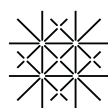


Frühjahrssemester 2023

Statistik 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

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.

14.02.2023

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 2023	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:

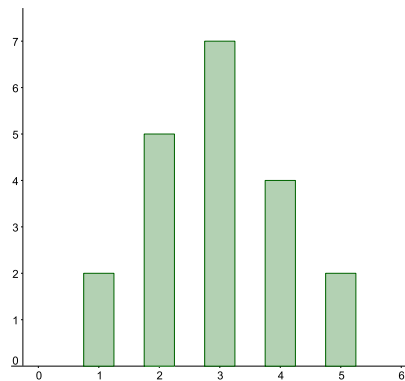
2, 4, 3, 4, 2, 3, 4, 5, 2, 1, 3, 2, 5, 3, 3, 4, 1, 2, 3, 3

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

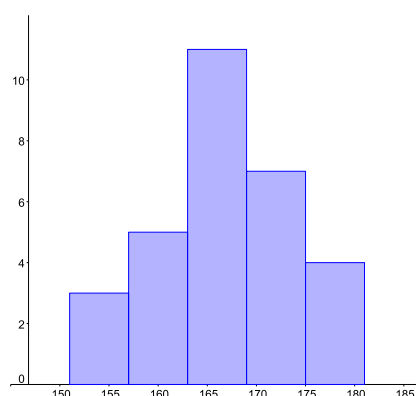
Von 30 (fiktiven) Studentinnen der Pharmazie wurden die Körperlängen (in cm) gemessen:

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

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

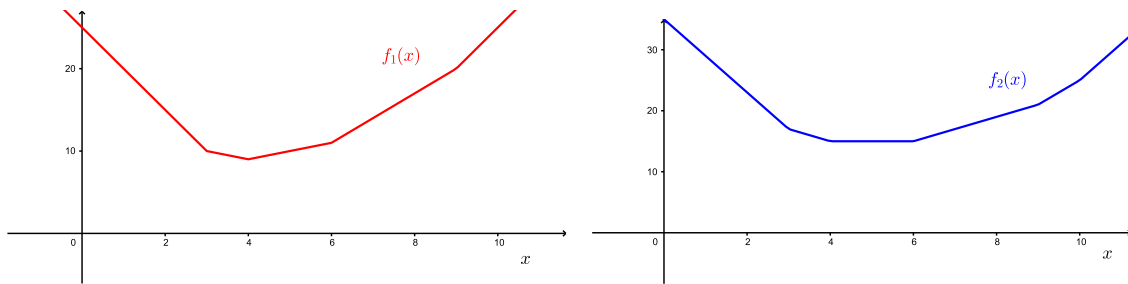
ist minimal für $x = \tilde{x}$.

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 . $\implies \tilde{x}_{0,25} \leq 3$

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. $\implies \tilde{x}_{0,25} \geq 3$

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 82 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, 1165, 1368, 676, 813, 849, 1081, 1342, 1425, 1597, 727, 1859, 1197, 761, 1019, 1978, 647, 795, 1050, 1573, 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
1	605
2	647
3	650
4	676
5	686
6	692
7	715
8	718
9	719
10	723
11	727
12	727
13	736
14	750
15	761
16	778
17	787
18	795
19	806
20	813
21	828

i	x_i
22	849
23	850
24	872
25	878
26	894
27	899
28	915
29	921
30	924
31	928
32	930
33	934
34	944
35	945
36	949
37	968
38	1019
39	1030
40	1030
41	1031

i	x_i
42	1050
43	1081
44	1089
45	1090
46	1140
47	1165
48	1197
49	1208
50	1249
51	1296
52	1304
53	1334
54	1336
55	1342
56	1368
57	1383
58	1425
59	1447
60	1461
61	1489
62	1505

i	x_i
63	1510
64	1521
65	1523
66	1537
67	1539
68	1573
69	1597
70	1629
71	1650
72	1691
73	1699
74	1727
75	1734
76	1855
77	1859
78	1878
79	1952
80	1978
81	2013
82	2024

Den Median können wir direkt ablesen.

Für die Quartile wenden wir Satz 1.1 an.

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

2. Wir betrachten die erzielten Punkte an der Prüfung Mathematik I vom 27.01.23. Es gab 168 Prüfungsteilnehmer*innen, wir haben also 168 ungeordnete Zahlen $x_1, \dots, x_{168}$, wobei jede Zahl x_i die Anzahl der erzielten Punkte der Person i angibt. Um die Quartile zu berechnen, ordnen wir zuerst diese 168 Zahlen der Grösse nach.

Für den Median $\tilde{x} = \tilde{x}_{0,5}$ rechnen wir $168 \cdot 0,5 = 84$. Der erste Punkt von Satz 1.1 sagt nun, dass der Median gleich dem arithmetischen Mittel der 84. und der 85. geordneten Zahl ist. Diese geordneten Zahlen sind beide gleich 33,5. Es gilt also $\tilde{x} = 33,5$ (Punkte).

Für das erste Quartil rechnen wir $168 \cdot 0,25 = 42$. Wieder der erste Punkt von Satz 1.1 sagt, dass $\tilde{x}_{0,25}$ gleich dem Durchschnitt der 42. und der 43. geordneten Zahl ist. Wie vorher sind diese beiden geordneten Zahlen gleich, nämlich 24,5. Es gilt also $\tilde{x}_{0,25} = 24,5$ (Punkte).

