

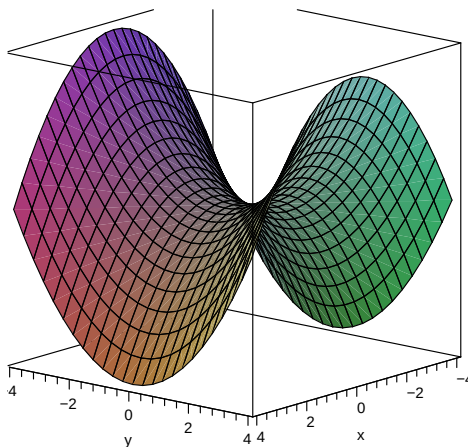
$$z = f(x, y) = 8 - 2x - y$$



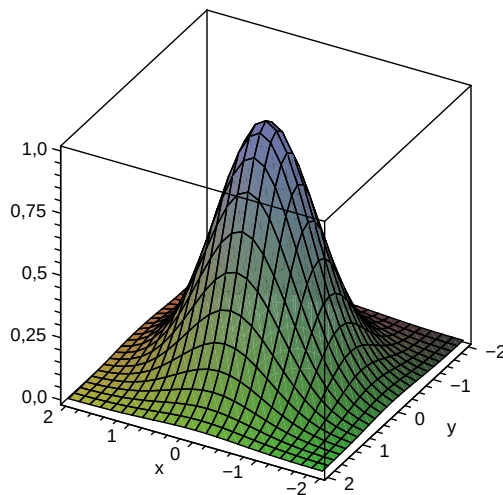
$$z = f(x, y) = 8 - x^2 - y^2$$



$$z = f(x, y) = y^2 - x^2$$



$$z = f(x, y) = e^{-(x^2+y^2)}$$



Wir können die Funktion $f : D \rightarrow \mathbb{R}$ in zwei Variablen auch durch Niveaulinien (wie die Höhenkurven auf Landkarten) veranschaulichen. Wir schneiden den Graphen von f mit horizontalen Ebenen, das heisst, parallel zur xy -Ebene in einer bestimmten Höhe $z = c$. Die Schnittkurve projizieren wir senkrecht in die xy -Ebene. Die *Niveaulinie* für $z = c$ ist also gegeben durch

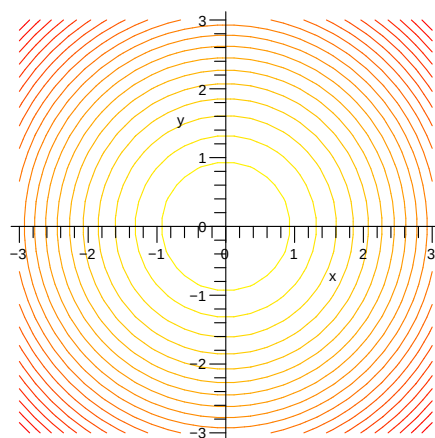
$$N_c = \{ (x, y) \in D \mid f(x, y) = c \} \subset \mathbb{R}^2.$$

Beispiele

1. $z = f(x, y) = 8 - x^2 - y^2$

Als Niveaulinien erhalten wir eine Familie konzentrischer Kreise:

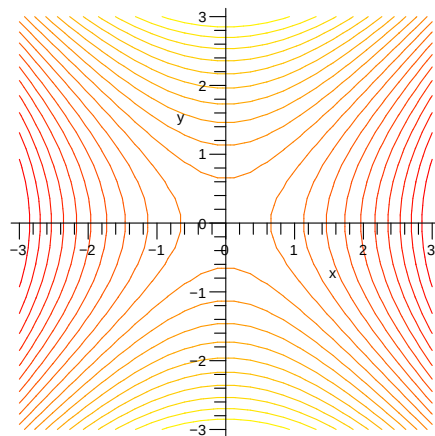
$$\begin{aligned} z > 8 & : \quad \text{keine Lösung} \\ 8 = z = 8 - x^2 - y^2 & \implies x^2 + y^2 = 0 \\ 4 = z = 8 - x^2 - y^2 & \implies x^2 + y^2 = 4 \\ 0 = z = 8 - x^2 - y^2 & \implies x^2 + y^2 = 8 \\ \dots & \quad \dots \end{aligned}$$



2. $z = f(x, y) = y^2 - x^2$

Als Niveaulinien erhalten wir eine Familie von Hyperbeln:

$$\begin{aligned} 0 = z = y^2 - x^2 & \implies y^2 - x^2 = 0 \text{ oder } y = \pm x \\ -4 = z = y^2 - x^2 & \implies y^2 - x^2 = -4 \\ 4 = z = y^2 - x^2 & \implies y^2 - x^2 = 4 \\ \dots & \quad \dots \end{aligned}$$



10.2 Partielle Ableitungen und Tangentialebenen

Eines unserer Ziele ist, Extremalstellen von Funktionen in mehreren Variablen zu finden und zu untersuchen. Wie für reelle Funktionen brauchen wir dazu Ableitungen.

Partielle Ableitungen

Sei $D \subset \mathbb{R}^2$ und $f : D \rightarrow \mathbb{R}$ eine Funktion.

Definition Die *partiellen Ableitungen* von f im Punkt (x_0, y_0) sind wie folgt definiert.

$$f_x(x_0, y_0) = \frac{\partial f}{\partial x}(x_0, y_0) = \lim_{h \rightarrow 0} \frac{f(x_0 + h, y_0) - f(x_0, y_0)}{h} = \lim_{x \rightarrow x_0} \frac{f(x, y_0) - f(x_0, y_0)}{x - x_0}$$

ist die *partielle Ableitung nach x* und

$$f_y(x_0, y_0) = \frac{\partial f}{\partial y}(x_0, y_0) = \lim_{h \rightarrow 0} \frac{f(x_0, y_0 + h) - f(x_0, y_0)}{h} = \lim_{y \rightarrow y_0} \frac{f(x_0, y) - f(x_0, y_0)}{y - y_0}$$

ist die *partielle Ableitung nach y* .

Beispiele

1. $f(x, y) = x^2 + 5xy + 3y^2 + 13$

2. $f(x, y) = x^2 e^{2y} + \ln(x)$

Für die partielle Ableitung f_x fixiert man also die Variable y (man behandelt y wie eine fixe reelle Zahl) und leitet wie gewohnt nach x ab. Analog für f_y .

Wie für reelle Funktionen brauchen wir zusätzlich höhere Ableitungen.

Definition Die *partiellen Ableitungen zweiter Ordnung* sind definiert durch

$$f_{xx} = \frac{\partial^2 f}{\partial x^2} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) \quad f_{yy} = \frac{\partial^2 f}{\partial y^2} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial y} \right)$$

und

$$f_{xy} = \frac{\partial^2 f}{\partial x \partial y} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial x} \right) \quad f_{yx} = \frac{\partial^2 f}{\partial y \partial x} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial y} \right).$$

Beispiel

$$f(x, y) = x^2 + 5xy + 3y^2 + 13 \quad \text{mit} \quad f_x(x, y) = 2x + 5y \quad \text{und} \quad f_y(x, y) = 5x + 6y$$

Satz 10.1 Sind die partiellen Ableitungen f_{xy} und f_{yx} stetige Funktionen, dann gilt

$$f_{xy} = f_{yx}.$$

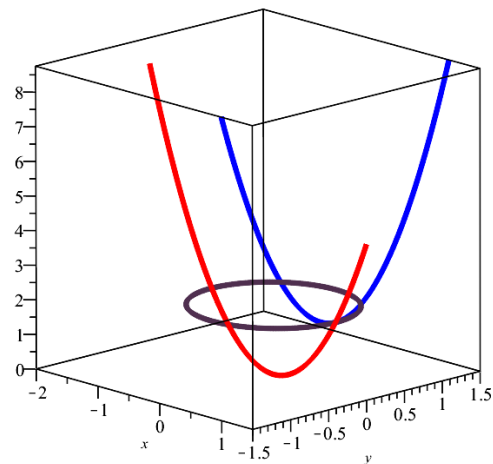
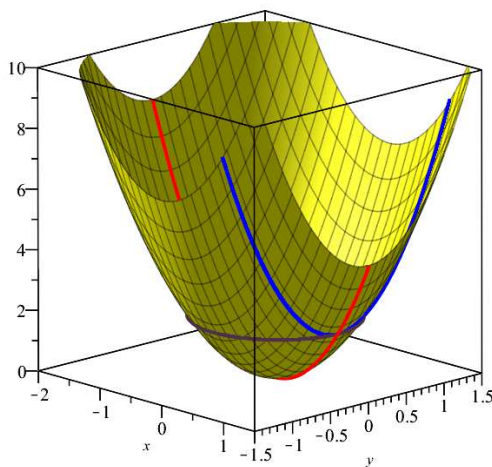
Die partiellen Ableitungen haben die folgende geometrische Bedeutung. Sei $z = f(x, y)$ eine Funktion mit Definitionsbereich D und $P = (x_0, y_0, z_0)$ mit $z_0 = f(x_0, y_0)$ ein Punkt auf dem Graphen von f . Durch diesen Punkt gibt es drei spezielle Kurven auf dem Graphen:

- xz -Kurve durch P : $\{ (x, y_0, z) \mid z = f(x, y_0) \text{ und } (x, y_0) \in D \}$
- yz -Kurve durch P : $\{ (x_0, y, z) \mid z = f(x_0, y) \text{ und } (x_0, y) \in D \}$
- xy -Kurve durch P : $\{ (x, y, z_0) \mid z_0 = f(x, y) \text{ und } (x, y) \in D \}$

Beispiel

Sei $z = f(x, y) = 2x^2 + 3y^2$ mit $D = \mathbb{R}^2$ und $P = (1, 0, 2)$.

- xz -Kurve durch P :
- yz -Kurve durch P :
- xy -Kurve durch P : $\{ (x, y, 2) \mid 2 = f(x, y) = 2x^2 + 3y^2 \text{ und } (x, y) \in \mathbb{R}^2 \}$



Bedeutung der partiellen Ableitungen im Punkt $P = (x_0, y_0, z_0)$:

- $f_x(x_0, y_0)$: Steigung der xz -Kurve in P
 - $f_x(x_0, y_0) = 0 \implies$ die xz -Kurve hat in P eine horizontale Tangente
 - $f_{xx}(x_0, y_0) < 0 \implies$ die xz -Kurve beschreibt eine Rechtskurve* in der Nähe von P
 - $f_{xx}(x_0, y_0) > 0 \implies$ die xz -Kurve beschreibt eine Linkskurve* in der Nähe von P
- $f_y(x_0, y_0)$: Steigung der yz -Kurve in P
 - $f_y(x_0, y_0) = 0 \implies$ die yz -Kurve hat in P eine horizontale Tangente
 - $f_{yy}(x_0, y_0) < 0 \implies$ die yz -Kurve beschreibt eine Rechtskurve* in der Nähe von P
 - $f_{yy}(x_0, y_0) > 0 \implies$ die yz -Kurve beschreibt eine Linkskurve* in der Nähe von P

* betrachtet als Funktion in einer Variablen $z = \tilde{f}(x)$, bzw. $z = \bar{f}(y)$.

Beispiel

Sei $z = f(x, y) = 2x^2 + 3y^2$ und $P = (1, 0, 2)$.

Tangentialebenen

Wir haben im letzten Semester (Kapitel 4, Abschnitt 4) gesehen, dass eine differenzierbare Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ in der Nähe eines Punktes $(x_0, f(x_0))$ durch eine Gerade, nämlich durch die Tangente an den Graphen von f , approximiert werden kann:

$$f(x) \approx f(x_0) + f'(x_0)(x - x_0)$$

Wir wollen nun analog eine Funktion $f : D \rightarrow \mathbb{R}$ in zwei Variablen in der Nähe des Punktes $P = (x_0, y_0, z_0)$, mit $z_0 = f(x_0, y_0)$, linear approximieren. Da der Graph von f eine Fläche ist, suchen wir eine Ebene

$$z = T(x, y) = c + a(x - x_0) + b(y - y_0),$$

welche

1. den Graphen in P berührt,
2. in P die gleiche Steigung wie f in x -Richtung hat,
3. in P die gleiche Steigung wie f in y -Richtung hat.

Es gibt also genau eine solche Ebene. Wir nennen sie *Tangentialebene*.

Satz 10.2 Die Tangentialebene an den Graphen der Funktion $z = f(x, y)$ an der Stelle (x_0, y_0) ist gegeben durch

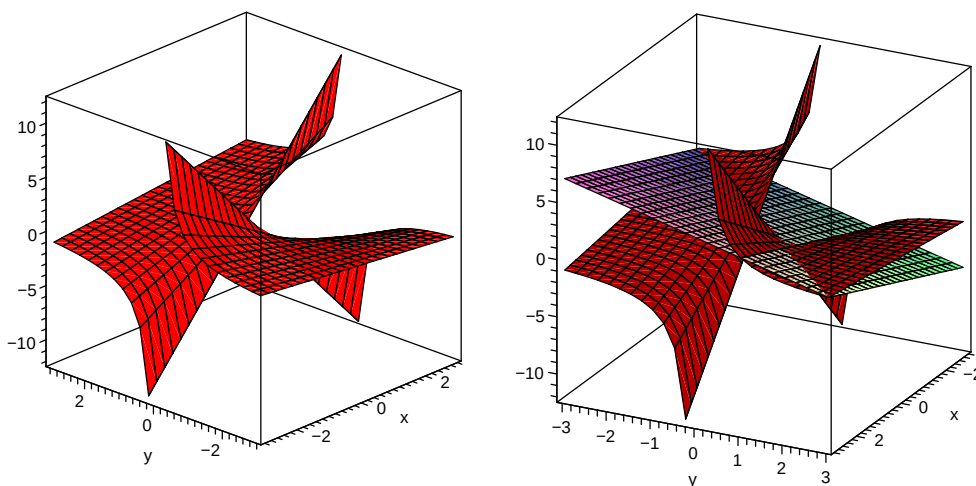
$$z = T(x, y) = f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) .$$

Beispiel

Gesucht ist die Tangentialebene an den Graphen der Funktion

$$z = f(x, y) = \frac{x}{y}$$

an der Stelle $(x_0, y_0) = (1, 1)$.



