

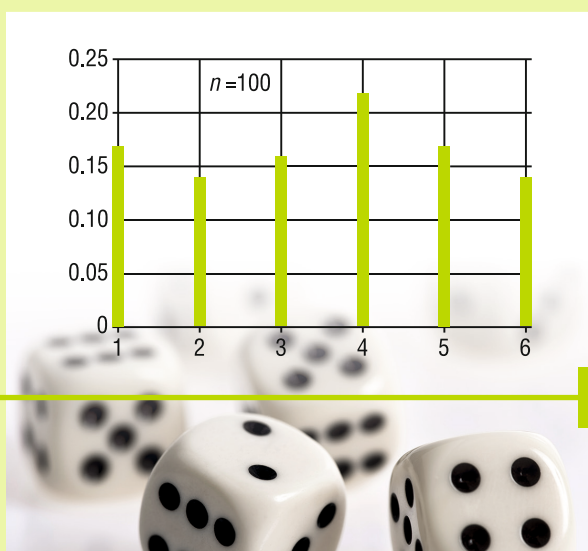
Michael Sachs



Mathematik-Studienhilfen

Wahrscheinlichkeitsrechnung und Statistik

für Ingenieurstudierende an Hochschulen



6., aktualisierte Auflage

HANSER

Sachs

Wahrscheinlichkeitsrechnung und Statistik



Ihr Plus – digitale Zusatzinhalte!

Auf unserem Download-Portal finden Sie zu diesem Titel kostenloses Zusatzmaterial. Geben Sie dazu einfach diesen Code ein:

plus-oq39k-1wt9n

plus.hanser-fachbuch.de



Bleiben Sie auf dem Laufenden!

Hanser Newsletter informieren Sie regelmäßig über neue Bücher und Termine aus den verschiedenen Bereichen der Technik. Profitieren Sie auch von Gewinnspielen und exklusiven Leseproben. Gleich anmelden unter

www.hanser-fachbuch.de/newsletter

Mathematik–Studienhilfen

Herausgegeben von

Prof. Dr. Bernd Engelmann

Hochschule für Technik, Wirtschaft und Kultur Leipzig,

Fachbereich Informatik, Mathematik und Naturwissenschaften

Zu dieser Buchreihe

Die Reihe Mathematik-Studienhilfen richtet sich vor allem an Studierende technischer und wirtschaftswissenschaftlicher Fachrichtungen an Fachhochschulen und Universitäten.

Die mathematische Theorie und die daraus resultierenden Methoden werden korrekt aber knapp dargestellt. Breiten Raum nehmen ausführlich durchgerechnete Beispiele ein, welche die Anwendung der Methoden demonstrieren und zur Übung zumindest teilweise selbstständig bearbeitet werden sollten.

In der Reihe werden neben mehreren Bänden zu den mathematischen Grundlagen auch verschiedene Einzelgebiete behandelt, die je nach Studienrichtung ausgewählt werden können. Die Bände der Reihe können vorlesungsbegleitend oder zum Selbststudium eingesetzt werden.

Bisher erschienen:

Dobner/Engelmann, *Analysis 1*

Dobner/Engelmann, *Analysis 2*

Dobner/Dobner, *Gewöhnliche Differenzialgleichungen*

Gramlich, *Lineare Algebra*

Gramlich, *Anwendungen der Linearen Algebra*

Knorrenschild, *Numerische Mathematik*

Knorrenschild, *Vorkurs Mathematik*

Martin, *Finanzmathematik*

Nitschke, *Geometrie*

Preuß, *Funktionaltransformationen*

Sachs, *Wahrscheinlichkeitsrechnung und Statistik*

Stingl, *OperationsResearch—Lineare Optimierung*

Tittmann, *Graphentheorie*

Michael Sachs

Wahrscheinlichkeitsrechnung und Statistik

für Ingenieurstudierende an Hochschulen

6., aktualisierte Auflage

HANSER

Autor:

Prof. Dr. Michael Sachs

Hochschule München

Fakultät für angewandte Naturwissenschaften und Mechatronik



Alle in diesem Buch enthaltenen Informationen wurden nach bestem Wissen zusammengestellt und mit Sorgfalt geprüft und getestet. Dennoch sind Fehler nicht ganz auszuschließen. Aus diesem Grund sind die im vorliegenden Buch enthaltenen Informationen mit keiner Verpflichtung oder Garantie irgendeiner Art verbunden. Autor(en, Herausgeber) und Verlag übernehmen infolgedessen keine Verantwortung und werden keine daraus folgende oder sonstige Haftung übernehmen, die auf irgendeine Weise aus der Benutzung dieser Informationen – oder Teilen davon – entsteht. Ebenso wenig übernehmen Autor(en, Herausgeber) und Verlag die Gewähr dafür, dass die beschriebenen Verfahren usw. frei von Schutzrechten Dritter sind. Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen usw. in diesem Werk berechtigt auch ohne besondere Kennzeichnung nicht zu der Annahme, dass solche Namen im Sinne der Warenzeichen- und Markenschutz-Gesetzgebung als frei zu betrachten wären und daher von jedermann benutzt werden dürften.