Für das dritte Quartil rechnen wir $168 \cdot 0,75 = 126$. Analog zum ersten Quartil ist $\tilde{x}_{0,75}$ gleich dem Durchschnitt der 126. und der 127. geordneten Zahl (40 und 40,5). Wir erhalten also $\tilde{x}_{0,75} = 40,25$ (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 $\tilde{x}_{0,25} = 24,5$ und $\tilde{x}_{0,75} = 40,25$. Und da der Median $\tilde{x} = 33,5$ 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 $x_{\min} = 4$ und die grösste Zahl $x_{\max} = 49$.

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.

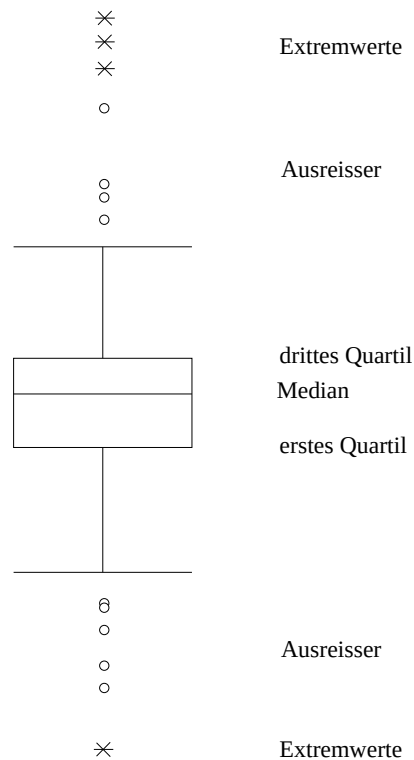
• innerhalb der Box

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.

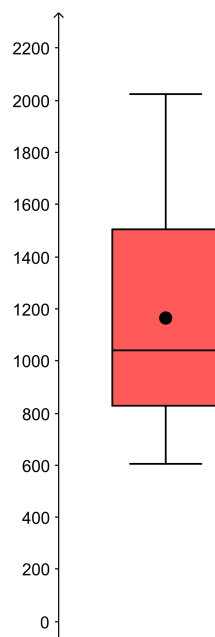
• ausserhalb der Box

- *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.

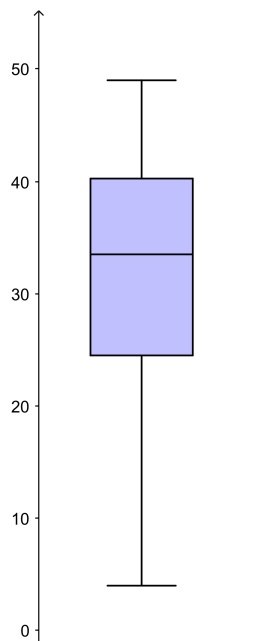


Beispiele

1. Im ersten Beispiel auf Seite 10 haben wir die folgenden Quartile berechnet: $\tilde{x}_{0,25} = 828$, $\tilde{x}_{0,5} = 1040,5$, $\tilde{x}_{0,75} = 1505$. Es gibt weder Ausreisser noch Extremwerte. Der kleinste Wert ist 605 und der grösste Wert 2024. Der Boxplot sieht wie folgt aus, wobei hier der schwarze Punkt in der Box die Lage des Mittelwerts $\bar{x} = 1164,60$ beschreibt.



2. Im zweiten Beispiel auf Seite 11 haben wir die folgenden Quartile erhalten: $\tilde{x}_{0,25} = 24,5$, $\tilde{x}_{0,5} = 33,5$, $\tilde{x}_{0,75} = 40,25$. Auch hier gibt es weder Ausreisser noch Extremwerte. Die grösste Zahl ist 49 und 4 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 HS22 haben 69 von den 168 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.

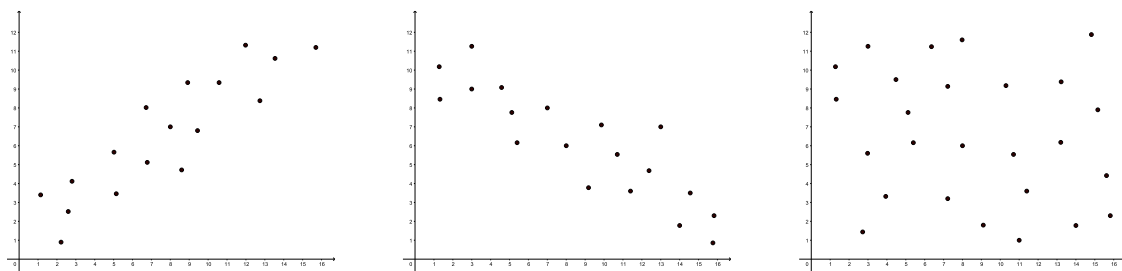
2 Korrelation und Regressionsgerade

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:

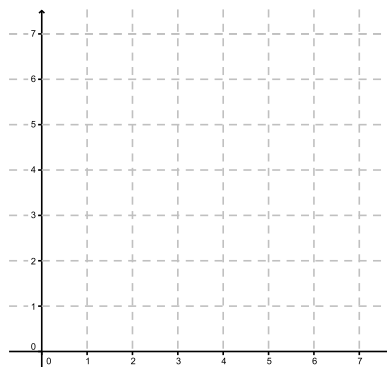
- Ist r_{xy} nahe bei 1, so gilt $y_i \approx ax_i + b$ für ein $a > 0$, das heisst, es besteht (beinahe) ein linearer Zusammenhang zwischen den Wertepaaren. Man spricht in diesem Fall von einer *starken positiven* Korrelation.
- Ist r_{xy} nahe bei -1 , so gilt $y_i \approx ax_i + b$ für ein $a < 0$, 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):

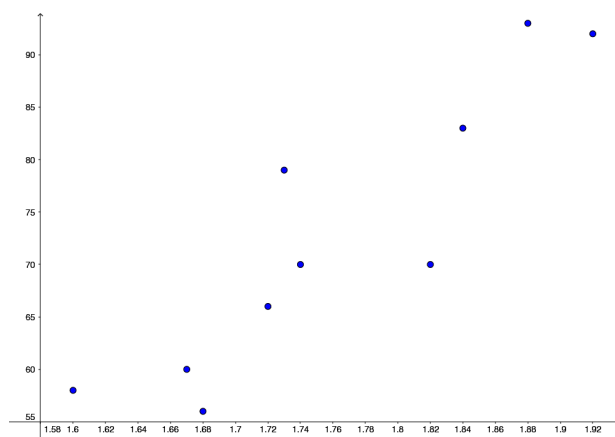
Korrelationskoeffizient:

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}

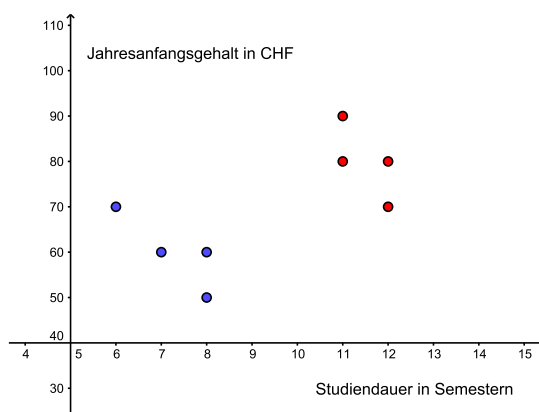
- Ist $r_{xy} \approx 0$, dann sagt dies nur, dass die beiden Datensätze keinen *linearen* Zusammenhang haben. Eventuell hängen sie jedoch quadratisch, exponentiell oder durch eine trigonometrische Funktion voneinander ab (vgl. Abschnitt 2.4).
- Falls r_{xy} nahe bei 1 oder -1 liegt, folgt lediglich, 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

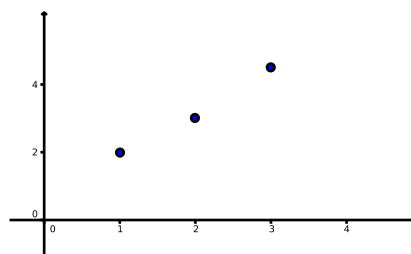
Der Umsatz einer Apotheke gibt einen wichtigen Hinweis auf ihre Wirtschaftlichkeit. Kann dieser Umsatz beispielsweise durch die Anzahl Kunden pro Tag abgeschätzt werden?

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:



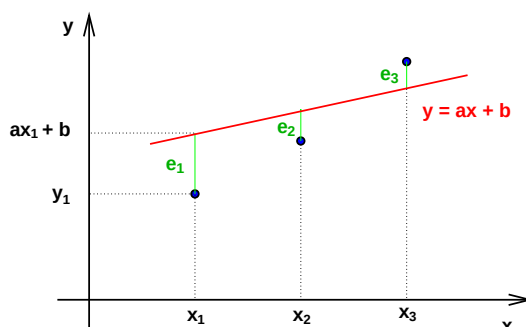
Die drei Punkte liegen fast auf einer Geraden. Es könnte also sein, dass ein linearer Zusammenhang zwischen den Messwerten x_1, x_2, x_3 und y_1, y_2, y_3 besteht, der jedoch durch verschiedene Einflüsse verfälscht wurde.

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 :

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

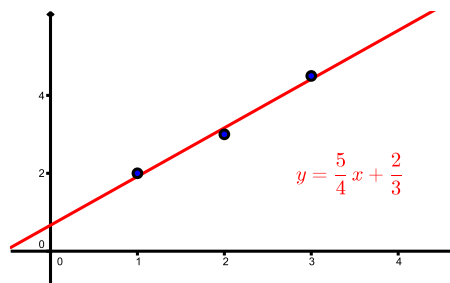
Wir werden im dritten Teil dieses Semesters lernen, dass eine notwendige Bedingung für ein Minimum das Verschwinden der Ableitungen von $F(a, b)$ nach den Variablen a und b ist:

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

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

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.