Wie für reelle Funktionen kann nun eine (ev. komplizierte) Funktion in zwei Variablen in der Nähe eines Punktes durch ihre Tangentialebene in diesem Punkt approximiert werden,

$$f(x, y) \approx T(x, y) = f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0)$$

für (x, y) in der Nähe von (x_0, y_0) .

Beispiel

Für f wie im vorhergehenden Beispiel bestimme man eine Näherung für $f(1,02; 0,94)$.

Das Differential

In Analogie zu Satz 4.10 vom letzten Semester für reelle Funktionen heisst eine Funktion $f : D \rightarrow \mathbb{R}$ in zwei Variablen (*total*) differenzierbar in $(x_0, y_0) \in D$, wenn

$$f(x, y) = f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) + r(x, y)$$

mit

$$\frac{r(x, y)}{\sqrt{(x - x_0)^2 + (y - y_0)^2}} \rightarrow 0 \quad \text{für} \quad (x, y) \rightarrow (x_0, y_0).$$

Eine in (x_0, y_0) differenzierbare Funktion ist also in (x_0, y_0) sehr gut durch die Tangentialebene approximierbar (für (x, y) gegen (x_0, y_0) geht der Restterm $r(x, y)$ schneller gegen 0 als der Abstand zwischen (x, y) und (x_0, y_0)). Zu beachten ist, dass alleine aus der Existenz der partiellen Ableitungen nicht folgt, dass f differenzierbar ist. Hingegen ist f differenzierbar, wenn die partiellen Ableitungen f_x und f_y stetige Funktionen sind.

Benutzen wir die Tangentialebene in (x_0, y_0) als Näherung von f in der Nähe von (x_0, y_0) , dann erhalten wir eine Näherung für die Änderung Δf von f , wenn sich x_0 um den kleinen Wert $\Delta x = dx$ und y_0 um den kleinen Wert $\Delta y = dy$ ändert,

$$\Delta f = f(x_0 + dx, y_0 + dy) - f(x_0, y_0) \approx f_x(x_0, y_0) dx + f_y(x_0, y_0) dy.$$

Definition Man nennt

$$df(x_0, y_0) = f_x(x_0, y_0) dx + f_y(x_0, y_0) dy$$

oder kurz

$$df = f_x dx + f_y dy$$

das (*totale*) Differential von f .

Im Fall einer Variablen gilt $f(x) = f(x_0) + f'(x_0)(x - x_0) + r(x)$ und damit ist $df = f'(x)dx$. Das Differential ist also die Verallgemeinerung der Ableitung auf mehrere Variablen.

Anwendung auf Fehlerabschätzungen

Sei $f(x, y)$ eine Funktion von zwei Messgrössen und dx, dy die Messfehler. Wie gross ist die Abweichung $\Delta f = f(x + dx, y + dy) - f(x, y)$? Ist f differenzierbar, dann können wir die Näherung

$$\Delta f \approx df = f_x dx + f_y dy$$

verwenden.

Beispiel

Sei $f(x, y) = xy$.

Messen wir also beispielsweise die Seitenlängen x , y eines Rechtecks mit je einem relativen Fehler von 1 %, dann ist der relative Fehler des aus x und y berechneten Flächeninhalts des Rechtecks gegeben durch

Kettenregel

Ist $x = x(t)$ und $y = f(x) = f(x(t))$ eine Funktion in einer Variablen, dann gilt die Kettenregel

$$y'(t) = \frac{df}{dt} = f'(x(t)) \cdot x'(t) = f'(x(t)) \frac{dx}{dt}.$$

Diese Regel kann auf zwei (und mehr) Variablen verallgemeinert werden.

Sei $z = f(x, y)$ eine Funktion und $x = x(t)$, $y = y(t)$ eine sogenannte *Parametrisierung* von x und y ; dies bedeutet, dass x und y Funktionen einer gemeinsamen Variablen (hier t , man nennt t den *Parameter*) sind. Dann ist $z = z(t) = f(x(t), y(t))$ eine Funktion von t und kann wie folgt abgeleitet werden.

Satz 10.3 (Kettenregel)

$$z'(t) = \frac{df}{dt} = f_x \frac{dx}{dt} + f_y \frac{dy}{dt} = f_x(x(t), y(t)) \cdot x'(t) + f_y(x(t), y(t)) \cdot y'(t)$$

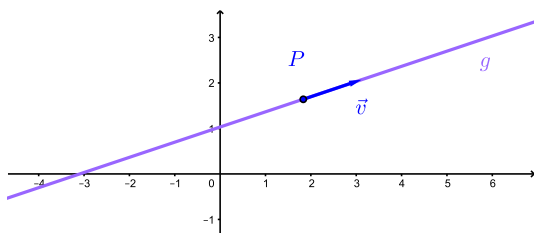
Beispiel

Sei $z = f(x, y) = xy^2$ mit $x = x(t) = e^{3t}$ und $y = y(t) = \sin(t)$.

10.3 Richtungsableitung, Gradient und Hesse-Matrix

Wir haben gesehen, dass f_x die Steigung in x -Richtung (d.h. der xz -Kurve) und f_y die Steigung in y -Richtung (d.h. der yz -Kurve) angibt. Wie sieht aber die Steigung entlang einer beliebigen Richtung aus?

Sei also $f : D \rightarrow \mathbb{R}$ eine Funktion in zwei Variablen und $P = (x_0, y_0)$ ein Punkt in D , in welchem f differenzierbar ist. Wir untersuchen $f(x, y)$, wobei wir (x, y) einschränken auf Punkte auf einer (beliebigen aber festen) Geraden g durch P . Sei $\vec{v} = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix}$ ein Richtungsvektor der Geraden g der Länge 1.



Die Punkte auf der Geraden g können also parametrisiert werden durch

$$(x(t), y(t)) = (x_0 + tx_1, y_0 + ty_1).$$

Damit ist f eingeschränkt auf g eine Funktion $f(x(t), y(t))$ von t und wir können sie mit der Kettenregel (Satz 10.3) ableiten. Wir erhalten

Definition Sei $\vec{v} = \begin{pmatrix} x_1 \\ y_1 \end{pmatrix}$ ein Vektor der Länge 1. Man nennt

$$\frac{\partial f}{\partial \vec{v}}(x_0, y_0) = f_x(x_0, y_0) \cdot x_1 + f_y(x_0, y_0) \cdot y_1$$

die *Richtungsableitung* von f an der Stelle (x_0, y_0) in Richtung des Vektors $\vec{v}$.

Die Richtungsableitung gibt die Steigung von f an der Stelle (x_0, y_0) in Richtung $\vec{v}$ an. Damit diese Steigung nur von der Richtung und nicht von der Länge von $\vec{v}$ abhängt, muss der Vektor $\vec{v}$ die Länge 1 haben.

Für die Spezialfälle $\vec{v} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ und $\vec{v} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ erhalten wir:

Beispiel

Wie gross ist die Steigung der Funktion $f(x, y) = 3xy - 2y^2$ an der Stelle $(5, 4)$ in Richtung $\begin{pmatrix} 1 \\ -2 \end{pmatrix}$?

Nun wollen wir die Richtung bestimmen, in welche die Steigung (oder das Wachstum) von f am grössten ist. Dazu ist es praktisch, den Gradienten von f zu definieren.

Definition Sei $f : D \rightarrow \mathbb{R}$ und $(x_0, y_0) \in D$. Der *Gradient* von f in (x_0, y_0) ist definiert durch

$$\text{grad} f(x_0, y_0) = \begin{pmatrix} f_x(x_0, y_0) \\ f_y(x_0, y_0) \end{pmatrix}, \quad \text{bzw. kurz: } \text{grad} f = \begin{pmatrix} f_x \\ f_y \end{pmatrix}.$$

Analog definiert man den Gradienten für $D \subset \mathbb{R}^3$ und $f = f(x, y, z)$.

Speziell in der Physik nutzt man für eine kürzere Schreibweise den *Nabla-Operator* ∇ :

$$\nabla = \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \end{pmatrix} \quad \Rightarrow \quad \nabla f = \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \end{pmatrix} f = \begin{pmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \end{pmatrix} = \text{grad} f$$

Analog für eine Funktion $f = f(x, y, z)$ in drei Variablen.

Die Richtungsableitung können wir nun mit Hilfe des Skalarproduktes schreiben. Sei γ der Zwischenwinkel der Vektoren ∇f und $\vec{v}$. Dann gilt

$$\frac{\partial f}{\partial \vec{v}} = f_x x_1 + f_y y_1 = \nabla f \cdot \vec{v} = \|\nabla f\| \|\vec{v}\| \cos \gamma = \|\nabla f\| \cos \gamma,$$

da $\vec{v}$ die Länge 1 hat. Es folgt:

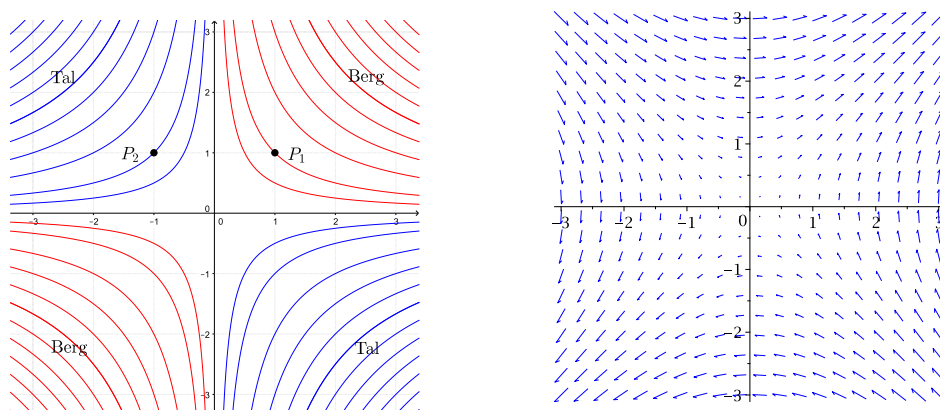
Eigenschaften des Gradienten

- Der Gradient ∇f zeigt in die Richtung der grössten Steigung von f .
- Der Gradient ∇f steht senkrecht auf den Niveaulinien.

Beispiel

Sei $f(x, y) = xy$. Wir betrachten die Punkte $P_1 = (1, 1)$ und $P_2 = (-1, 1)$.

Hier sind die Niveaulinien und die Gradienten (die Gradienten sind verkürzt eingezeichnet):



Quadratische Approximation

Als Vorbereitung für das Bestimmen von Extremalstellen einer Funktion f in zwei Variablen wollen wir f durch eine quadratische Funktion annähern.

Wir erinnern uns (letztes Semester, Abschnitt 4.5), dass eine in x_0 zweimal differenzierbare reelle Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ in der Nähe von x_0 durch das Taylorpolynom

$$p_2(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{f''(x_0)}{2}(x - x_0)^2$$

approximiert werden kann. Dieses Taylorpolynom hat die Eigenschaft, dass $p_2(x_0) = f(x_0)$, $p_2'(x_0) = f'(x_0)$ und $p_2''(x_0) = f''(x_0)$.

Sei nun $f : D \rightarrow \mathbb{R}$ eine Funktion in zwei Variablen, welche an der Stelle (x_0, y_0) zweimal stetig differenzierbar ist. Wir suchen ein quadratisches Polynom $p(x, y) = p_2(x, y)$, so dass $p(x_0, y_0) = f(x_0, y_0)$ und alle partiellen Ableitungen erster und zweiter Ordnung von $p(x, y)$ und $f(x, y)$ in (x_0, y_0) übereinstimmen. Mit diesen Bedingungen erhalten wir das Polynom

$$p(x, y) = f(x_0, y_0) + f_x(x_0, y_0)(x - x_0) + f_y(x_0, y_0)(y - y_0) + \frac{1}{2} \left(f_{xx}(x_0, y_0)(x - x_0)^2 + 2f_{xy}(x_0, y_0)(x - x_0)(y - y_0) + f_{yy}(x_0, y_0)(y - y_0)^2 \right).$$

Der lineare Teil davon ist die Tangentialebene in (x_0, y_0) an f .

Wir können dieses Polynom kompakter schreiben mit Hilfe der Bezeichnungen

$$\vec{x} = \begin{pmatrix} x \\ y \end{pmatrix} \quad \text{und} \quad \vec{a} = \begin{pmatrix} x_0 \\ y_0 \end{pmatrix}.$$

Es folgt

$$\begin{pmatrix} x - x_0 \\ y - y_0 \end{pmatrix} = \vec{x} - \vec{a} \quad \text{und} \quad (x - x_0, y - y_0) = (\vec{x} - \vec{a})^T.$$

Wir fassen also f als Funktion von Vektoren auf (der Ortsvektoren der Punkte in D).