Bibliografische Information der Deutschen Nationalbibliothek:

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet über <http://dnb.d-nb.de> abrufbar.

Dieses Werk ist urheberrechtlich geschützt.

Alle Rechte, auch die der Übersetzung, des Nachdruckes und der Vervielfältigung des Buches, oder Teilen daraus, vorbehalten. Kein Teil des Werkes darf ohne schriftliche Genehmigung des Verlages in irgendeiner Form (Fotokopie, Mikrofilm oder ein anderes Verfahren) – auch nicht für Zwecke der Unterrichtsgestaltung – reproduziert oder unter Verwendung elektronischer Systeme verarbeitet, vervielfältigt oder verbreitet werden.

© 2021 Carl Hanser Verlag München;

Internet: www.hanser-fachbuch.de

Lektorat: Dipl.-Ing. Natalia Silakova-Herzberg

Herstellung: Anne Kurth

Satz: Michael Sachs

Titelbild: © stock.adobe.com/Claudia Paulussen

Covergestaltung: Max Kostopoulos

Coverkonzept: Marc Müller-Bremer, www.rebranding.de, München

Druck und Binden: Friedrich Pustet GmbH & Co. KG, Regensburg

Printed in Germany

Print-ISBN: 978-3-446-46943-3

E-Book-ISBN: 978-3-446-46962-4

Vorwort zur 6. Auflage

Für die 6. Auflage habe ich zahlreiche Beispiele und Aufgaben mithilfe aktueller Veröffentlichungen auf den neuesten Stand gebracht und hartnäckige Druckfehler beseitigt. Neu aufgenommen wurden Boxplots und aus aktuellem Anlass die Begriffe Sensitivität und Spezifität von Schnelltests. In der Wahrscheinlichkeitsrechnung habe ich die bedingten Wahrscheinlichkeiten vorgezogen und daraus den Begriff der Unabhängigkeit von Ereignissen abgeleitet, weil dieser Zugang nach meiner langjährigen Erfahrung den Studierenden plausibler ist. Ich danke allen Leserinnen und Lesern, die mir Verbesserungshinweise gegeben haben, sowie den Mitarbeiterinnen des Hanser Verlags für die stets angenehme Zusammenarbeit.

München, im März 2021

Michael Sachs

Aus dem Vorwort zur 5. Auflage

„Wozu brauchen wir das alles?“

So lautet eine häufig gestellte Frage von Studierenden, besonders an Hochschulen für angewandte Wissenschaften, wo immer der Aspekt der Anwendung im Vordergrund steht. Der vorliegende Band aus der Reihe „Studienhilfen Mathematik“ versucht, für den Bereich der Statistik eine befriedigende Antwort auf diese Frage zu geben. In allen technischen Studiengängen machen Studierende bei der Durchführung von Versuchen die Erfahrung von zufälligen Einflüssen. Die Erforschung von deren Gesetzen ist Gegenstand der **Wahrscheinlichkeitsrechnung**, ihre Anwendung Gegenstand der **schließenden Statistik**. Am Anfang steht ein Kapitel über **beschreibende Statistik**, das vollkommen ohne den Wahrscheinlichkeitsbegriff auskommt.

Es kommen keine schwierigen Beweise oder umfangreiche Theorien vor. Ingenieurinnen und Ingenieure dürfen sich hier getrost auf die gesicherten Ergebnisse der Mathematik verlassen und sollen vielmehr die Aussagen verstehen und richtig einschätzen lernen, die hinter solchen Sätzen stehen, sowie die Methoden sinnvoll anwenden und ihre Ergebnisse korrekt interpretieren können. Aus der höheren Mathematik werden Kenntnisse der elementaren Funktionen einer reellen Veränderlichen und ihrer Ableitungen sowie des Riemannschen Integrals vorausgesetzt.

Jeder Abschnitt stellt in Lehrsätzen und gelösten Aufgaben die Hilfsmittel bereit, mit denen man die anschließenden Übungsaufgaben lösen kann. Die Ergebnisse der Aufgaben sind zur Kontrolle im Anhang angegeben.

Inhaltsverzeichnis

1	Wozu Statistik?	7
2	Beschreibende Statistik	10
2.1	Grundbegriffe	10
2.2	Eindimensionale Häufigkeitsverteilungen	15
2.3	Kumulierte Häufigkeiten und empirische Verteilungsfunktion	20
2.4	Lageparameter	26
2.5	Streuungsparameter	39
2.6	Zweidimensionale Häufigkeitsverteilungen	47
2.7	Korrelationsrechnung	54
2.8	Regressionsrechnung	59
3	Wahrscheinlichkeitsrechnung	67
3.1	Kombinatorische Grundlagen	67
3.2	Zufall, Ereignisalgebra	69
3.3	Wahrscheinlichkeit und Satz von Laplace	77
3.4	Bedingte Wahrscheinlichkeiten und unabhängige Ereignisse	84
3.5	Zufällige Variable und Wahrscheinlichkeitsverteilungen	94
3.6	Erwartungswert und Varianz einer Verteilung	104
3.7	Wichtige diskrete Verteilungen	114
3.8	Die Normalverteilung	123
4	Schließende Statistik	135
4.1	Problemstellung, Zufallsstichproben	135
4.2	Punktschätzungen	137
4.3	Intervallschätzungen	148
4.4	Hypothesentests	165
A	Tabellen	182
B	Lösungen der Übungsaufgaben	185
	Literaturverzeichnis	196
	Sachwortverzeichnis	197