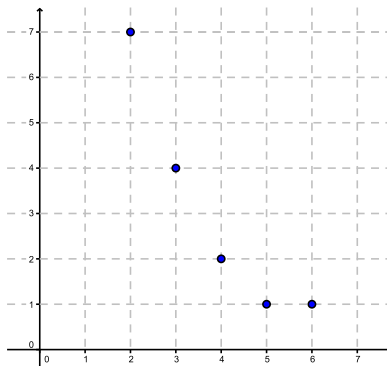
Satz 2.2 Die Regressionsgerade zu den Wertepaaren $(x_1, y_1), \dots, (x_n, y_n)$ hat die Gleichung $y = ax + b$ mit

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

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:

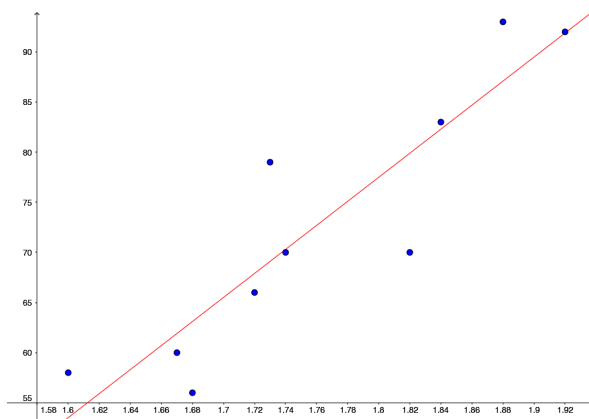
$$\bar{x} = 4, \quad \bar{y} = 3, \quad c_{xy} = \frac{-15}{4}, \quad s_x^2 = \frac{10}{4}$$

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 $(y_i - f(x_i))^2$.

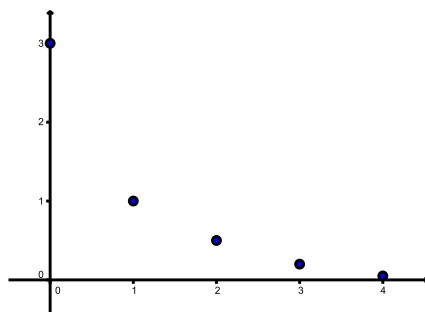
Im Fall einer Exponentialfunktion kann dieses Minimierungsproblem auf eine lineare Regression zurückgeführt werden.

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

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

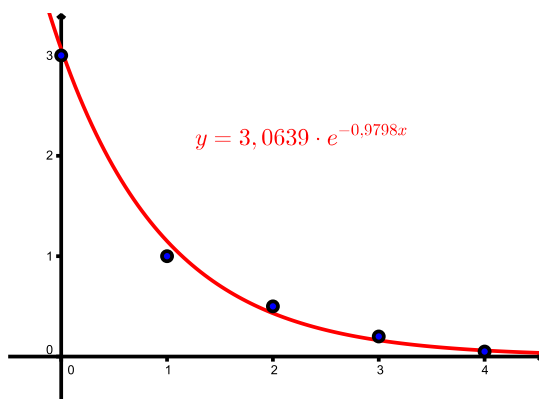
Es ist also $a = -0,9798$ und für c finden wir

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

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 $\overline{A}$ bezeichnet und entspricht dem Komplement $\overline{A} = \Omega \setminus A$.
- 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 $\overline{A}$ 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 $A =$ (gerade Augenzahl) und $B =$ (Augenzahl durch 3 teilbar) ?

2. Aus einem gut gemischten Kartenspiel von 52 Karten wird eine einzige Karte gezogen. Wie gross ist die Wahrscheinlichkeit, ein As oder die Herz-Dame zu ziehen?

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*

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

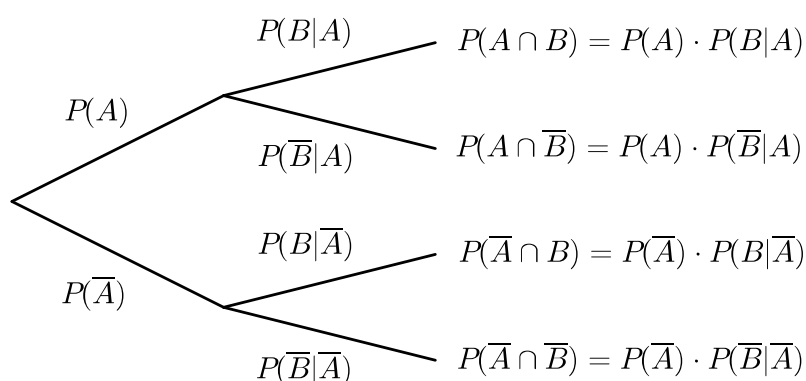
Die zweite Gleichheit im Satz folgt, indem wir die Rollen von A und B in der Definition vertauschen.

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?

Die eben berechnete Wahrscheinlichkeit ist überraschend klein! Mit Hilfe eines konkreten Zahlenbeispiels sollte diese Wahrscheinlichkeit jedoch einleuchten. Betrachten wir nämlich 10 000 zufällig ausgewählte Personen in diesem Land. Dann ist eine Person (0,01 % von 10 000) HIV positiv und diese erhält ein positives Testresultat. Von den restlichen 9 999 HIV negativen Personen erhält jedoch auch eine Person ein positives Testresultat. Damit erhalten zwei Personen ein positives Testresultat, aber nur eine davon ist HIV positiv.