Die Tangentialebene lässt sich damit schreiben als

$$T(\vec{x}) = f(\vec{a}) + \nabla f(\vec{a})^T (\vec{x} - \vec{a}).$$

Der quadratische Teil von $p(x, y)$ ist eine quadratische Form und lässt sich schreiben als

$$\frac{1}{2}(\vec{x} - \vec{a})^T H_f(\vec{a})(\vec{x} - \vec{a})$$

wobei

$$H_f(\vec{a}) = \begin{pmatrix} f_{xx}(\vec{a}) & f_{xy}(\vec{a}) \\ f_{yx}(\vec{a}) & f_{yy}(\vec{a}) \end{pmatrix}$$

die *Hesse-Matrix* von f an der Stelle $\vec{a}$ ist. Unter der Voraussetzung, dass f in $\vec{a} = (x_0, y_0)^T$ zweimal stetig differenzierbar ist, gilt $f_{xy}(\vec{a}) = f_{yx}(\vec{a})$. Die Hesse-Matrix ist in diesem Fall also symmetrisch.

Satz 10.4 *Die quadratische Approximation der Funktion f an der Stelle $\vec{a}$ ist gegeben durch*

$$p(\vec{x}) = f(\vec{a}) + \nabla f(\vec{a})^T (\vec{x} - \vec{a}) + \frac{1}{2}(\vec{x} - \vec{a})^T H_f(\vec{a})(\vec{x} - \vec{a}).$$

Der Vergleich mit dem Taylorpolynom $p_2(x)$ einer reellen Funktion zeigt, dass die Hesse-Matrix die Verallgemeinerung der zweiten Ableitung ist.

Beispiel

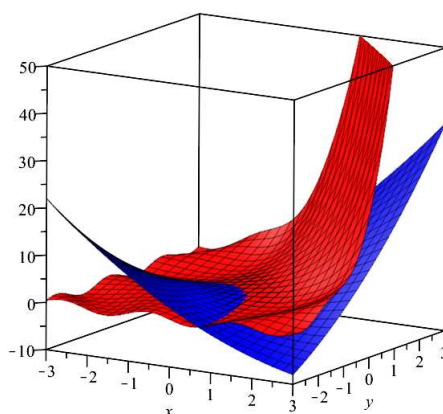
Gesucht ist die quadratische Approximation der Funktion $f(x, y) = e^{x+y} + \sin(xy)$ an der Stelle $(x_0, y_0) = (0, 0)$.

Weiter ist

$$\begin{aligned} f_{xx} &= e^{x+y} - y^2 \sin(xy) && \implies f_{xx}(0, 0) = 1 \\ f_{yy} &= e^{x+y} - x^2 \sin(xy) && \implies f_{yy}(0, 0) = 1 \\ f_{xy} &= e^{x+y} - xy \sin(xy) + \cos(xy) = f_{yx} && \implies f_{xy}(0, 0) = 2 \end{aligned}$$

Damit erhalten wir

Graphen von $f(x, y)$ und $p(x, y)$



10.4 Lokale und globale Extrema

In diesem Abschnitt wollen wir lokale und globale Extremalstellen von Funktionen in zwei Variablen finden und untersuchen. Auch diese Methoden lassen sich problemlos auf Funktionen von mehr als zwei Variablen übertragen.

Lokale Extrema

Definition Die Funktion $f : D \rightarrow \mathbb{R}$ besitzt in $(x_0, y_0) \in D$ ein *lokales Minimum* bzw. ein *lokales Maximum*, falls

$$f(x, y) \geq f(x_0, y_0) \quad \text{bzw.} \quad f(x, y) \leq f(x_0, y_0)$$

für alle (x, y) in der Nähe von (x_0, y_0) . Ein lokales Minimum bzw. Maximum (x_0, y_0) heisst *isoliert*, falls $f(x, y) \neq f(x_0, y_0)$ für alle $(x, y) \neq (x_0, y_0)$ in der Nähe von (x_0, y_0) .

Von differenzierbaren Funktionen in einer Variablen wissen wir, dass deren Ableitung an einer lokalen Extremalstelle verschwindet. Ist nun (x_0, y_0) eine lokale Extremalstelle einer (in (x_0, y_0) differenzierbaren) Funktion $f(x, y)$, so haben die xz - und die yz -Kurven auf dem Graphen von f ebenfalls ein lokales Extremum in x_0 bzw. y_0 , also gilt (vgl. Seite 119)

$$f_x(x_0, y_0) = f_y(x_0, y_0) = 0.$$

Satz 10.5 Hat f in (x_0, y_0) eine lokale Extremalstelle, dann gilt

$$\nabla f(x_0, y_0) = \vec{0} \quad \text{d.h.} \quad f_x(x_0, y_0) = 0 \quad \text{und} \quad f_y(x_0, y_0) = 0.$$

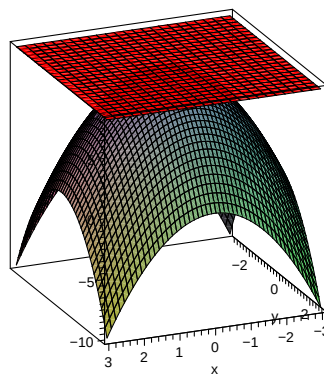
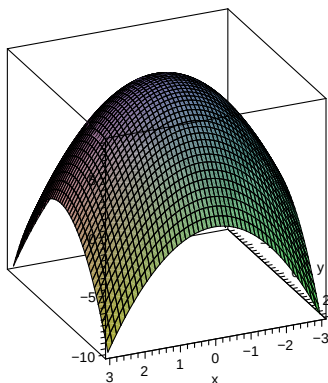
Geometrisch bedeutet dies, dass die Tangentialebene in (x_0, y_0) an den Graphen von f horizontal ist,

$$z = T(x, y) = f(x_0, y_0) .$$

Die Umkehrung von Satz 10.5 ist im Allgemeinen falsch.

Beispiele

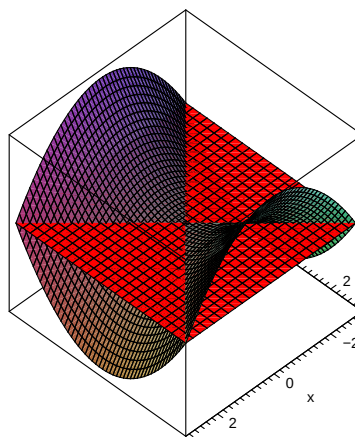
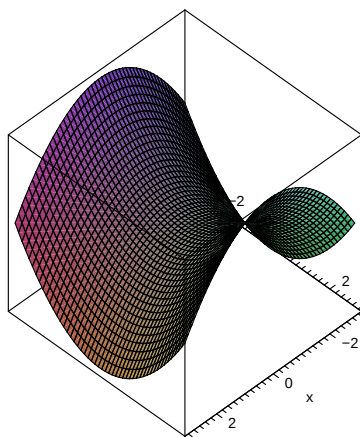
1. Sei $f(x, y) = 8 - x^2 - y^2$. An der Stelle $(x_0, y_0) = (0, 0)$ hat f ein isoliertes Maximum und tatsächlich gilt



2. Sei $f(x, y) = y^2 - x^2$. An der Stelle $(0, 0)$ gilt $\nabla f(0, 0) = \vec{0}$, aber f hat dort kein lokales Extremum. Es ist $f(0, 0) = 0$, und in der Nähe von $(0, 0)$ nimmt f sowohl positive als auch negative Werte an,

$$f(0, y) > 0 \text{ für alle } y \neq 0 \quad \text{und} \quad f(x, 0) < 0 \text{ für alle } x \neq 0 .$$

Die Funktion hat in $(0, 0)$ einen *Sattelpunkt*.



Definition Eine Funktion $f(x, y)$ hat in (x_0, y_0) einen *Sattelpunkt*, wenn $\nabla f(x_0, y_0) = \vec{0}$, aber (x_0, y_0) keine Extremalstelle ist.

Um entscheiden zu können, ob ein lokales Minimum, Maximum oder ein Sattelpunkt vorliegt, müssen wir die zweiten partiellen Ableitungen von f betrachten, das heisst die Hesse-Matrix von f . Wie vorher gehen wir davon aus, dass $f_{xy} = f_{yx}$ und daher die Hesse-Matrix symmetrisch ist.

Sei f eine Funktion mit $\nabla f(x_0, y_0) = \nabla f(\vec{a}) = \vec{0}$. Dies bedeutet also, dass f in (x_0, y_0) ein lokales Minimum, ein lokales Maximum oder einen Sattelpunkt hat.

Mit der quadratischen Approximation von Satz 10.4 gilt dann

$$f(\vec{x}) \approx f(\vec{a}) + \frac{1}{2}(\vec{x} - \vec{a})^T H_f(\vec{a})(\vec{x} - \vec{a})$$

falls $\vec{x} - \vec{a}$ "klein" ist, das heisst, falls $\vec{x}^T = (x, y)$ nahe bei $\vec{a}^T = (x_0, y_0)$ ist. Das lokale Verhalten der Funktion f hängt also in der Nähe von (x_0, y_0) hauptsächlich ab von den Werten der quadratischen Form

$$(\vec{x} - \vec{a})^T H_f(\vec{a})(\vec{x} - \vec{a}).$$

Nehmen wir nun an, dass die Hesse-Matrix $H_f(\vec{a}) = H_f(x_0, y_0)$ positiv definit ist. Dann gilt

$$(\vec{x} - \vec{a})^T H_f(\vec{a})(\vec{x} - \vec{a}) > 0$$

für alle $(x, y) \neq (x_0, y_0)$ in der Nähe von (x_0, y_0) . Dies bedeutet, dass $f(x_0, y_0) = f(\vec{a})$ ein isoliertes lokales Minimum ist.

Analog bedeutet eine negativ definite Hesse-Matrix $H_f(\vec{a})$ ein isoliertes lokales Maximum von f in (x_0, y_0) . Die dritte Möglichkeit, nämlich dass $H_f(\vec{a})$ indefinit ist, bedeutet geometrisch einen Sattelpunkt in (x_0, y_0) .

Für die Zusammenfassung dieser drei Möglichkeiten nutzen wir Satz 9.6 zur Bestimmung der Definitheit der Hesse-Matrix, wobei wir die folgende Abkürzung verwenden:

$$\Delta = \det(H_f(x_0, y_0)) = f_{xx}(x_0, y_0)f_{yy}(x_0, y_0) - f_{xy}(x_0, y_0)^2$$

Satz 10.6 Sei $f(x, y)$ in (x_0, y_0) zweimal stetig differenzierbar und

$$\nabla f(x_0, y_0) = \vec{0}.$$

Dann gilt:

- $\Delta > 0, f_{xx}(x_0, y_0) > 0$ (d.h. die Hesse-Matrix $H_f(\vec{a})$ ist positiv definit)
 $\implies f(x_0, y_0)$ ist ein isoliertes lokales Minimum
- $\Delta > 0, f_{xx}(x_0, y_0) < 0$ (d.h. die Hesse-Matrix $H_f(\vec{a})$ ist negativ definit)
 $\implies f(x_0, y_0)$ ist ein isoliertes lokales Maximum
- $\Delta < 0$ (d.h. die Hesse-Matrix $H_f(\vec{a})$ ist indefinit)
 $\implies f$ hat in (x_0, y_0) einen Sattelpunkt

Gilt $\nabla f(x_0, y_0) = \vec{0}$ und $\Delta = 0$, dann ist keine allgemeine Aussage möglich.

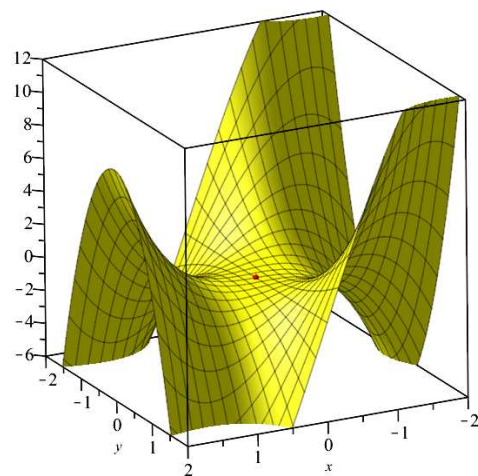
Beispiele

1. Sei $f(x, y) = 8 - x^2 - y^2$. Wir haben schon berechnet, dass $f_x = -2x$, $f_y = -2y$ und dass $\nabla f(0, 0) = \vec{0}$. Die Hesse-Matrix ist nun

Also ist $f(0, 0) = 8$ tatsächlich ein isoliertes lokales Maximum.

2. Sei $f(x, y) = xy$.

3. Sei $f(x, y) = x^3 - 3xy^2$.



4. Sei $f(x, y) = (x^2 - 1)^2 + y^2 - 1$.

Wir müssen also die Stellen $(0, 0)$, $(1, 0)$ und $(-1, 0)$ untersuchen.

Globale Extrema

Wie bei Funktionen in einer Variablen können lokale Extremalstellen auch *globale* Extremalstellen sein, müssen aber nicht. Um neben den lokalen auch die globalen Extremalstellen zu finden, nehmen wir an, dass f auf einer “anständigen” Menge D definiert ist, zum Beispiel auf einem Kreis, einem Vieleck, einem Streifen oder auf ganz $\mathbb{R}^2$.

In diesem Fall können wir wie folgt vorgehen:

- (1) Untersuche die Stellen im Inneren von D , wo $\nabla f = \vec{0}$ gilt;
- (2) Untersuche f an denjenigen Stellen, wo f nicht stetig differenzierbar ist;
- (3) Untersuche f auf dem Rand von D (falls dieser zu D gehört).