1 Wozu Statistik?

„Statistik informiert. Alle.“

Mit diesen Worten beginnt eine Broschüre aus dem Jahr 1990, herausgegeben vom Bayerischen Landesamt für Statistik und Datenverarbeitung¹. Die dort getroffenen Aussagen gelten vorwiegend für den wirtschaftlichen und sozialen Bereich, lassen sich jedoch unschwer auch auf das Arbeitsumfeld von Ingenieurinnen und Ingenieuren übertragen. In einer sich rasch verändernden Welt, in der Massendaten anfallen und mit Computerhilfe auch schnell verarbeitet werden können, ist es wichtig

- zu wissen, welchen Zustand wir heute erreicht haben,
- mit gestern zu vergleichen,
- zukünftige Zustände abzuschätzen,
- Aussagen zu machen über die Auswirkungen getroffener Maßnahmen.

Diese Aufgaben müssen gelöst werden auf der Basis meist riesiger Datenmengen, die wegen ihres Umfanges keinem Entscheidungsträger unbearbeitet vorgelegt werden können. Aufgabe der **beschreibenden** oder **deskriptiven Statistik** (Kapitel 2) ist es daher, große und unübersichtliche Datenmengen so aufzubereiten, dass wenige aussagekräftige Kenngrößen und/oder Grafiken entstehen, mit denen dann die genannten Fragestellungen gelöst werden können. In diesen Kenngrößen ist dann die gesamte Datenmenge gewissermaßen „fokussiert“.

Beispiel 1.1

Eine Zeugnisnote in einem Diplomzeugnis errechnet sich aus sehr vielen Einzelnoten, die evtl. sogar mit unterschiedlicher Gewichtung in die Gesamtnote eingehen. Die beschreibende Statistik liefert die Formel, mit der man aus den Einzelnoten die Gesamtnote berechnet. ■

Die zur Berechnung der Kenngrößen zugrunde gelegte Datenmenge kann zwar sehr groß werden, ist aber prinzipiell immer endlich und spricht nur für sich selbst. Es werden also keinerlei weiterführende Rückschlüsse gezogen auf irgendwelche Einheiten, die nicht untersucht wurden. Bereits im 19. Jahrhundert hat man jedoch erkannt, dass sehr große statistische Massen, wie z. B. die Bevölkerung eines Landes, nicht mehr durch eine vollständige Erhebung zu erfassen sind. Der Aufwand hierfür wäre technisch und organisatorisch viel

¹heute: Bayerisches Landesamt für Statistik

zu hoch. Während hier eine Totalerhebung zwar schwierig, aber prinzipiell immerhin noch möglich wäre, verhält es sich ganz anders, wenn man Aussagen treffen will z. B. über alle Teile, die eine bestimmte Maschine jemals produziert, oder gar über die Menge aller möglichen Würfe mit einem bestimmten Würfel. In beiden Fällen ist die Grundgesamtheit, also die Menge aller Teile bzw. aller Würfe, von vorneherein überhaupt nicht bekannt, sie liegt nicht so greifbar vor uns wie im Falle der Noten eines Faches.

In solchen Fällen muss man sich auf **Stichproben**, also Teilerhebungen beschränken. Die gesuchten Kenngrößen können dann nicht mehr exakt bestimmt, sondern nur noch geschätzt werden. Hier kommt aber der Begriff des **Zufalls** ins Spiel: Ziehen zwei Leute eine Stichprobe aus der gleichen Grundgesamtheit und berechnet jeder das arithmetische Mittel $\bar{x}$ seiner Stichprobe, so werden sie im Allgemeinen zu unterschiedlichen Ergebnissen kommen. Welche Elemente in die jeweilige Stichprobe gelangen, hängt nämlich von Faktoren ab, die der Stichprobenzieher nicht bestimmen kann.

Kapitel 4 ist daher der **schließenden** oder **induktiven Statistik** gewidmet, die Methoden entwickelt, um von einer Stichprobe auf die Gesamtheit schließen zu können.