Dass die betrachtete Wahrscheinlichkeit in diesem Beispiel nicht höher ist, liegt vor allem daran, dass der Anteil der HIV positiven Personen in der Bevölkerung dieses Landes sehr klein ist (0,01 %). In der Schweiz beispielsweise beträgt dieser Anteil etwa 0,2 % (Stand November 2022). Dies führt zur Wahrscheinlichkeit $P(\text{HIV pos.} \mid \text{Test pos.}) \approx 95 \%$.

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 eine ungerade Augenzahl 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 Ergebnisraum. 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

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

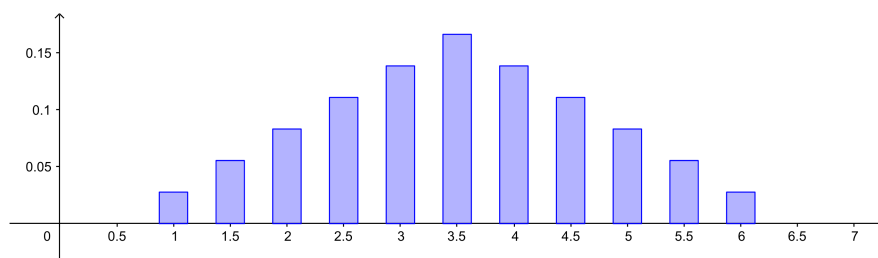
genau wie im 4. Beispiel. Die Varianz und die Streuung sind nun allerdings kleiner als im 4. Beispiel. Es gilt

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

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

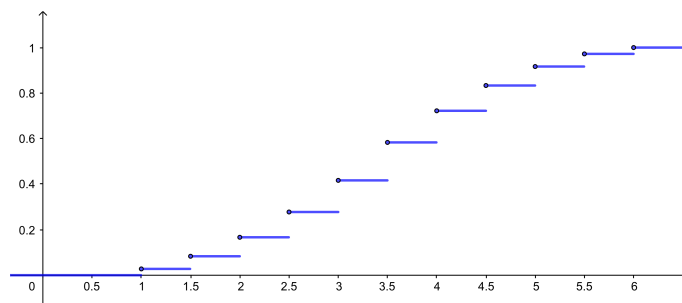
$$\{ (x_1, p_1), (x_2, p_2), (x_3, p_3), \dots \}$$

die *Wahrscheinlichkeitsverteilung* oder *Verteilung* von X . Die Abbildung $F : \mathbb{R} \rightarrow [0, 1]$ definiert durch

$$F(x) = P(X \leq x) = P(\{\omega \in \Omega \mid X(\omega) \leq x\}) = \sum_{x_k \leq x} p_k$$

heisst *Verteilungsfunktion* von X .

Eine diskrete Zufallsgrösse kann man also nicht nur mit einem Stabdiagramm sondern auch durch den Graphen der Verteilungsfunktion (eine sogenannte “Treppenfunktion”) graphisch darstellen. Hier der Graph von $F(x)$ für das letzte Beispiel:



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 (Beweis durch Nachrechnen).

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

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

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 $E(X) = 4$, $E(Y) = 2$ und

$$E(X \cdot Y) = \frac{1}{6} (3 + 9 + 10 + 15 + 5 + 6) = \frac{48}{6} = 8 = E(X) \cdot E(Y).$$

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.

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:

$$(1) \quad E(aX + bY + c) = aE(X) + bE(Y) + c$$

$$(2) \quad Var(aX + c) = a^2 Var(X)$$

Falls X und Y stochastisch unabhängig sind, gilt weiter:

$$(3) \quad E(X \cdot Y) = E(X) \cdot E(Y)$$

$$(4) \quad Var(X + Y) = Var(X) + Var(Y)$$

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

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

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 Ihr Taschenrechner keine Taste zur direkten Berechnung von Binomialkoeffizienten hat, sollten Sie den ersten 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

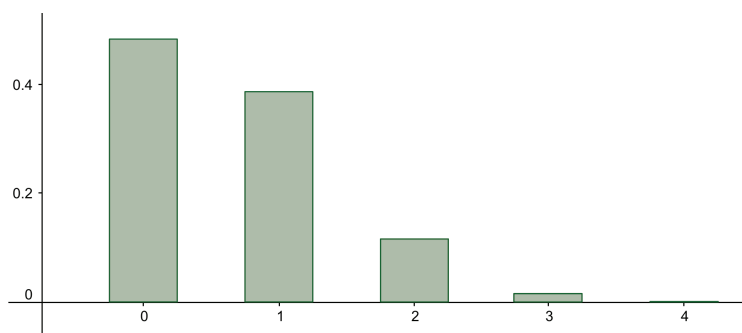
Definiert man im vorhergehenden Beispiel die Zufallsgrösse

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

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 . Kurzschreibweise: $X \sim B(n, p)$.

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.