Beispiel

Sei $D \subset \mathbb{R}^2$ das Dreieck mit den Ecken $(1, 0)$, $(0, 0)$, $(0, 1)$, also gegeben durch $x \geq 0$, $y \geq 0$ und $x + y \leq 1$ und sei $f(x, y) = xy$.

- (1) Wir haben schon gesehen, dass $\nabla f \neq \vec{0}$ im Inneren von D ;
- (2) f ist auf ganz $\mathbb{R}^2$ stetig differenzierbar.

(3)

Die Funktion f hat also ein globales Maximum in $(\frac{1}{2}, \frac{1}{2})$ mit $f(\frac{1}{2}, \frac{1}{2}) = \frac{1}{4}$ und ein globales Minimum für $x = 0$ oder $y = 0$ mit $f(0, y) = f(x, 0) = 0$.

10.5 Extrema mit Nebenbedingung

In vielen Anwendungsbeispielen sucht man nach Extremalstellen einer Funktion $f : D \rightarrow \mathbb{R}$, $D \subset \mathbb{R}^2$, entlang einer Kurve in D .

Darstellung von Kurven

Eine Kurve in der xy -Ebene kann durch eine Gleichung der Form

$$\phi(x, y) = 0$$

beschrieben werden. Diese Darstellung einer Kurve nennt man *implizit*.

Beispiele

1. $\phi(x, y) = x^2 - y = 0$ beschreibt eine Parabel.
2. $\phi(x, y) = x^2 + y^2 - 1 = 0$ beschreibt den Einheitskreis.
3. $\phi(x, y) = 5x^2 - 4xy + 8y^2 - 36 = 0$ beschreibt eine Ellipse (vgl. das Beispiel auf Seite 109).

Die Kurve des ersten Beispiels könnte auch in einer anderen Form beschrieben werden. Die Gleichung $x^2 - y = 0$ kann nach y aufgelöst werden: $y = x^2$. Die Parabel ist also der Graph einer Funktion $y = f(x)$, nämlich $f(x) = x^2$. Die Darstellung einer Kurve in der Form $y = f(x)$ oder $x = f(y)$ nennt man *explizit*.

Eine explizite Darstellung ist jedoch nicht immer möglich (eher die Ausnahme). Die Gleichung $\phi(x, y) = 0$ des Einheitskreises (oder auch der Ellipse) kann nicht nach y oder x aufgelöst werden. Der Einheitskreis ist demnach nicht der Graph einer Funktion $y = f(x)$ oder $x = f(y)$. Zum Beispiel erhält man durch $y = f(x) = \sqrt{1 - x^2}$ nur den oberen Halbkreis oder $x = f(y) = -\sqrt{1 - y^2}$ ergibt nur den linken Halbkreis.

Die implizite Kurvengleichung $\phi(x, y) = 0$ kann jedoch als Niveaulinie

$$\phi(x, y) = c = 0$$

der Funktion $\phi(x, y)$ aufgefasst werden.

Extrema entlang einer Kurve

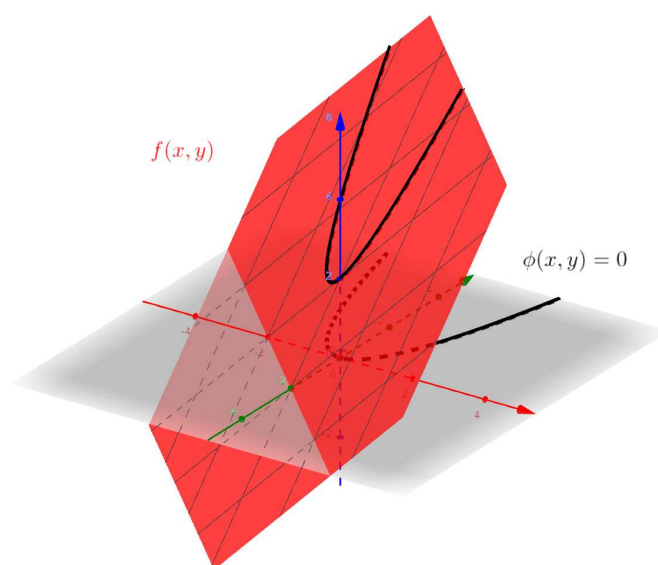
Wir suchen nun nach (lokalen oder globalen) Extremalstellen (x_0, y_0) einer Funktion $f(x, y)$ unter der *Nebenbedingung*

$$\phi(x, y) = 0.$$

Das heisst, wir untersuchen f eingeschränkt auf die Kurve definiert durch $\phi(x, y) = 0$.

Beispiel

Die Funktion $f(x, y) = x + y + 2$ hat auf $D = \mathbb{R}^2$ weder lokale noch globale Extrema. Eingeschränkt auf die Kurve $\phi(x, y) = x^2 - y = 0$ (eine Parabel) hat f jedoch ein globales Minimum, nämlich $f(-\frac{1}{2}, \frac{1}{4}) = \frac{7}{4}$.



1. Fall: Die Gleichung der Kurve kann in *expliziter Form* gegeben werden.

Das heisst, die Gleichung $\phi(x, y) = 0$ kann nach y oder nach x aufgelöst werden zu $y = g(x)$, bzw. $x = h(y)$. Die Funktion $f(x, y)$ wird damit zu einer Funktion in nur einer Variablen:

$$f(x, y) = f(x, g(x)) = \tilde{f}(x), \quad \text{bzw.} \quad f(x, y) = f(h(y), y) = \bar{f}(y).$$

Nun können wir mit den gewohnten Methoden für Funktionen in einer Variablen die Extremalstellen von $\tilde{f}$, bzw. $\bar{f}$ bestimmen.

Beispiel

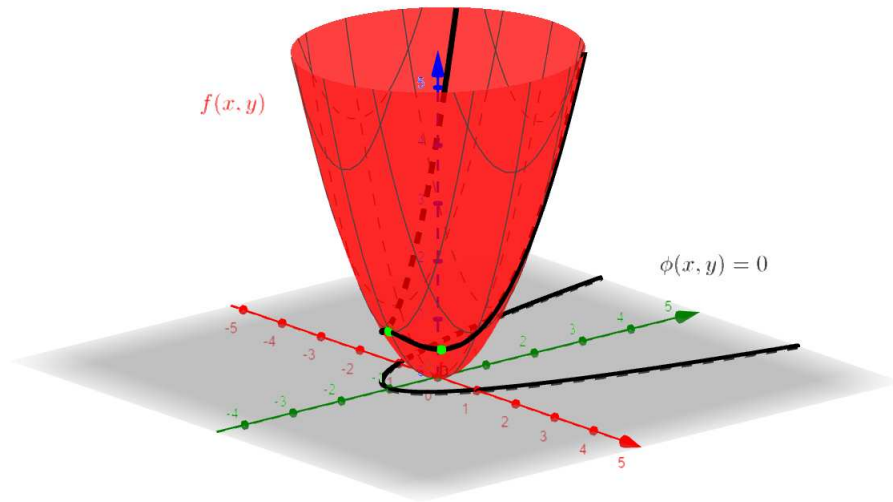
Sei $f(x, y) = x^2 + y^2$ und die Nebenbedingung

$$\phi(x, y) = x^2 - y - 1 = 0.$$

Gesucht sind also die Extrema von f eingeschränkt auf die Parabel $y = g(x) = x^2 - 1$.

Die Extremalstellen x von $\tilde{f}$ erfüllen die Bedingung $\tilde{f}'(x) = 4x^3 - 2x = 2x(2x^2 - 1) = 0$. Damit

finden wir sofort die Extremalstellen $x_1 = 0$, $x_{2,3} = \pm \frac{1}{\sqrt{2}}$ von $\tilde{f}$, bzw. $(0, -1)$ und $(\pm \frac{1}{\sqrt{2}}, -\frac{1}{2})$ von f . Der Funktionswert $f(0, -1) = 1$ ist ein lokales Maximum und $f(\pm \frac{1}{\sqrt{2}}, -\frac{1}{2}) = \frac{3}{4}$ sind lokale (und globale) Minima von f unter der Nebenbedingung $\phi(x, y) = 0$.



2. Fall: Die Gleichung der Kurve ist in *impliziter Form* gegeben.

Nehmen wir an, wir hätten eine lokale Maximalstelle (x_0, y_0) von $f(x, y)$ unter der Nebenbedingung $\phi(x, y) = 0$ gefunden. Wir betrachten die Niveaulinie N_c zum Niveau $c = f(x_0, y_0)$. Diese Linie N_c trennt die beiden Gebiete $f(x, y) > c$ und $f(x, y) < c$. Die Kurve C definiert durch $\phi(x, y) = 0$ geht ebenfalls durch den Punkt (x_0, y_0) . Da (x_0, y_0) eine (lokale) Maximalstelle von f unter der Nebenbedingung $\phi = 0$ ist, liegt diese Kurve (in der Nähe von (x_0, y_0)) ganz auf der Seite $f(x, y) \leq c$ der Niveaulinie N_c .

Die Niveaulinie N_c und die Kurve C sind daher im Punkt (x_0, y_0) tangential (bzw. die Tangenten an die Niveaulinie N_c und an die Kurve C sind identisch). Die Kurve C fassen wir nun auf als Niveaulinie $\phi(x, y) = c = 0$ der Funktion $\phi(x, y)$. Da der Gradient einer Funktion senkrecht auf den Niveaulinien steht, sind die Gradienten von f und ϕ in (x_0, y_0) folglich parallel.

Satz 10.7 Sei (x_0, y_0) eine lokale Extremalstelle von f unter der Nebenbedingung $\phi(x, y) = 0$. Ist $\nabla\phi(x_0, y_0) \neq 0$, dann gilt

$$\nabla f(x_0, y_0) = \lambda \cdot \nabla\phi(x_0, y_0)$$

für ein $\lambda \in \mathbb{R}$.

Zur Bestimmung von (x_0, y_0) bilden wir die Hilfsfunktion

$$F(x, y, \lambda) = f(x, y) - \lambda\phi(x, y)$$

und suchen (x, y, λ) mit $\nabla F(x, y, \lambda) = \vec{0}$.

Es gilt nämlich

$$\nabla F = \begin{pmatrix} F_x \\ F_y \\ F_\lambda \end{pmatrix} = \begin{pmatrix} f_x - \lambda \phi_x \\ f_y - \lambda \phi_y \\ -\phi \end{pmatrix} .$$

Also ist $\nabla F = \vec{0}$ äquivalent zu den beiden Bedingungen

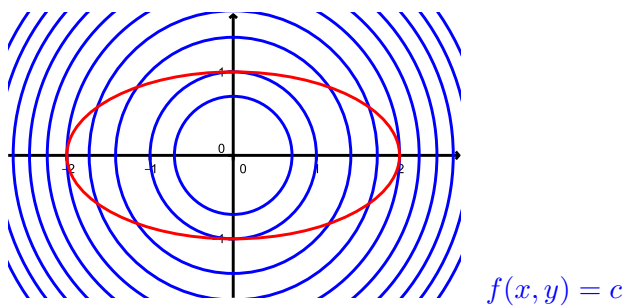
$$\nabla f = \lambda \cdot \nabla \phi \quad \text{und} \quad \phi = 0 .$$

Beispiel

Sei $f(x, y) = x^2 + y^2$ und die Nebenbedingung

$$\phi(x, y) = \frac{x^2}{4} + y^2 - 1 = 0 .$$

Wir suchen also die Extremalstellen auf der Ellipse $\phi(x, y) = 0$ mit den Halbachsen 2 und 1.



Wir setzen

$$F(x, y, \lambda) = f(x, y) - \lambda \phi(x, y) = x^2 + y^2 - \lambda \left(\frac{x^2}{4} + y^2 - 1 \right) .$$

11 Vektorfelder und Wegintegrale

Das Ziel dieses Kapitels ist die Integration entlang einer Kurve in der Ebene oder im Raum, wobei der Integrand eine reellwertige Funktion oder ein sogenanntes Vektorfeld ist.

11.1 Vektorfelder

Der Gradient ∇f einer Funktion $f(x, y) : D \rightarrow \mathbb{R}$ ordnet jedem Punkt $(x_0, y_0) \in D$ einen Vektor, nämlich $\nabla f(x_0, y_0) \in \mathbb{R}^2$, zu. Eine solche Zuordnung nennt man Vektorfeld. Das Vektorfeld definiert durch einen Gradienten nennt man auch *Gradientenfeld*.

Definition Sei $D \subset \mathbb{R}^2$. Ein *Vektorfeld auf D* ist eine Abbildung $F : D \rightarrow \mathbb{R}^2$, welche jedem Punkt $(x, y) \in D$ einen Vektor $F(x, y) \in \mathbb{R}^2$ zuordnet.

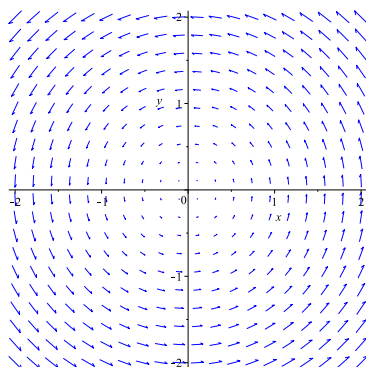
Analog ist ein Vektorfeld auf einer Menge $D \subset \mathbb{R}^3$ als eine Abbildung $F : D \rightarrow \mathbb{R}^3$ definiert.