Beispiel 1.2

Eine große Firma kauft 100 Personal Computer bei Händler *A* und 200 bei Händler *B*. Im Laufe des ersten Jahres gehen bei den *A*-Computern zwölf Netzteile kaputt, bei den *B*-Computern 16 Netzteile. Der Beitrag der beschreibenden Statistik ist hier eher gering: Sie gibt uns die Formel, um die Anteile p_A und p_B der Computer mit defektem Netzteil zu berechnen:

$$p_A = \frac{12}{100} = 12 \% \quad \text{und} \quad p_B = \frac{16}{200} = 8 \%.$$

A hat also schlechtere Qualität in Bezug auf das Merkmal „Lebensdauer des Netzteils“ geliefert als *B*. Aber kann man daraus schließen, dass *A* **generell** einen höheren Ausschussanteil hervorbringt als *B*, oder könnte das schlechte Ergebnis für *A* auch einfach Zufall gewesen sein? Diese Frage kann nur die schließende Statistik beantworten. Wir werden auf dieses Beispiel im Abschnitt 4.3 zurückkommen. ■

Beispiel 1.3

Um eine Prognose über den Ausgang einer Wahl treffen zu können, werden 100 ausgewählte Personen befragt. 45 antworten, sie würden die *A*-Partei wählen. Die beschreibende Statistik kann dann nur sagen,

dass 45 % der Befragten *A*-Anhänger sind. Die schließende Statistik dagegen versucht, ausgehend von diesem Ergebnis Aussagen über **alle** Wahlberechtigten zu machen (die berühmte Hochrechnung). Kann man ausgehend von der Stichprobe mit Umfang 100 schon sagen, dass Partei *A* unter allen Wahlberechtigten 45 % der Stimmen hat? Rein gefühlsmäßig erscheinen 100 als zu wenig, aber mit welcher Genauigkeit kennen wir den wahren Anteil der *A*-Wähler, wenn wir 1 000 oder gar 10 000 Personen befragen? Hängt die Anzahl der zu befragenden Personen von der Gesamtzahl der Wahlberechtigten ab oder nicht? Was heißt überhaupt in diesem Zusammenhang „Genauigkeit“? Dies sind Fragen, die die schließende Statistik präzisieren und beantworten muss. ■

Als Grundlage für die schließende Statistik müssen wir also zunächst den Begriff des Zufalls präzisieren. Hierfür steht eine mathematische Theorie zur Verfügung, die **Wahrscheinlichkeitsrechnung**. Ihre wichtigsten Erkenntnisse stehen in Kapitel 3.

Beschreibende und schließende Statistik werden auch zusammengefasst unter dem Begriff **statistische Methodenlehre**. Darüber hinaus wird der Begriff **Statistik** aber auch angewendet für die einzelne Tabelle, („die Arbeitslosenstatistik“, „die Unfallstatistik“). Schließlich meint Statistik auch noch die praktische Tätigkeit des Registrierens und Auswertens von Daten, z. B. in dem Satz „Die amtliche Statistik in Bayern wird vom Landesamt für Statistik wahrgenommen“.

Wir fassen zusammen:

Definition 1.1

Die statistische Methodenlehre ist eine Hilfswissenschaft. Ihre Aufgabe ist es, Methoden für die Erhebung, Aufbereitung, Analyse und Interpretation von Daten bereitzustellen mit dem Ziel, Strukturen in Massenerscheinungen zu erkennen.

2 Beschreibende Statistik

2.1 Grundbegriffe

Die **statistische Masse** besteht aus allen denjenigen Elementen, denen das Interesse des Statistikers gilt. Ihre einzelnen Mitglieder heißen **statistische Einheiten**, **statistische Elemente** oder **Merkmalsträger**. Vor jeder ernst zu nehmenden statistischen Tätigkeit muss die statistische Masse in räumlicher, zeitlicher und sachlicher Hinsicht präzise definiert werden.

Beispiel 2.1

Mögliche statistische Massen sind:

- Natürliche Personen (z. B. Studierende, Kunden, Mitarbeiter),
- Sachen (z. B. Maschinen, Produkte),
- Institutionen, Körperschaften (z. B. Betriebe, Städte, Länder),
- Ereignisse (z. B. Maschinenausfälle, Geburten, Todesfälle, Firmengründungen, Konkurse).

Eine korrekt beschriebene statistische Masse wäre etwa: „Alle Studierenden des Studienganges Bioingenieurwesen an der Hochschule München im WS 2017/2018“. Räumliche Eingrenzung: Hochschule München. Zeitliche Eingrenzung: WS 2017/2018. Sachliche Eingrenzung: Studierende des Bioingenieurwesens. ■

An den statistischen Elementen interessieren uns bestimmte Eigenschaften oder **Merkmale**. Die Werte, die ein Merkmal annehmen kann, heißen **Merkmalsausprägungen**. Ein sinnvolles Merkmal muss mindestens zwei verschiedene Ausprägungen haben.

Beispiel 2.2

Eine Computerzeitschrift testet verschiedene Drucker. Die statistische Masse sind also die getesteten Drucker. Bei jedem Drucker interessieren den Leser folgende Merkmale:

- Hersteller (alle Herstellernamen),
- Bezeichnung des Gerätes (alle Bezeichnungen),
- Preis (0 € – 10 000 €),
- Gewicht (0 kg – 10 kg),

- Seitenzahl pro Minute (0 – 100 Seiten),
- Gesamturteil (sehr gut, gut, mittel, schlecht, sehr schlecht).