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

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

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

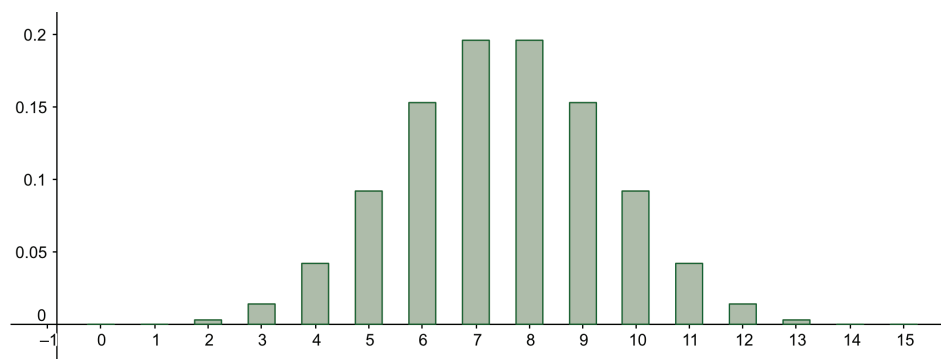
- $P(4 \leq k \leq 8) =$

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

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 dieses Jahr genau 6 Verkehrsunfälle mit derselben Ursache geben?

Die Technikanlage im Hörsaal 1 des Pharmazentrums der Universität Basel funktioniert durchschnittlich 10 Mal pro Semester nicht auf Anhieb. Mit welcher Wahrscheinlichkeit versagt die Technikanlage genau 6 Mal dieses Semester?

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 n der Experimente (Anzahl Autofahrten, bzw. Anzahl Benutzungen der Technikanlage). 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

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

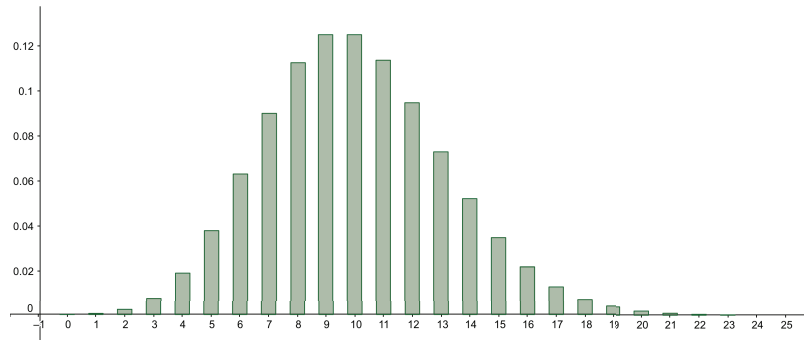
Für die beiden Beispiele finden wir also die Wahrscheinlichkeit

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*. Kurzschreibweise: $X \sim Po(\lambda)$

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?

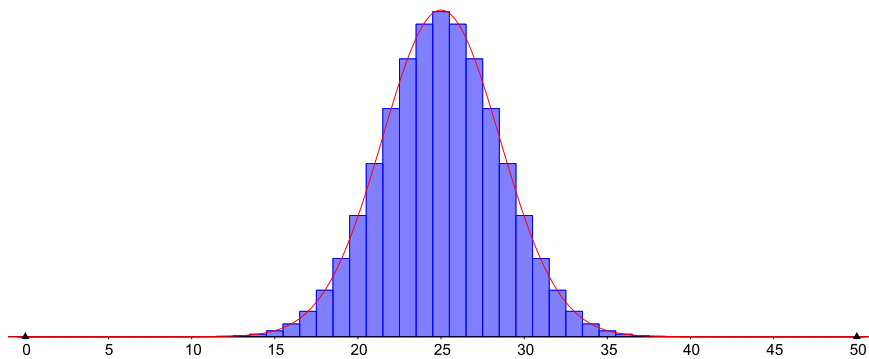
1. Exakte Berechnung mit der Binomialverteilung:

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 400 mg und 600 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*. Ist die Verteilungsfunktion $F(x) = P(X \leq x)$ einer (stetigen) Zufallsgrösse X gegeben durch

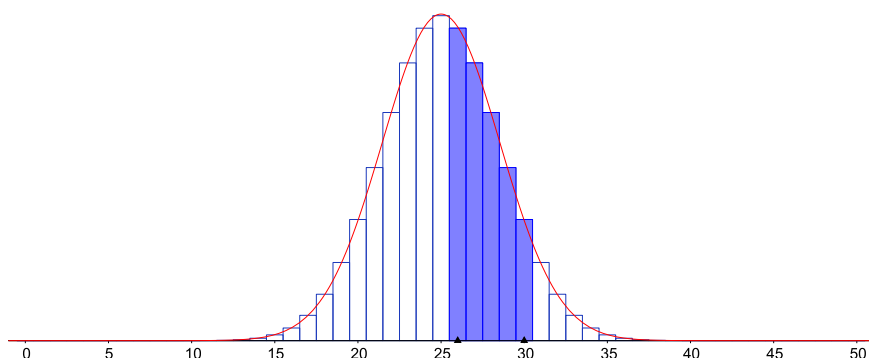
$$F(x) = P(X \leq x) = \int_{-\infty}^x f(t) dt = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2} dt,$$

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.

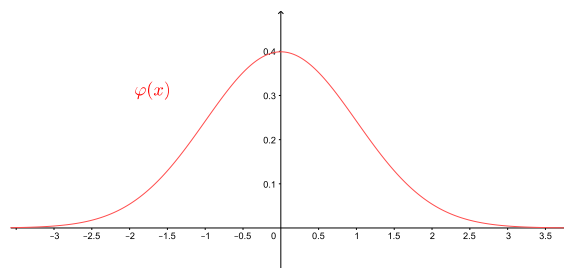
Satz 6.1 *Die Funktion $f(x)$ hat eine (lokale und globale) Maximalstelle in $x = \mu$ und zwei Wendestellen in $x = \mu \pm \sigma$.*

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 $\lim_{x \rightarrow \pm\infty} \varphi(x) = 0$.