In den Naturwissenschaften treffen wir beispielsweise auf Vektorfelder definiert durch elektrische, magnetische oder Gravitationskräfte oder durch Windgeschwindigkeiten.

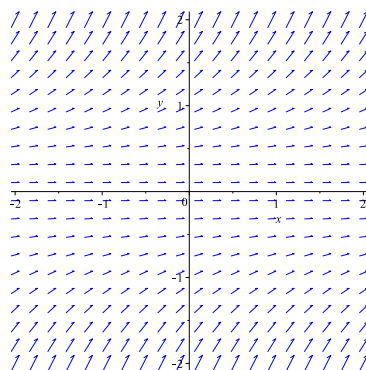
Bei graphischen Darstellungen von Vektorfeldern wird oft nur die Richtung der Vektoren richtig wiedergegeben, deren Betrag jedoch skaliert.

Beispiele

1. Sei $F(x, y) = \begin{pmatrix} -y \\ x \end{pmatrix}$ definiert auf $D = \mathbb{R}^2$.



2. Sei $F(x, y) = \begin{pmatrix} 2 \\ y^2 \end{pmatrix}$ definiert auf $D = \mathbb{R}^2$.



Jedes Gradientenfeld ist ein Vektorfeld, aber nicht jedes Vektorfeld ist ein Gradientenfeld.

Definition Ein Vektorfeld $F(x, y) : D \rightarrow \mathbb{R}^2$, das ein Gradientenfeld ist, heisst *konservativ*. Das heisst, $F(x, y)$ ist konservativ, falls es eine Funktion $f(x, y) : D \rightarrow \mathbb{R}$ gibt mit

$$F = \nabla f.$$

Eine solche Funktion f heisst *Potentialfunktion* des Vektorfeldes F .

Nach Satz 10.1 erfüllt eine “anständige” Funktion $f(x, y)$ die Beziehung $f_{xy} = f_{yx}$. Für das Gradientenfeld $\nabla f = \begin{pmatrix} f_x \\ f_y \end{pmatrix}$ einer solchen Funktion gilt daher

$$\frac{\partial f_x}{\partial y} = \frac{\partial f_y}{\partial x} \quad \text{bzw.} \quad \frac{\partial f_y}{\partial x} - \frac{\partial f_x}{\partial y} = 0.$$

Für ein Vektorfeld $F(x, y) = \begin{pmatrix} u(x, y) \\ v(x, y) \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix}$ können wir deshalb folgern:

$$F \text{ konservativ} \quad \implies \quad \frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} = 0$$

Diese für ein Gradientenfeld notwendige Bedingung nennt man *Integrabilitätsbedingung*.

Diese Bedingung ist auch hinreichend, wenn das Vektorfeld F auf $D = \mathbb{R}^2$ definiert und *stetig differenzierbar* ist. Letzteres bedeutet, dass $u(x, y)$ und $v(x, y)$ differenzierbar und die partiellen Ableitungen stetig sind.

Satz 11.1 Sei $F(x, y) = \begin{pmatrix} u \\ v \end{pmatrix}$ ein stetig differenzierbares Vektorfeld auf $D = \mathbb{R}^2$. Dann gilt:

$$F \text{ konservativ} \quad \iff \quad \frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} = 0$$

Beispiele

$$1. F(x, y) = \begin{pmatrix} 4x^3y^2 \\ 2x^4y + 2y \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix}$$

$$2. F(x, y) = \begin{pmatrix} -y \\ x \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix}$$

Analog ist der Gradient einer Funktion $f(x, y, z) : D \rightarrow \mathbb{R}$ mit $D \subset \mathbb{R}^3$ gegeben durch $\nabla f = \begin{pmatrix} f_x \\ f_y \\ f_z \end{pmatrix}$ und für f "anständig" gilt nun

$$f_{yz} = f_{zy}, f_{zx} = f_{xz}, f_{xy} = f_{yx}.$$

Für ein Vektorfeld $F(x, y, z) = \begin{pmatrix} u(x, y, z) \\ v(x, y, z) \\ w(x, y, z) \end{pmatrix} = \begin{pmatrix} u \\ v \\ w \end{pmatrix}$ können wir deshalb folgern:

$$F \text{ konservativ} \quad \implies \quad \frac{\partial w}{\partial y} - \frac{\partial v}{\partial z} = 0, \quad \frac{\partial u}{\partial z} - \frac{\partial w}{\partial x} = 0, \quad \frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} = 0$$

Wir haben hier also drei *Integrabilitätsbedingungen*.

Definition Sei $F(x, y, z) = \begin{pmatrix} u \\ v \\ w \end{pmatrix}$ ein Vektorfeld auf $D \subset \mathbb{R}^3$. Die *Rotation* von F ist das Vektorfeld auf D definiert durch

$$\text{rot } F = \begin{pmatrix} \frac{\partial w}{\partial y} - \frac{\partial v}{\partial z} \\ \frac{\partial u}{\partial z} - \frac{\partial w}{\partial x} \\ \frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \end{pmatrix}.$$

Symbolisch kann die Rotation mit Hilfe des Nabla-Operators ∇ und des Vektorprodukts geschrieben werden:

$$\text{rot } F = \nabla \times F = \begin{pmatrix} \frac{\partial}{\partial x} \\ \frac{\partial}{\partial y} \\ \frac{\partial}{\partial z} \end{pmatrix} \times \begin{pmatrix} u \\ v \\ w \end{pmatrix}$$

Die Integrabilitätsbedingungen für ein Vektorfeld $F : D \rightarrow \mathbb{R}^3$ sind also gleichbedeutend mit $\text{rot } F = \vec{0}$. Analog zu Satz 11.1 gilt nun die folgende Äquivalenz.

Satz 11.2 Sei F ein stetig differenzierbares Vektorfeld auf $D \subset \mathbb{R}^3$. Dann gilt:

$$F \text{ konservativ} \quad \iff \quad \text{rot } F = \vec{0}$$

Mit Hilfe der Sätze 11.1 und 11.2 können wir also feststellen, ob ein auf $\mathbb{R}^2$ oder $\mathbb{R}^3$ definiertes Vektorfeld konservativ ist. Ist dies der Fall, wie finden wir eine zugehörige Potentialfunktion? Schauen wir dazu zwei typische Beispiele an.

Beispiele

1. Gegeben sei das Vektorfeld

$$F(x, y) = \begin{pmatrix} 12xy^3 \\ 18x^2y^2 + 7y^6 \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix} \quad \text{auf } D = \mathbb{R}^2.$$

Die Stetigkeitsbedingungen von Satz 11.1 sind erfüllt. Also überprüfen wir die Integrabilitätsbedingung:

Nach Satz 11.1 ist F konservativ.

Nun wollen wir eine Funktion $f(x, y)$ mit $\nabla f = F$ finden.

(1) Integration von u nach x :

(2) Ableiten von f nach y und Gleichsetzen mit v :

(3) Integration von $g_y(y)$ nach y :

(4) Einsetzen von $g(y)$ in f :

Mit dieser Methode kann zu jedem konservativen Vektorfeld F auf $\mathbb{R}^2$ eine Potentialfunktion f gefunden werden. Zwei verschiedene Potentialfunktionen zum gleichen Vektorfeld unterscheiden sich dabei nur um eine Konstante.

Diese Methode kann auf Vektorfelder auf $\mathbb{R}^3$ angepasst werden.

2. Gegeben sei das Vektorfeld

$$F(x, y, z) = \begin{pmatrix} e^x y + 1 \\ e^x + z \\ y \end{pmatrix} = \begin{pmatrix} u \\ v \\ w \end{pmatrix} \quad \text{auf } D = \mathbb{R}^3.$$

Die Stetigkeitsbedingungen von Satz 11.2 sind erfüllt. Also berechnen wir $\operatorname{rot} F$:

Nach Satz 11.2 ist F konservativ.

Auch hier bestimmen wir eine Funktion $f(x, y, z)$ mit $\nabla f = F$.

(1) Integration von u nach x :

(2) Ableiten von f nach y und Gleichsetzen mit v :

(3) Integration von $g_y(y, z)$ nach y und Einsetzen von $g(y, z)$ in f :

(4) Ableiten von f nach z und Gleichsetzen mit w :

(5) Integration von $h_z(z)$ nach z und Einsetzen von $h(z)$ in f aus (3):

Allgemeiner gelten die Sätze 11.1 und 11.2 für Vektorfelder F auf $D \subset \mathbb{R}^2$, bzw. $D \subset \mathbb{R}^3$, falls D *offen* (d.h. ohne Rand) und *einfach zusammenhängend* (d.h. je zwei Punkte in D können mit einer "regulären" Kurve verbunden werden und jede geschlossene Kurve in D ist auf einen Punkt stetig zusammenziehbar, ohne D zu verlassen) ist. Zum Beispiel ist jeder Kreis oder jede Halbebene in $\mathbb{R}^2$ einfach zusammenhängend, aber nicht $\mathbb{R}^2 \setminus \{(0, 0)\}$.

Ist also der Definitionsbereich D eines Vektorfeldes F nicht einfach zusammenhängend, dann sind die Sätze 11.1 und 11.2 im Allgemeinen falsch, das heisst genauer, die Pfeile $\Leftarrow$ sind falsch.

Beispiel

Wir betrachten auf $D = \mathbb{R}^2 \setminus \{(0, 0)\}$ (damit ist D nicht einfach zusammenhängend) das Vektorfeld

$$F(x, y) = \frac{1}{x^2 + y^2} \begin{pmatrix} -y \\ x \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix}.$$

Die Integrabilitätsbedingung ist erfüllt, denn

$$\frac{\partial v}{\partial x} = \frac{y^2 - x^2}{(x^2 + y^2)^2} = \frac{\partial u}{\partial y}.$$

Für $x \neq 0$ finden wir

$$f(x, y) = \arctan\left(\frac{y}{x}\right)$$

mit $\nabla f = F$. Diese Funktion f lässt sich jedoch nicht stetig auf ganz $\mathbb{R}^2 \setminus \{(0, 0)\}$ fortsetzen. Das Vektorfeld F hat auf ganz $\mathbb{R}^2 \setminus \{(0, 0)\}$ keine Potentialfunktion und ist deshalb nicht

konservativ auf D .

Auf der einfach zusammenhängenden Menge $\tilde{D} = \{(x, y) \in \mathbb{R}^2 \mid x > 0\}$ hingegen ist F konservativ und die obige Funktion f ist eine Potentialfunktion.

Wir werden im Abschnitt 11.3 über Wegintegrale noch auf eine andere Weise sehen, dass das obige Vektorfeld auf D kein Gradientenfeld ist.

11.2 Wege und Kurven

Um im nächsten Abschnitt über Kurven integrieren zu können, müssen wir diese zunächst genau definieren und ihre Eigenschaften untersuchen.

Definition Ein *Weg* in $\mathbb{R}^n$ ist eine Abbildung

$$\vec{x} : I \longrightarrow \mathbb{R}^n \quad t \mapsto \vec{x}(t) = \begin{pmatrix} x_1(t) \\ \vdots \\ x_n(t) \end{pmatrix}$$

eines Intervalls $I = [a, b] \subset \mathbb{R}$ in den $\mathbb{R}^n$, wobei die Funktionen $x_i : I \longrightarrow \mathbb{R}$ stetig sind.

Der Weg heisst *(stetig) differenzierbar*, wenn die Funktionen x_i (stetig) differenzierbar sind (stetig differenzierbar bedeutet, dass die Ableitungen $x'_i(t)$ wieder stetig sind).

Das Bild $C = \vec{x}(I)$ nennt man eine *Kurve* in $\mathbb{R}^n$ und $\vec{x}$ eine *Parametrisierung* von C .

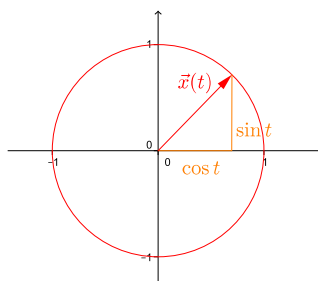
Eine Parametrisierung einer Kurve ist nicht eindeutig. Von ihr hängt ab, mit welcher Geschwindigkeit die Kurve durchlaufen wird. Ist $\vec{x}$ differenzierbar, dann nennt man den Vektor der Ableitungen

$$\dot{\vec{x}}(t) = \begin{pmatrix} x'_1(t) \\ \vdots \\ x'_n(t) \end{pmatrix}$$

den *Geschwindigkeitsvektor* von $\vec{x}$ an der Stelle t . Gilt $\dot{\vec{x}}(t_0) \neq \vec{0}$ dann ist $\dot{\vec{x}}(t_0)$ tangential an die Kurve im Punkt $\vec{x}(t_0)$.

Beispiele

1. Sei $\vec{x}(t) = \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}$ für $t \in I = [0, 2\pi]$.



Für $\vec{x}(t)$ wie oben und $I = [-\pi, \pi]$ erhalten wir ebenfalls den Einheitskreis, der Weg beginnt und endet nun aber im Punkt $(-1, 0)$.

Eine andere Parametrisierung des Einheitskreises wäre zum Beispiel $\vec{x}(t) = \begin{pmatrix} \cos(t^2) \\ \sin(t^2) \end{pmatrix}$ für $t \in I = [0, \sqrt{2\pi}]$. Damit wird der Einheitskreis schneller durchlaufen.