In Klammern stehen jeweils die möglichen Merkmalsausprägungen. ■

Bezüglich der Art ihrer Merkmalsausprägungen (Werte) lassen sich Merkmale wie folgt in Kategorien einteilen:

- Ein **qualitatives Merkmal** liegt vor, wenn die Werte keine physikalische Einheit brauchen. Man unterscheidet dabei **qualitativ-nominale** Merkmale von **qualitativ-ordinalen** Merkmalen: Bei einem nominalen Merkmal sind die Merkmalsausprägungen nur dem Namen nach unterscheidbar, sie drücken aber keinerlei Wertung oder Intensität aus. Bei einem ordinalen Merkmal¹ unterscheiden sich die Ausprägungen nicht nur hinsichtlich ihrer Namen, sondern können zusätzlich noch in eine (inhaltlich sinnvolle) Rangordnung gebracht werden.
- Ein **quantitatives Merkmal**² liegt vor, wenn die Nennung des Wertes allein noch keinen Sinn ergibt, weil die Einheit fehlt. Man unterscheidet hier zwischen **quantitativ-diskret** und **quantitativ-stetig**: Diskrete Merkmale haben Werte, die durch einen Zählprozess entstehen. Die Werte sind dann meist die natürlichen Zahlen 0, 1, 2, ..., die Einheit ist 1 oder bei größeren Anzahlen auch oft 1 000. Zwischen den einzelnen Ausprägungen können keine Werte angenommen werden. Stetige Merkmale dagegen werden durch Messung gewonnen und können (theoretisch, je nach Genauigkeit des Messgerätes) jeden Wert innerhalb eines sinnvollen Intervalls annehmen.

Beispiel 2.3

Im letzten Beispiel sind die Merkmale „Hersteller“, „Bezeichnung“ und „Gesamturteil“ qualitativ, denn ihre Werte benötigen keinerlei Einheit. „Gesamturteil“ ist darüber hinaus noch ordinal, denn die Werte können in eine Rangskala von „sehr gut“ bis „sehr schlecht“ geordnet werden. Die übrigen Merkmale sind quantitativ, „Preis“ und „Gewicht“ stetig, „Seitenzahl“ diskret. (Streng genommen ist „Preis“ ebenfalls ein diskretes Merkmal, denn Preise können keine Werte zwischen zwei aufeinanderfolgenden Cent annehmen. Hier ist die Einteilung aber so fein, dass man „Preis“ meist als stetiges Merkmal behandelt.) ■

¹auch *Rang-Merkmal*

²auch *metrisches* oder *kardinales* Merkmal

Die Kategorie eines Merkmals hat Einfluss auf die Gestaltung von Fragebögen und grafischen Bildschirmoberflächen: Bei einem qualitativen Merkmal mit nur wenigen Ausprägungen (z. B. Familienstand, Geschlecht) wird man alle Alternativen auflisten mit der Möglichkeit, die zutreffende anzukreuzen. Ein qualitatives Merkmal mit sehr vielen möglichen Ausprägungen (z. B. Name, Wohnort, Staatsangehörigkeit) wird man als Freitextfeld zum Ausfüllen realisieren. Quantitative Merkmale schließlich haben sehr viele Werte und werden daher als Freitextfeld dargestellt, oder man fasst die Werte zu wenigen Intervallen zusammen (Klassierung, siehe Abschnitt 2.2) mit der Möglichkeit anzukreuzen.

Zum Zwecke der effizienten Abspeicherung von Merkmalsausprägungen in DV-Anlagen werden die Ausprägungen oft **verschlüsselt**, d. h., den einzelnen Werten werden Zahlen zugeordnet, z. B. beim qualitativ-nominalen Merkmal Familienstand: „0“ für „ledig“, „1“ für „verheiratet“, „2“ für „verwitwet“ usw. Das heißt aber **nicht**, dass daraus jetzt ein quantitatives Merkmal entstanden ist, was man schon allein daran erkennt, dass „verwitwet“ nicht doppelt soviel ist wie „verheiratet“.

Werden die benötigten Daten durch eine eigens für statistische Zwecke organisierte Erhebung gewonnen, sprechen wir von einer **Primärstatistik**, werden sie dagegen von bereits vorhandenen Verwaltungs- oder Unternehmensdateien für die Statistik „abgezweigt“, von einer **Sekundärstatistik**.