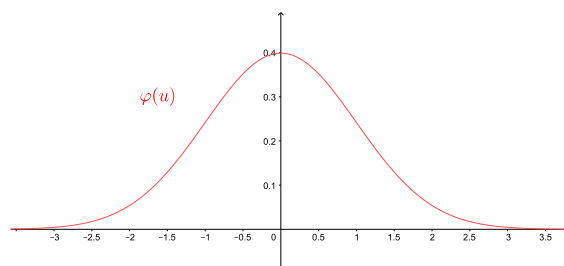
In der Tabelle (Seite 11 der Tabellen von H. Walser oder in jeder Formelsammlung) sind die Werte der Stammfunktion

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

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:

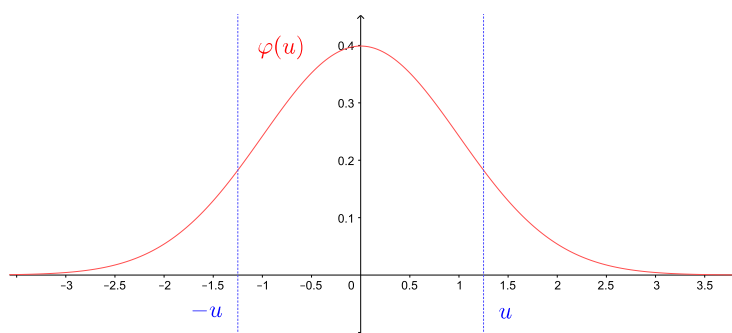
Eine weitere wichtige Eigenschaft der Standardglockenkurve ist, dass der Flächeninhalt der gesamten Fläche unter der Kurve gleich 1 ist.

Satz 6.3

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

Es folgen sofort zwei weitere Eigenschaften:

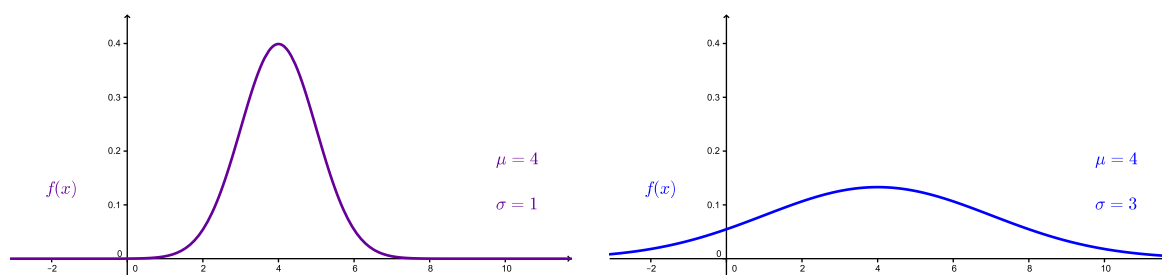
$$\begin{aligned}\Phi(0) &= \frac{1}{2} \\ \Phi(-u) &= 1 - \Phi(u)\end{aligned}$$



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

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

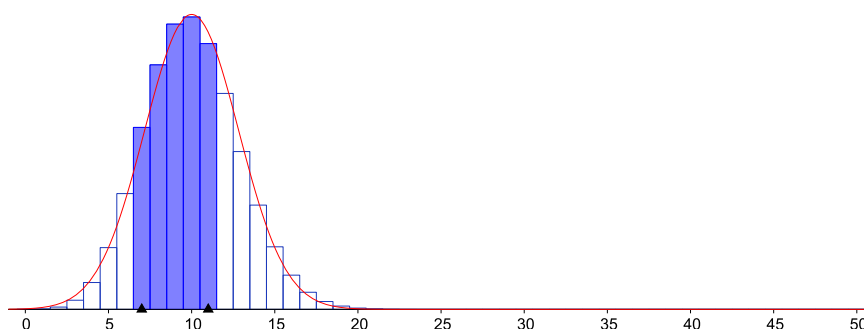
Beispiel

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 *Verteilungsfunktion* $F(x) = P(X \leq x)$ gegeben ist durch

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

Die Verteilungsfunktion ist damit eine stetige Funktion (im Gegensatz zur Verteilungsfunktion einer diskreten Zufallsgrösse). 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 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 (bzw. σ^2) sind der Erwartungswert und die Standardabweichung (bzw. Varianz) der Verteilung. Kurzschreibweise: $X \sim N(\mu, \sigma^2)$.