Nun ist $\vec{x}(t) = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ schon für $t = \sqrt{\frac{\pi}{2}}$ und der Geschwindigkeitsvektor ist

2. Sei $\vec{x}(t) = \begin{pmatrix} t \\ t^2 \\ t^3 \end{pmatrix}$ für $t \in I = [-1, 1]$. Für $t = 0$ ist $\vec{x}(0) = \vec{0}$.

Die Kurve $C = \vec{x}(I)$ geht also durch den Ursprung und in diesem Punkt ist die x -Achse die Tangente an C .

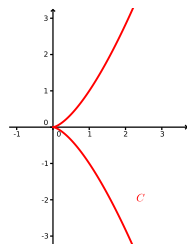
Definition Eine Kurve heisst *einfach*, wenn sie sich nicht überkreuzt und nicht berührt, ausser eventuell am Anfangs- und Endpunkt.

Eine einfache Kurve heisst *geschlossen*, wenn Anfangs- und Endpunkt einer Parametrisierung übereinstimmen.

Eine Kurve heisst *regulär*, falls es eine Parametrisierung $\vec{x}(t)$ der Kurve gibt, die stetig differenzierbar ist mit $\dot{\vec{x}}(t) \neq \vec{0}$ für alle t .

Beispiel

Die Kurve C parametrisiert durch $\vec{x}(t) = \begin{pmatrix} t^2 \\ t^3 \end{pmatrix}$ für $t \in [-2, 2]$ ist nicht regulär, da $\dot{\vec{x}}(0) = \vec{0}$.



11.3 Wegintegrale

Das Wegintegral (oder Kurvenintegral) ist eine Verallgemeinerung des bestimmten Integrals

$$\int_a^b f(x) dx ,$$

wobei nun nicht über ein Intervall $I = [a, b]$ auf der x -Achse sondern über einen Weg, bzw. eine Kurve in der Ebene oder im Raum integriert wird.

Wegintegral von reellwertigen Funktionen

Definition Sei C in $\mathbb{R}^n$ eine einfache, reguläre Kurve parametrisiert durch $\vec{x} : [a, b] \rightarrow C$ und sei $f : C \rightarrow \mathbb{R}$ eine stetige Funktion. Dann ist das *Wegintegral* von f über C definiert durch

$$\int_C f ds = \int_a^b f(\vec{x}(t)) \|\dot{\vec{x}}(t)\| dt .$$

Hier bezeichnet

$$s = s(t) = \int_a^t \|\dot{\vec{x}}(u)\| du$$

die Bogenlänge von C zwischen den Punkten $\vec{x}(a)$ und $\vec{x}(t)$. Damit ist

$$\int_C ds = \int_a^b \|\dot{\vec{x}}(t)\| dt = \text{Länge der Kurve } C .$$

Das Wegintegral ist unabhängig von der Wahl der Parametrisierung der Kurve C . Ähnlich wie das bestimmte Integral $\int_a^b f(x) dx$ kann $\int_C f ds$ als Flächeninhalt “zwischen” $f(x, y)$ und C interpretiert werden, falls $f(x, y) \geq 0$ für alle (x, y) .

Beispiele

1. Wir berechnen die Länge des Einheitskreises C in $\mathbb{R}^2$. Wir haben schon gesehen, dass C parametrisiert werden kann durch

$$\vec{x}(t) = \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix} \quad \text{mit} \quad \dot{\vec{x}}(t) = \begin{pmatrix} -\sin(t) \\ \cos(t) \end{pmatrix} \quad \text{für } t \in [0, 2\pi] .$$

Damit folgt

Zum Vergleich wählen wir eine andere Parametrisierung von C . Der Weg

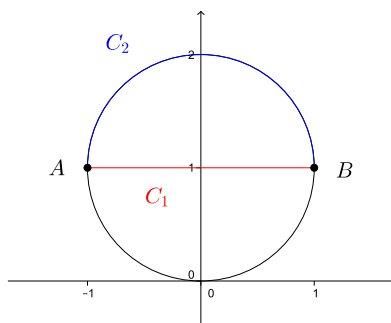
$$\vec{x}(t) = \begin{pmatrix} \cos(2t + \pi) \\ \sin(2t + \pi) \end{pmatrix} \quad \text{mit} \quad \dot{\vec{x}}(t) = \begin{pmatrix} -2\sin(2t + \pi) \\ 2\cos(2t + \pi) \end{pmatrix} \quad \text{für } t \in [0, \pi]$$

durchläuft den Einheitskreis doppelt so schnell wie vorher und der Anfangs- und Endpunkt ist nun $(-1, 0)$ und nicht $(1, 0)$. Wir erhalten damit

2. Gegeben sei ein Draht der Form C . Die Massendichte sei gegeben durch die Funktion $f(x, y)$. Dann ist $\int_C f \, ds$ die Gesamtmasse des Drahtes.

3. Sei $f(x, y) = \frac{1}{x^2 + y^2}$.

Diese Funktion kann als Intensität einer Strahlung mit Strahlenquelle im Ursprung interpretiert werden. Die Strahlung nimmt mit dem Quadrat der Entfernung von der Strahlenquelle ab. Wir wollen die Strahlenbelastung auf zwei verschiedenen Wegen von $A = (-1, 1)$ nach $B = (1, 1)$ berechnen, wobei die Durchlaufgeschwindigkeit konstant und gleich ist.



- C_1 : Strecke von A nach B . Es gilt

Es folgt $\|\dot{\vec{x}}(t)\| = 1$ für alle t und wir erhalten

- C_2 : Oberer Halbkreis von A nach B . Es gilt

$$\vec{x}(t) = \begin{pmatrix} -\cos(t) \\ \sin(t) + 1 \end{pmatrix} \quad \text{mit} \quad \dot{\vec{x}}(t) = \begin{pmatrix} \sin(t) \\ \cos(t) \end{pmatrix} \quad \text{für } t \in [0, \pi].$$

Es folgt $\|\dot{\vec{x}}(t)\| = 1$ für alle t und wir erhalten

$$\int_{C_2} f \, ds = \int_0^\pi \frac{1}{\cos^2(t) + (\sin(t) + 1)^2} dt = \int_0^{\frac{\pi}{2}} \frac{1}{\sin(t) + 1} dt = \int_0^{\frac{\pi}{2}} \frac{1}{\cos(t) + 1} dt.$$

Da $\cos(t) + 1 = 2\cos^2(\frac{t}{2})$, finden wir

$$\int_{C_2} f \, ds = \frac{1}{2} \int_0^{\frac{\pi}{2}} \frac{1}{\cos^2(\frac{t}{2})} dt = \tan\left(\frac{t}{2}\right) \Big|_0^{\frac{\pi}{2}} = \tan\left(\frac{\pi}{4}\right) = 1.$$

Auf dem zweiten Weg ist die Strahlenbelastung also kleiner als auf dem ersten Weg, obwohl der zweite Weg länger ist als der erste Weg. Allgemein kann man zeigen, dass die Strahlenbelastung minimal ist auf dem die Strahlenquelle nicht enthaltenden Bogen AB des Kreises durch A , B und die Strahlenquelle.

Wegintegral von Vektorfeldern

Wir können auch über ein Vektorfeld entlang einer Kurve integrieren.

Definition Sei C in $\mathbb{R}^n$ eine einfache, reguläre Kurve parametrisiert durch $\vec{x} : [a, b] \rightarrow C$ und sei $F : C \rightarrow \mathbb{R}^n$ ein stetiges Vektorfeld. Dann ist das *Wegintegral* von F über C definiert durch

$$\int_C F \cdot d\vec{s} = \int_a^b F(\vec{x}(t)) \cdot \dot{\vec{x}}(t) dt.$$

Der Punkt im Integranden auf der rechten Seite bedeutet dabei das Skalarprodukt. Er sollte deshalb auch auf der linken Seite in der Bezeichnung geschrieben werden. Dieses Wegintegral ist bis aufs Vorzeichen unabhängig von der Wahl der Parametrisierung von C .

Das vektorielle Wegintegral hat die folgende physikalische Bedeutung. Sei $F(\vec{x}(t))$ die Kraft, die auf ein Teilchen an der Stelle $\vec{x}(t)$ wirkt (zum Beispiel in einem Gravitationsfeld oder elektrischen Kraftfeld). Dann liefert das Wegintegral von F über C die Arbeit, die aufgewendet wird, um das Teilchen längs C zu bewegen.

Beispiele

1. Sei $F(x, y, z) = \begin{pmatrix} xy \\ x - z \\ xz \end{pmatrix}$ und C die Kurve parametrisiert durch $\vec{x}(t) = \begin{pmatrix} t \\ t^3 \\ 3 \end{pmatrix}$, für $t \in [0, 2]$.

2. Sei $F(x, y) = \begin{pmatrix} 3x \\ 0 \end{pmatrix}$ und C der Einheitskreis.

In diesem Beispiel ist das Wegintegral über jede geschlossene Kurve gleich Null. Dies hängt mit dem folgenden Satz zusammen.

Satz 11.3 *Sei F ein konservatives Vektorfeld auf $D \subset \mathbb{R}^n$ mit zugehöriger Potentialfunktion f und sei $C \subset D$ eine einfache, reguläre Kurve parametrisiert durch $\vec{x} : [a, b] \rightarrow C$. Dann gilt*

$$\int_C F \cdot d\vec{s} = f(\vec{x}(b)) - f(\vec{x}(a)).$$

Insbesondere hängt das Wegintegral nicht vom gewählten Weg ab, sondern nur vom Anfangs- und Endpunkt. Es gilt also

$$\int_C F \cdot d\vec{s} = 0$$

falls die Kurve C geschlossen ist.

Wir überprüfen den Satz in $\mathbb{R}^2$. Wegen $F = \nabla f$ gilt

$$\begin{aligned} \int_C F \cdot d\vec{s} &= \int_a^b \begin{pmatrix} f_x(\vec{x}(t)) \\ f_y(\vec{x}(t)) \end{pmatrix} \cdot \begin{pmatrix} x'(t) \\ y'(t) \end{pmatrix} dt = \int_a^b (f_x(\vec{x}(t)) x'(t) + f_y(\vec{x}(t)) y'(t)) dt \\ &= \int_a^b \frac{d}{dt} f(\vec{x}(t)) dt = f(\vec{x}(t)) \Big|_a^b = f(\vec{x}(b)) - f(\vec{x}(a)). \end{aligned}$$

Beispiele

1. Wir betrachten nochmals das 2. Beispiel von vorher mit $F(x, y) = \begin{pmatrix} 3x \\ 0 \end{pmatrix}$. Dieses Vektorfeld ist konservativ, denn zum Beispiel ist $f(x, y) = \frac{3}{2}x^2$ eine Potentialfunktion von F . Nach Satz 11.3 ist also das Wegintegral über jede geschlossene Kurve gleich Null.

2. Sei $F(x, y, z) = \begin{pmatrix} e^x y + 1 \\ e^x + z \\ y \end{pmatrix}$ und C die Kurve parametrisiert durch $\vec{x}(t) = \begin{pmatrix} t \\ \sqrt{t+1} \\ 0 \end{pmatrix}$, für $t \in [0, 3]$.

3. Wir betrachten nochmals das 1. Beispiel von vorher mit $F(x, y, z) = \begin{pmatrix} xy \\ x - z \\ xz \end{pmatrix}$. Die dort gegebene Kurve C hat den Anfangspunkt $(0, 0, 3)$ und den Endpunkt $(2, 8, 3)$. Wir wählen nun eine andere Kurve, nämlich $\tilde{C}$ parametrisiert durch $\vec{x}(t) = \begin{pmatrix} t \\ 4t \\ 3 \end{pmatrix}$, für $t \in [0, 2]$, die denselben Anfangs- und Endpunkt hat.

4. Sei $F(x, y) = \frac{1}{x^2+y^2} \begin{pmatrix} -y \\ x \end{pmatrix}$ wie auf Seite 140 und sei C der Einheitskreis. Damit gilt

$$F(\vec{x}(t)) \cdot \dot{\vec{x}}(t) = \begin{pmatrix} -\sin(t) \\ \cos(t) \end{pmatrix} \cdot \begin{pmatrix} -\sin(t) \\ \cos(t) \end{pmatrix} = \sin^2(t) + \cos^2(t) = 1$$

und wir erhalten

$$\int_C F \cdot d\vec{s} = \int_0^{2\pi} dt = 2\pi \neq 0.$$

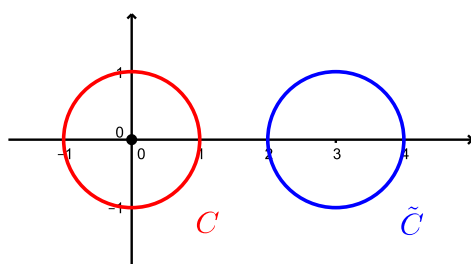
Mit Satz 11.3 können wir also folgern, dass F nicht konservativ auf $D = \mathbb{R}^2 \setminus \{(0, 0)\}$ ist.

Nun verschieben wir den Kreis so, dass er den Ursprung nicht umläuft. Zum Beispiel betrachten wir den Kreis $\tilde{C}$ parametrisiert durch

$$\vec{x}(t) = \begin{pmatrix} 3 + \cos(t) \\ \sin(t) \end{pmatrix} \quad \text{für } t \in [0, 2\pi] .$$

Damit liegt der Kreis und die vom Kreis umschlossene Fläche in der einfach zusammenhängenden Menge $\tilde{D} = \{ (x, y) \in \mathbb{R}^2 \mid x > 0 \}$. Auf $\tilde{D}$ ist F konservativ und es gilt

$$\int_{\tilde{C}} F \cdot d\vec{s} = 0 .$$



Dies bestätigt unsere Überlegungen auf den Seiten 140–141.