Beispiel 2.4

Ermittelt die „Arbeitsgemeinschaft Fernsehforschung“ (AGF) die Fernseh-Einschaltquoten, so entsteht eine Primärstatistik, dasselbe gilt z. B. für die Qualitätskontrolle in Betrieben. Die monatliche Arbeitslosenstatistik der Bundesagentur für Arbeit dagegen ist eine Sekundärstatistik, da hier auf das bereits in den Arbeitsämtern vorliegende Material zurückgegriffen wird. ■

Werden alle Elemente einer statistischen Masse in die Erhebung einbezogen, liegt eine **Totalerhebung** oder **Vollerhebung** vor, ansonsten eine **Teilerhebung** oder **Stichprobe**. Elemente in Stichproben können zufällig ausgewählt werden (**Zufallsstichprobe**, sie wird im Abschnitt 4.1 behandelt), oder bewusst zusammengesetzt werden, sodass in der Stichprobe die Werte gewisser Merkmale mit den gleichen relativen Häufigkeiten („Quoten“) repräsentiert sind wie in der Gesamtheit (**repräsentative Stichprobe**). Um repräsentative Stichproben zusammensetzen zu können, braucht man Informationen über die zugrunde liegende Grundgesamtheit aus früheren Untersuchungen.

Beispiel 2.5

Die AGF führt Teilerhebungen zur Fernsehzuschauerforschung durch. Dieses Fernsehpanel umfasst ca. 5 400 täglich berichtende Haushalte, in denen rund 11 000 Personen leben (Stichprobe). Damit wird die Fernsehnutzung von 75,3 Mio. Personen ab drei Jahren (Grundgesamtheit) abgebildet (Stand 04.01.2021, Quelle: [12]). Sind also etwa in der Gesamtbevölkerung 12 % zwischen 14 und 29 Jahre alt, so müssen in einer repräsentativen Stichprobe von 10 000 Testpersonen ca. 1 200 Testpersonen zwischen 14 und 29 Jahre alt sein. ■

Für die anschließende Veröffentlichung der Daten in Form einer Tabelle gibt es die DIN-Norm 55 301 ([9]). Auch wenn sie inzwischen zurückgezogen wurde, ist es sicher nicht verkehrt, sich an einige Grundregeln dieser Norm zu halten. Bild 2.1 zeigt die vier Hauptbestandteile einer Tabelle. Der Tabellenkopf

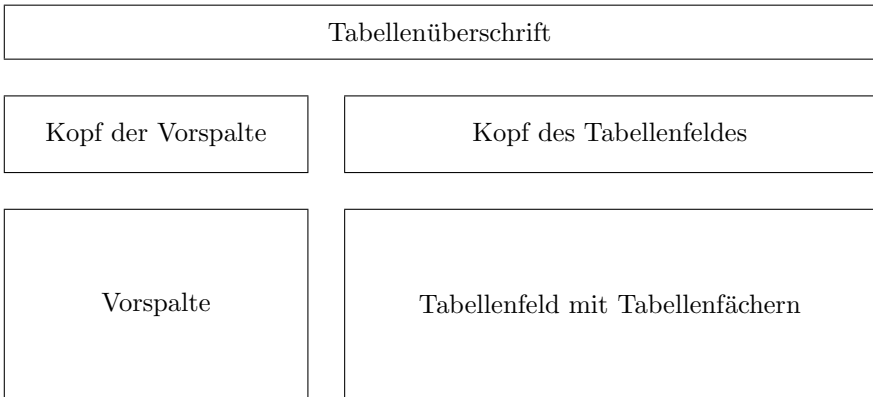


Bild 2.1: Bestandteile einer Tabelle

kann dabei Unterstrukturen enthalten. Bei quantitativ-stetigen Merkmalen müssen die Einheiten im Kopf oder in der Überschrift stehen. Tabellenfächer sollten niemals leer sein, es können aber u.a. folgende Symbole verwendet werden (siehe [9], 10.6):

- nichts vorhanden
- ... Angabe fällt später an
- / Zahlenwert nicht sicher genug
- . Zahlenwert unbekannt oder geheimzuhalten
- x Tabellenfach gesperrt, weil Aussage nicht sinnvoll

Die sachliche, räumliche und zeitliche Eingrenzung der statistischen Masse soll aus dem Titel der Veröffentlichung, der Tabellenüberschrift, der Vorspalte, dem Kopf und ggf. aus Fußnoten hervorgehen. Bild 2.2 zeigt eine nach diesen Richtlinien gestaltete Tabelle (Auszug aus [13], S. 110, Tab. 38):

Absolventen und Abgänger 2018/19 in Bayern nach Abschlussarten

Schulart	Absolventen und Abgänger insgesamt	und zwar			
		ohne Abschluss	mit		
			Abschluss der Mittelschule	dar. mit Quali ¹	allgemeine Hochschulreife
Haupt-/ Mittelschule	39 717	3 347	22 159	14 986	x
Realschule	36 865	337	642	141	x
Gymnasium	40 124	99	383	59	36 279

¹ Qualifizierender Abschluss der Mittelschule

Bild 2.2: Nach DIN 55 301 gestaltete Tabelle (Auszug)

Aufgaben

2.1 Nennen Sie zu den folgenden Merkmalsträgern (statistischen Elementen) und Merkmalen jeweils einige sinnvolle Merkmalsausprägungen und geben Sie an, um welche Merkmalskategorie es sich dabei handelt (qualitativ-nominal, qualitativ-ordinal, quantitativ-diskret, quantitativ-stetig):

Merkmalsträger	Merkmal
Personen	Hochschulstudium? (Ja/Nein)
Personen	Familienstand
Fernsehzuschauer	Gefallen an einer Sendung
Betriebe	Anzahl Mitarbeiter
Motoren	Leistung
Telefongespräche	Dauer
Klavierstücke	Schwierigkeitsgrad

2.2 Was sind die statistischen Elemente in Bild 2.2 und welche Merkmale wurden erhoben? Nennen Sie die dazugehörigen Ausprägungen. Um welche Merkmalskategorien handelt es sich? Liegt eine Primär- oder Sekundärstatistik vor, eine Teil- oder Totalerhebung?