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, das heisst, wenn $Z \sim N(0, 1)$. 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 %	68,27 %	95,45 %	99,73 %	≈ 100 %

Diese Tabelle liest sich nun so: Die Wahrscheinlichkeit, dass eine $N(\mu, \sigma^2)$ -verteilte Zufallsgrösse 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 Überlagerung 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 $\sigma^2 = \text{Var}(X_i) = \frac{1}{4}$. Dann ist

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

die Anzahl Zahl bei n Würfen 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

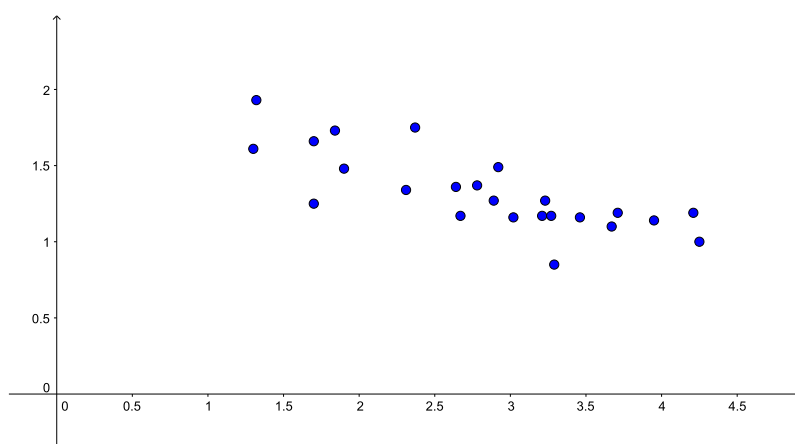
Blätter von Bäumen, aus denen ein bestimmter Wirkstoff gewonnen werden kann, sollen geerntet werden. Wir überlegen uns, ob der Wirkstoffgehalt in einem Blatt davon abhängt, wie hoch das Blatt am Baum hängt. Wenn nicht, könnte man einfach die leicht zugänglichen Blätter in niedriger Höhe ernten, ohne 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_{krit}^2 aus der Tabelle für den Freiheitsgrad 1.

Entscheid: $\chi^2 > \chi_{\text{krit}}^2 \implies A$ 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 HS22 aufgelistet (ohne Wiederholer*innen), getrennt nach Studiengang und Resultat:

	Chemie	Bio	Geo	Pharma	Total
Note ≥ 5	10	6	9	41	66
Note < 5	21	9	14	43	87
Total	31	15	23	84	153

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	13,37	6,47	9,92	36,24	66
Note < 5	17,63	8,53	13,08	47,76	87
Total	31	15	23	84	153

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

	Chemie	Bio	Geo	Pharma
Note ≥ 5	-3,37	-0,47	-0,92	4,76
Note < 5	3,37	0,47	0,92	-4,76

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

$$\chi^2 = 2,80.$$

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. (Auch auf dem 5 %-Niveau müssen wir die Nullhypothese beibehalten, denn $\chi_{\text{krit},5\%}^2 = 7,81$.)

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

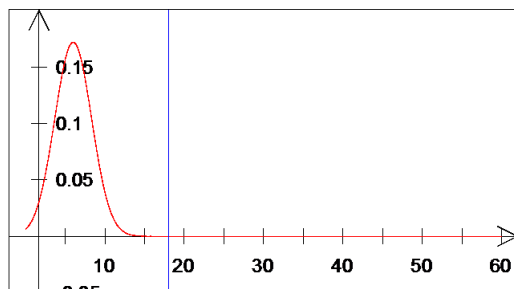
$$h_{\text{krank}} = \frac{18}{60} = 0,3.$$

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?

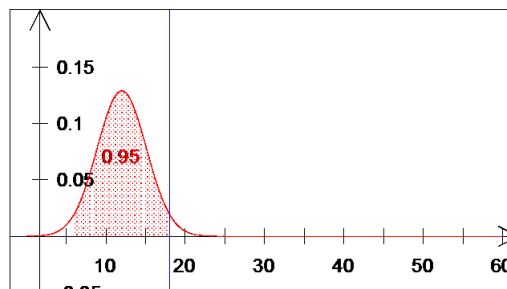
Gesucht ist ein *zweiseitiges Vertrauensintervall* für p_{krank} zum Niveau $1 - \alpha = 95\%$.

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.

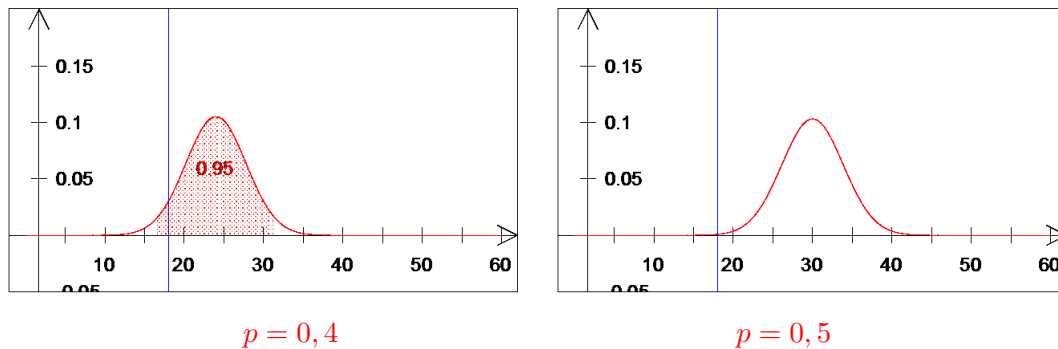
Wir können verschiedene Werte für p ausprobieren, um eine Idee zu bekommen.



$$p = 0,1$$



$$p = 0,2$$



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.