12 Integration in mehreren Variablen

Das bestimmte Integral $\int_a^b f(x) dx$ liefert den Inhalt der Fläche, die zwischen dem Intervall $[a, b]$ und dem Graphen von f eingeschlossen ist. Bei einem *Bereichsintegral*

$$\iint_D f(x, y) dx dy$$

wird das Volumen des Körpers bestimmt, der zwischen dem Bereich $D \subset \mathbb{R}^2$ und dem Graphen von f eingeschlossen ist.

12.1 Bereichsintegrale

Wir betrachten zunächst Rechtecke als Bereiche D in $\mathbb{R}^2$, das heisst

$$D = [a, b] \times [c, d] = \{ (x, y) \mid a \leq x \leq b, c \leq y \leq d \} \subset \mathbb{R}^2.$$

Sei $f : D \rightarrow \mathbb{R}$ eine stetige Funktion. Der Graph von f schliesst mit dem Rechteck D einen Körper ein, dessen Volumen wir nun berechnen werden.

Betrachten wir eine feste Zahl $y_0 \in [c, d]$, so ist das Integral

$$F(y_0) = \int_{x=a}^b f(x, y_0) dx$$

der Flächeninhalt des *Querschnitts* $\{ (x, y_0, f(x, y_0)) \mid x \in [a, b] \}$ des eingeschlossenen Körpers. Durch Integration von $F(y)$ über das Intervall $[c, d]$ erhalten wir das Volumen V dieses Körpers,

$$V = \int_{y=c}^d F(y) dy = \int_{y=c}^d \left(\int_{x=a}^b f(x, y) dx \right) dy.$$

Dabei spielt es keine Rolle, ob man zuerst nach x und dann nach y oder umgekehrt integriert, solange f stetig auf D ist.

Analog kann man das Bereichsintegral für Quader $[a, b] \times [c, d] \times [k, \ell] \subset \mathbb{R}^3$ für Funktionen $f(x, y, z)$ in drei Variablen erklären.

Beispiele

1. Sei $D = [a, b] \times [c, d]$ und $f : D \rightarrow \mathbb{R}$ mit $f(x, y) = 1$.

Allgemein erhält man durch Integration von $f = 1$ über den Bereich D den Flächeninhalt von D . Man ermittelt nämlich das Volumen V des Körpers der Höhe 1 über dem Bereich D , was mit dem Flächeninhalt von D übereinstimmt, da "Volumen = Grundfläche $\cdot$ Höhe" gilt. Dies wird demnächst von Nutzen sein, wenn wir über kompliziertere Bereiche integrieren.

2. Sei $D = [0, 1] \times [1, 2] \times [2, 4]$ und $f : D \rightarrow \mathbb{R}$ mit $f(x, y, z) = 2x + z + 1$.

Manchmal kürzt man die Schreibweise der Integrale ab und schreibt

$$\int_D f = \int_D f dF = \iint_D f dx dy \quad \text{bzw.} \quad \int_D f = \int_D f dV = \iiint_D f dx dy dz$$

für einen Bereich D in $\mathbb{R}^2$, bzw. $\mathbb{R}^3$.

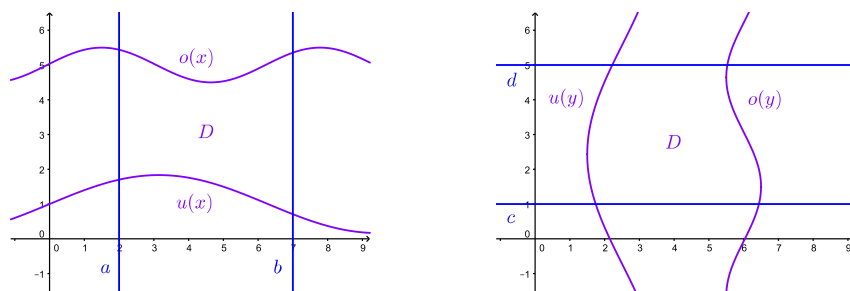
Integration über Normalbereiche

Allgemeiner als Rechtecke in $\mathbb{R}^2$ sind Bereiche des $\mathbb{R}^2$ der Form

$$\begin{aligned} D &= \{ (x, y) \mid a \leq x \leq b, u(x) \leq y \leq o(x) \} \quad \text{bzw.} \\ D &= \{ (x, y) \mid c \leq y \leq d, u(y) \leq x \leq o(y) \}, \end{aligned}$$

wobei u und o reelle Funktionen sind (u steht für *untere Grenze* und o für *obere Grenze*). Man nennt einen solchen Bereich D einen *Normalbereich*.

Normalbereiche in $\mathbb{R}^2$ sehen wie folgt aus:



Analog ist ein Bereich D in $\mathbb{R}^3$ ein *Normalbereich*, wenn er von der Form

$$D = \{ (x, y, z) \mid a \leq x \leq b, u(x) \leq y \leq o(x), \tilde{u}(x, y) \leq z \leq \tilde{o}(x, y) \}$$

ist, wobei die Rollen der Koordinaten x, y, z vertauscht sein können.

Definition Ist D ein Normalbereich in $\mathbb{R}^2$, bzw. $\mathbb{R}^3$, dann nennt man

$$\iint_D f(x, y) dy dx = \int_{x=a}^b \int_{y=u(x)}^{o(x)} f(x, y) dy dx \quad \text{bzw.}$$

$$\iiint_D f(x, y, z) dz dy dx = \int_{x=a}^b \int_{y=u(x)}^{o(x)} \int_{z=\tilde{u}(x,y)}^{\tilde{o}(x,y)} f(x, y, z) dz dy dx$$

das *Doppelintegral*, bzw. *Dreifachintegral* über D .

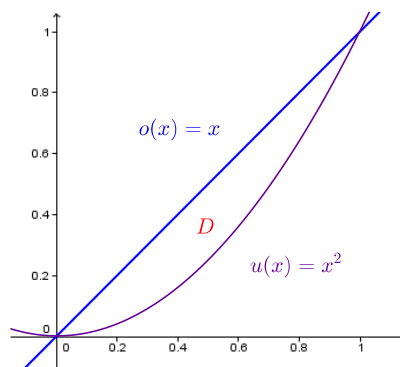
Wie schon nach dem 1. Beispiel auf Seite 150 bemerkt, erhält man durch Integration von $f(x, y) = 1$, bzw. $f(x, y, z) = 1$ den Flächeninhalt, bzw. das Volumen von D .

Satz 12.1 Für einen Normalbereich D gilt

$$\int_D 1 = \text{Flächeninhalt, bzw. Volumen von } D$$

Beispiel

Sei $D = \{ (x, y) \mid 0 \leq x \leq 1, x^2 \leq y \leq x \}$.



- (a) Wir berechnen den Flächeninhalt von D . Nach Satz 12.1 ist er gleich

Die rechte Seite ist nun genau das, was wir aus der Schule kennen. Wir finden den Flächeninhalt $\frac{1}{2} - \frac{1}{3} = \frac{1}{6}$.

- (b) Wir berechnen das Doppelintegral von $f(x, y) = 2xy$ über D .

12.2 Koordinatentransformationen

Oft sind Integrale über Normalbereiche schwierig zu berechnen, da die unteren und oberen Grenzen $u(x)$ und $o(x)$ nach dem Einsetzen in die Stammfunktionen zu komplizierten Integranden führen. In diesen Fällen kann der Wechsel zu einem anderen Koordinatensystem hilfreich sein. Das heisst, man wechselt zu Polar-, Zylinder oder Kugelkoordinaten.

- *Polarkoordinaten* (r, φ)

Die Polarkoordinaten bilden ein Koordinatensystem von $\mathbb{R}^2$. Die Umrechnung lautet

$$\begin{aligned}x &= r \cos \varphi \\y &= r \sin \varphi\end{aligned}$$

für $r \geq 0$ und $\varphi \in [0, 2\pi)$.

- *Zylinderkoordinaten* (r, φ, z)

Die Zylinderkoordinaten ergänzen die Polarkoordinaten um die z -Koordinate zu einem Koordinatensystem von $\mathbb{R}^3$. Die Umrechnung ist dementsprechend

$$\begin{aligned}x &= r \cos \varphi \\y &= r \sin \varphi \\z &= z\end{aligned}$$

für $r \geq 0$, $\varphi \in [0, 2\pi)$ und $z \in \mathbb{R}$.

- *Kugelkoordinaten* (r, φ, ϑ)

Die Kugelkoordinaten sind ein Koordinatensystem von $\mathbb{R}^3$. Die Umrechnung lautet

$$\begin{aligned}x &= r \cos \varphi \sin \vartheta \\y &= r \sin \varphi \sin \vartheta \\z &= r \cos \vartheta\end{aligned}$$

für $r \geq 0$, $\varphi \in [0, 2\pi)$ und $\vartheta \in [0, \pi]$.

Ebenfalls üblich ist es, anstatt des Winkels ϑ den Winkel $\tilde{\vartheta} = \frac{\pi}{2} - \vartheta$ zu benutzen, wobei dann $\tilde{\vartheta} \in [-\frac{\pi}{2}, \frac{\pi}{2}]$ ist und $\sin \vartheta$ (bzw. $\cos \vartheta$) in den Umrechnungsformeln durch $\cos \tilde{\vartheta}$ (bzw. $\sin \tilde{\vartheta}$) zu ersetzen ist. Bei der Erdkugel entspricht φ dem Längengrad und $\tilde{\vartheta}$ dem Breitengrad.

Integriert man nun über eine Funktion und wechselt das Koordinatensystem, dann braucht es im Integral bezüglich der neuen Koordinaten einen Korrekturfaktor.

Satz 12.2 *Sei $D \subset \mathbb{R}^2$, bzw. $D \subset \mathbb{R}^3$ ein Bereich im kartesischen Koordinatensystem. Für eine Funktion $f : D \rightarrow \mathbb{R}$ gelten die folgenden Transformationsformeln.*

- *Integration in Polarkoordinaten:*

$$\iint_D f(x, y) \, dx \, dy = \int_r \int_\varphi f(r \cos \varphi, r \sin \varphi) \, r \, d\varphi \, dr$$

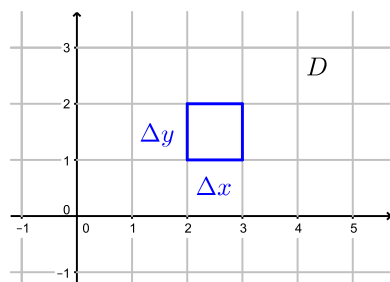
- *Integration in Zylinderkoordinaten:*

$$\iiint_D f(x, y, z) \, dx \, dy \, dz = \int_r \int_\varphi \int_z f(r \cos \varphi, r \sin \varphi, z) \, r \, dz \, d\varphi \, dr$$

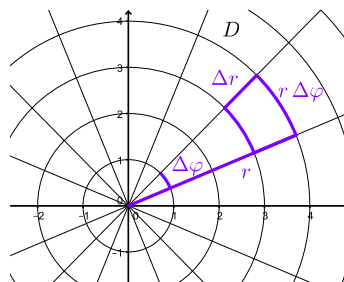
- *Integration in Kugelkoordinaten:*

$$\iiint_D f(x, y, z) \, dx \, dy \, dz = \int_r \int_\varphi \int_\vartheta f(r \cos \varphi \sin \vartheta, r \sin \varphi \sin \vartheta, r \cos \vartheta) \, r^2 \sin \vartheta \, d\vartheta \, d\varphi \, dr$$

Woher kommt beispielsweise der Korrekturfaktor r bei der Integration in Polarkoordinaten? Bei der Integration in kartesischen Koordinaten unterteilt man den Integrationsbereich D in kleine Rechtecke mit Seitenlängen Δx und Δy . Das heisst, man summiert über Rechtecke mit Flächeninhalt $\Delta x \Delta y$, bzw. integriert über infinitesimal kleine Rechtecke mit Flächeninhalt $dx \, dy$. Bei der Integration in Polarkoordinaten wird der Bereich D in kleine Ringteilflächen unterteilt, deren Flächeninhalt ungefähr $r \Delta r \Delta \varphi$ ist. Dies führt zu $r \, dr \, d\varphi$ im Integral.



kartesische Koordinaten



Polarkoordinaten

Die Korrekturfaktoren in den anderen beiden Fällen sind analog erklärbar. Allgemeiner kann eine beliebige Koordinatentransformation durchgeführt werden (ähnlich der Substitution bei einem Integral einer reellen Funktion). Der Korrekturfaktor berechnet sich dann durch die sogenannte Jacobideterminante.

Beispiele

1. Sei $D \subset \mathbb{R}^2$ der Kreisring mit Aussenkreisradius 2 und Innenkreisradius 1 und sei $f(x, y) = x(x^2 + y^2)$.

2. Das Volumen V des Zylinders vom Radius R und der Höhe h ist gegeben durch

$$V = \int_{r=0}^R \int_{\varphi=0}^{2\pi} \int_{z=0}^h 1 \cdot r \, dz d\varphi dr = 2\pi h \int_{r=0}^R r \, dr = \pi R^2 h.$$

12.3 Flächenintegrale

Das bestimmte Integral über einem Intervall in $\mathbb{R}$ haben wir verallgemeinert zu einem Wegintegral über einer Kurve in $\mathbb{R}^2$. Ähnlich können wir das Bereichsintegral über einem Rechteck in $\mathbb{R}^2$ verallgemeinern zu einem Flächenintegral über einer Fläche in $\mathbb{R}^3$.