- 2.3** Entwerfen Sie eine Tabellenstruktur, in der der Energieverbrauch in Bayern in den Jahren 2018, 2019 und 2020 dargestellt werden kann, aufgeschlüsselt nach Energieträgern (Mineralölprodukte, Gase, Strom usw.). Der Verbrauch soll sowohl in TJ (Tera-Joule) als auch in 1 000 t SKE (1 000-Tonnen-Steinkohleeinheiten) ausgewiesen werden.

2.2 Eindimensionale Häufigkeitsverteilungen

Der Begriff **eindimensional** meint die Betrachtung und Auswertung eines **einzelnen** Merkmals. Mehrere Merkmale und ihre Zusammenhänge werden in Abschnitt 2.6 behandelt.

Die **Urliste** ist diejenige Liste, die direkt bei der Datenerhebung entsteht. Sie ist meist sehr unübersichtlich und für Veröffentlichungszwecke in der Regel nicht geeignet. Um sich einen Überblick über die Verteilung der statistischen Masse auf die einzelnen Merkmalsausprägungen zu verschaffen, stellt man daher zunächst die **Häufigkeitsverteilung** des Merkmals X in Form einer unklassierten oder klassierten **Häufigkeitstabelle** dar.

Definition 2.1

Es sei n die Anzahl der statistischen Elemente. Diese seien von 1 bis n durchnummeriert, dann ist

$$x_i := \text{Ausprägung des Merkmals } X \text{ bei Element } i, \quad i = 1, \dots, n.$$

Es sei m die Anzahl der **verschiedenen** Merkmalsausprägungen von X bei einer konkreten statistischen Masse, also $m \leq n$. Dann ist

$$a_j := j\text{-te Ausprägung des Merkmals } X, \quad j = 1, \dots, m.$$

Die **absolute Häufigkeit** h_j und die **relative Häufigkeit** f_j sind erklärt durch

$$h_j := \text{Anzahl der } x_i \text{ mit Ausprägung } a_j, \quad (2.1)$$

$$f_j := \frac{h_j}{n}. \quad (2.2)$$

Relative Häufigkeiten werden oft in Prozent angegeben. DIN 55 301 empfiehlt jedoch in 10.3, bei Umfängen von $n < 100$ keine Prozentzahlen zu bilden. Meint man also acht Elemente einer statistischen Masse vom Umfang $n = 10$, so sollte man nicht von „80 %“ sprechen, da die Prozentzahl eine höhere Genauigkeit vortäuscht, als die Bezugszahl (hier 10) besitzt.

Offensichtlich müssen folgende Beziehungen gelten:

$$0 \leq h_j \leq n \quad \text{und} \quad \sum_{j=1}^m h_j = n, \quad (2.3)$$

$$0 \leq f_j \leq 1 \quad \text{und} \quad \sum_{j=1}^m f_j = 1. \quad (2.4)$$

Beispiel 2.6

Bei einer Befragung von 60 erfolgreichen Studienabsolventen wird u.a. das Merkmal X : „Anzahl Fachsemester bis zum Bachelorabschluss“ erhoben. Es entsteht folgende Urliste:

9 8 7 7 8 10 6 8 8 7 9 7
 10 8 8 9 7 8 9 10 6 10 8 9
 9 7 7 8 8 7 8 7 7 8 8 8
 10 7 10 9 8 6 9 7 8 7 9 12
 9 8 9 6 12 8 7 8 9 7 8 7

Erstellen Sie eine Tabelle der absoluten und relativen Häufigkeiten.