Eine Kurve ist das Bild einer Abbildung (Parametrisierung) in einer Variablen t . Eine Fläche ist das Bild einer Abbildung (Parametrisierung) in zwei Variablen u und v .

Beispiele von Flächen sind die Kugeloberfläche (Sphäre), die Oberfläche eines Zylinders, die Oberfläche eines Kegels oder der Graph einer Funktion $f(x, y) : D \rightarrow \mathbb{R}$.

Definition Eine Teilmenge $\mathcal{F} \subset \mathbb{R}^3$ heisst *Fläche*, wenn es eine (stückweise) stetig differenzierbare Abbildung

$$\vec{x} : B \subset \mathbb{R}^2 \rightarrow \mathbb{R}^3, \quad (u, v) \mapsto \vec{x}(u, v) = \begin{pmatrix} x(u, v) \\ y(u, v) \\ z(u, v) \end{pmatrix}$$

mit $\vec{x}(B) = \mathcal{F}$ gibt. Die Fläche heisst *regulär*, wenn

$$\vec{x}_u(u, v) \times \vec{x}_v(u, v) \neq \vec{0}, \quad \text{wobei } \vec{x}_u(u, v) = \begin{pmatrix} x_u(u, v) \\ y_u(u, v) \\ z_u(u, v) \end{pmatrix}, \quad \vec{x}_v(u, v) = \begin{pmatrix} x_v(u, v) \\ y_v(u, v) \\ z_v(u, v) \end{pmatrix},$$

für alle $(u, v) \in B$ (bis auf endlich viele Ausnahmen).

Die Vektoren $\vec{x}_u(u, v)$ und $\vec{x}_v(u, v)$ sind Tangentialvektoren an die Fläche $\mathcal{F}$ im Punkt $\vec{x}(u, v)$. Die Bedingung $\vec{x}_u \times \vec{x}_v = \vec{x}_u(u, v) \times \vec{x}_v(u, v) \neq \vec{0}$ bedeutet, dass diese beiden Vektoren linear unabhängig sind. Der Vektor $\vec{x}_u \times \vec{x}_v$ steht senkrecht zu $\vec{x}_u$ und $\vec{x}_v$, das heisst senkrecht auf der Fläche $\mathcal{F}$. Man nennt $\vec{x}_u \times \vec{x}_v$ *Normalenvektor* von $\mathcal{F}$.

Dank den Zusätzen “stückweise” und “bis auf endlich viele Ausnahmen” in der Definition können auch zusammengesetzte Flächenstücke (wie zum Beispiel die Oberfläche des Zylinders, zusammengesetzt aus Mantel, Boden und Deckel) als reguläre Fläche betrachtet werden.

Beispiele

1. Die Kugeloberfläche $\mathcal{F}$ vom Radius 2 ist eine reguläre Fläche parametrisiert durch

$$\vec{x} : [0, 2\pi] \times [0, \pi] \longrightarrow \mathcal{F}, \quad \vec{x}(\varphi, \vartheta) = \begin{pmatrix} 2 \cos \varphi \sin \vartheta \\ 2 \sin \varphi \sin \vartheta \\ 2 \cos \vartheta \end{pmatrix}$$

2. Der Graph der Funktion $f(x, y) = 8 - x^2 - y^2$ ist eine reguläre Fläche parametrisiert durch

$$\vec{x} : \mathbb{R}^2 \longrightarrow \text{Graph}(f), \quad \vec{x}(u, v) = \begin{pmatrix} u \\ v \\ 8 - u^2 - v^2 \end{pmatrix}.$$

Es gilt

Wir definieren nun Flächenintegrale für Funktionen f und Vektorfelder F analog zu den Wegintegralen. Die Rolle des Geschwindigkeitsvektors $\dot{\vec{x}}(t)$ bei den Wegintegralen übernimmt nun der Vektor $\vec{x}_u \times \vec{x}_v$ bei den Flächenintegralen.

Definition Sei $\mathcal{F} \subset \mathbb{R}^3$ eine reguläre Fläche parametrisiert durch $\vec{x}(u, v)$ für $(u, v) \in B$.

- Sei $f(x, y, z) : \mathcal{F} \longrightarrow \mathbb{R}$ eine stetige Funktion. Das *Flächenintegral* von f über $\mathcal{F}$ ist definiert durch

$$\iint_{\mathcal{F}} f \, dS = \iint_B f(\vec{x}(u, v)) \|\vec{x}_u \times \vec{x}_v\| \, du \, dv$$

- Sei $F(x, y, z) : \mathcal{F} \longrightarrow \mathbb{R}^3$ ein stetiges Vektorfeld. Das *Flächenintegral* von F über $\mathcal{F}$ ist definiert durch

$$\iint_{\mathcal{F}} F \cdot d\vec{S} = \iint_B F(\vec{x}(u, v)) \cdot (\vec{x}_u \times \vec{x}_v) \, du \, dv$$

Beim Wegintegral haben wir durch $\int_C ds$ die Länge der Kurve C erhalten. Analog ist nun

$$\iint_{\mathcal{F}} dS = \text{Flächeninhalt der Fläche } \mathcal{F}$$

Ist F das Geschwindigkeitsfeld einer strömenden Flüssigkeit, dann ist $\iint_{\mathcal{F}} F \cdot d\vec{S}$ die Flüssigkeitsmenge, die pro Zeiteinheit die Fläche $\mathcal{F}$ durchströmt. Man nennt dieses Integral deshalb auch *Fluss von F durch $\mathcal{F}$ in Richtung $\vec{x}_u \times \vec{x}_v$* .

Beispiel

Sei $f(x, y, z) = x^2 + y^2 + 2z$ und $\mathcal{F} = \{ (x, y, z) \mid x^2 + y^2 = 1, 0 \leq z \leq 1 \}$ die Mantelfläche des Zylinders vom Radius 1 und der Höhe 1. Diese Fläche kann parametrisiert werden durch

$$\vec{x}(\varphi, z) = \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ z \end{pmatrix} \quad \text{für } (\varphi, z) \in [0, 2\pi) \times [0, 1].$$

Für das Flächenintegral erhalten wir

12.4 Integralsätze

In diesem letzten Abschnitt werden die Integralsätze von Green, Gauß und Stokes kurz vorgestellt. Für konkrete Berechnungen sind diese Sätze sehr nützlich.

Der Divergenzsatz von Gauß

Der Divergenzsatz von Gauß führt das Flächenintegral $\iint_{\mathcal{F}} F \cdot d\vec{S}$ eines Vektorfeldes F auf ein Dreifachintegral zurück.

Definition Sei $F(x, y, z) = \begin{pmatrix} u \\ v \\ w \end{pmatrix}$ ein Vektorfeld auf $D \subset \mathbb{R}^3$. Die *Divergenz* von F ist definiert durch

$$\operatorname{div} F = \frac{\partial u}{\partial x} + \frac{\partial v}{\partial y} + \frac{\partial w}{\partial z}.$$

Analog ist die *Divergenz* eines Vektorfeldes $F(x, y) = \begin{pmatrix} u \\ v \end{pmatrix}$ definiert durch

$$\operatorname{div} F = \frac{\partial u}{\partial x} + \frac{\partial v}{\partial y}.$$

Symbolisch kann man die Divergenz mit Hilfe des Skalarprodukts

$$\operatorname{div} F = \nabla \cdot F$$

schreiben. Die Divergenz ist also eine reellwertige Funktion, $\operatorname{div} F : D \rightarrow \mathbb{R}$.

Sei F das Geschwindigkeitsfeld einer strömenden Flüssigkeit. Die Divergenz gibt an, ob an einer Stelle $(x, y, z) \in D$ Flüssigkeit entsteht oder verloren geht oder ob Gleichgewicht besteht. Es gilt

- $\operatorname{div} F(x, y, z) > 0 \implies$ *Quelle*: Es fließt mehr ab als zu.
- $\operatorname{div} F(x, y, z) < 0 \implies$ *Senke*: Es fließt mehr zu als ab.
- $\operatorname{div} F(x, y, z) = 0 \implies$ *Quellenfrei*: Es fließt genauso viel zu wie ab.

Wir nennen einen Bereich $D \subset \mathbb{R}^3$ *regulär*, falls D eine geschlossene, reguläre Oberfläche $\mathcal{F}_D$ hat. Typische Beispiele sind Kugel, Zylinder, Kegel oder ein Quader.

Satz 12.3 (Divergenzsatz von Gauß) *Sei $D \subset \mathbb{R}^3$ regulär und seine Oberfläche $\mathcal{F}_D$ so parametrisiert, dass der Normalenvektor $\vec{x}_u \times \vec{x}_v$ nach aussen zeigt. Sei F ein stetig differenzierbares Vektorfeld auf D . Dann gilt*

$$\iiint_D \operatorname{div} F \, dV = \iint_{\mathcal{F}_D} F \cdot d\vec{S}.$$

Ist also beispielsweise das Flächenintegral auf der rechten Seite schwierig zu berechnen, so kann stattdessen das eventuell einfachere Bereichsintegral auf der linken Seite berechnet werden.

Beispiele

1. Sei D die Kugel vom Radius 1 und $F(x, y, z) = \begin{pmatrix} 0 \\ 0 \\ -1 \end{pmatrix}$.

Mit dem Satz von Gauß folgt, dass der Fluss durch die gesamte Kugeloberfläche nach aussen gleich Null ist.

2. Sei D wieder die Kugel vom Radius 1 und $F(x, y, z) = \begin{pmatrix} x \\ y \\ z \end{pmatrix}$. Wir benutzen, dass das Volumen der Kugel gleich $\frac{4}{3}\pi$ ist.

Der Fluss durch die gesamte Kugeloberfläche beträgt also 4π .

Der Satz von Stokes

Beim Satz von Stokes gehen wir von einer regulären Fläche $\mathcal{F} \subset \mathbb{R}^3$ aus, die zwei Seiten hat und deren Rand eine (einfache, reguläre) geschlossene Kurve ist. Ist $\vec{x}(u, v)$ eine Parametrisierung von $\mathcal{F}$, dann muss die Randkurve $C_{\mathcal{F}}$ so parametrisiert werden, dass die Fläche $\mathcal{F}$ zu unserer Linken ist, wenn wir $C_{\mathcal{F}}$ durchlaufen und unser Kopf in Richtung von $\vec{x}_u \times \vec{x}_v$ zeigt.

Satz 12.4 (Satz von Stokes) *Sei F ein stetig differenzierbares Vektorfeld auf der Fläche $\mathcal{F}$. Dann gilt*

$$\iint_{\mathcal{F}} \operatorname{rot} F \cdot d\vec{S} = \int_{C_{\mathcal{F}}} F \cdot d\vec{s}.$$

Anstelle des Flächenintegrals von $\operatorname{rot} F$ können wir also das Wegintegral von F über den Rand von $\mathcal{F}$ berechnen. Insbesondere ist das Flächenintegral von $\operatorname{rot} F$ für alle Flächen mit derselben Randkurve gleich.

Der Satz von Green

Der Satz von Green entspricht dem Satz von Stokes in der Ebene.

Wir betrachten einen Bereich $D \subset \mathbb{R}^2$, der von einer (einfachen, regulären) geschlossenen Kurve C_D berandet ist. Wir parametrisieren die Randkurve C_D so, dass der Bereich zu unserer Linken ist, wenn wir C_D durchlaufen.

Satz 12.5 (Satz von Green) *Sei $F(x, y) = \begin{pmatrix} u \\ v \end{pmatrix}$ ein stetig differenzierbares Vektorfeld auf D . Dann gilt*

$$\iint_D \left(\frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \right) dx dy = \int_{C_D} F \cdot d\vec{s}.$$

Wir können also ein Doppelintegral mit Hilfe eines eventuell einfacheren Wegintegrals berechnen.

Ist der Bereich D einfach zusammenhängend, dann ist F konservativ auf D , genau dann wenn die Integrabilitätsbedingung erfüllt ist (vgl. Satz 11.1). Dies ist genau dann der Fall, wenn das Doppelintegral auf der linken Seite gleich Null ist. Mit dem Satz von Green folgt also, dass F konservativ auf D ist, genau dann wenn das Wegintegral über den Rand von D (eine geschlossene Kurve!) Null ist, in Übereinstimmung mit Satz 11.3.

Betrachten wir nun das spezielle Vektorfeld $F(x, y) = \begin{pmatrix} -y \\ x \end{pmatrix} = \begin{pmatrix} u \\ v \end{pmatrix}$ auf D . Das Doppelintegral auf der linken Seite der Gleichung im Satz von Green wird damit zu

$$\iint_D \left(\frac{\partial v}{\partial x} - \frac{\partial u}{\partial y} \right) dx dy = \iint_D (1 + 1) dx dy = 2 \iint_D dx dy = 2 \text{Flächeninhalt}(D) .$$

Der Satz von Green führt damit zum folgenden praktischen Satz.

Satz 12.6

$$\text{Flächeninhalt}(D) = \frac{1}{2} \int_{C_D} \begin{pmatrix} -y \\ x \end{pmatrix} \cdot d\vec{s} = \int_{C_D} \begin{pmatrix} 0 \\ x \end{pmatrix} \cdot d\vec{s} = \int_{C_D} \begin{pmatrix} -y \\ 0 \end{pmatrix} \cdot d\vec{s}$$

Um also den Flächeninhalt eines Bereiches D zu berechnen, genügt es, entlang des Randes von D zu integrieren!