Lösung: Es ist $n = 60$, aber es gibt nur wenige verschiedene Ausprägungen von X , nämlich 6, 7, 8, 9, 10 und 12. Wir erstellen eine Tabelle, die in der Vorspalte diese Werte aufweist:

Semester- zahl a_j	Strichliste	Häufigkeit	
		absolut h_j	relativ f_j
6		4	0,0667
7		16	0,2667
8		20	0,3333
9		12	0,2000
10		6	0,1000
12		2	0,0333
Σ		60	1,0000

Die Häufigkeitsverteilung gibt also an, wie sich die (absoluten oder relativen) Häufigkeiten auf die verschiedenen Merkmalsausprägungen a_j verteilen. ■

Im vorigen Beispiel wurden in der Vorspalte alle erhaltenen Merkmalsausprägungen aufgelistet. Man spricht von einer **unklassierten Häufigkeitsverteilung**. Ihre grafische Darstellung liefert das **Stabdiagramm**. Die Merkmalsausprägungen werden dabei immer auf der waagerechten Achse angetragen, die Höhe der Stäbe entspricht den Häufigkeiten. Dabei ist es egal, ob man absolute oder relative Häufigkeiten betrachtet, der Unterschied drückt sich nur in der Beschriftung der senkrechten Achse aus. Bild 2.3 zeigt das Stabdiagramm zum vorigen Beispiel. Deutlich erkennt man nun eine eingipf-

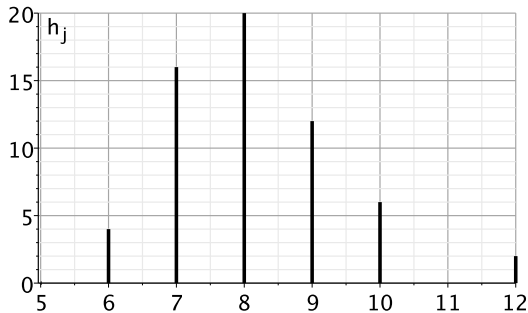


Bild 2.3: Stabdiagramm zu den Fachsemestern

lige Verteilung, der man sofort entnehmen kann, dass Studierende mit acht Semestern am häufigsten vorkommen.

Bei einem **stetigen** Merkmal erhalten wir in der Regel sehr viele verschiedene Merkmalsausprägungen, bei hinreichend großer Genauigkeit des Messgerätes ist sogar zu erwarten, dass bei allen n statistischen Elementen verschiedene Werte x_i auftreten. In diesem Fall ist es sinnvoll, die erhaltenen Merkmalsausprägungen zu Klassen der Gestalt „von ... bis unter ...“ zusammenzufassen. Es entsteht eine **klassierte Häufigkeitsverteilung**. Dabei gelten folgende Empfehlungen:

- Man vermeide nach unten und nach oben offene Klassen am Anfang und am Ende.
- Die Anzahl der Klassen sollte den Wert $\sqrt{n}$ (n =Umfang der statistischen Masse) nicht überschreiten.

Der Preis für den Gewinn an Übersichtlichkeit, den man durch die Klassenbildung erhält, ist ein Informationsverlust, denn über die Verteilung der Werte innerhalb einer Klasse ist nun nichts mehr bekannt.

Definition 2.2

Gegeben sei ein stetiges Merkmal X mit Werten im Intervall $[a; b)$ und eine Folge von reellen Zahlen

$$a = a_1 < b_1 = a_2 < b_2 = a_3 < b_3 \dots a_m < b_m = b.$$

Das Intervall

$$[a_j; b_j[= \{x | a_j \leq x < b_j\}$$

heißt j -te **Klasse** K_j . Die Klassen bilden also eine lückenlose und sich nicht überlappende Zerlegung des gesamten Wertebereiches von X . Die untere Klassengrenze gehört dabei immer zur Klasse selbst, die obere zur nächsten Klasse. Die **absolute Häufigkeit** h_j und die **relative Häufigkeit** f_j der Klasse K_j sind erklärt durch

$$h_j := \text{Anzahl der Ausprägungen } x \text{ mit } a_j \leq x < b_j, \quad (2.5)$$

$$f_j := \frac{h_j}{n}. \quad (2.6)$$

Weiterhin seien

$$d_j := b_j - a_j \text{ die } \mathbf{Klassenbreite}, \quad (2.7)$$

$$m_j := \frac{a_j + b_j}{2} \text{ die } \mathbf{Klassenmitte} \text{ und} \quad (2.8)$$

$$f_j^* := \frac{f_j}{d_j} \text{ die } \mathbf{Klassendichte} \text{ oder } \mathbf{Besetzungsdichte} \quad (2.9)$$

Der Begriff „Dichte“ steht dabei in Einklang mit der Bedeutung in der Physik: Von Dichte spricht man immer dann, wenn eine Größe durch den von ihr eingenommenen Raum (hier die Klassenbreite) dividiert wird. Ist der Raum groß, so ist die Dichte klein und umgekehrt.

Beispiel 2.7

Von 5 000 Telefongesprächen wurden an der zentralen Telefonanlage die Dauern in Minuten gemessen. Die folgende Tabelle gibt die klassierte

absolute und relative Häufigkeitsverteilung an, wobei nicht alle Klassen gleich lang sind, sowie die Klassenbreiten, -dichten und -mitten:

j	Klasse K_j von a_j min bis unter b_j min	h_j	f_j	d_j	f_j^*	m_j
1	0– 2	1 650	0,3300	2	0,1650	1,0
2	2– 4	1 111	0,2222	2	0,1111	3,0
3	4– 6	720	0,1440	2	0,0720	5,0
4	6– 8	508	0,1016	2	0,0508	7,0
5	8–10	332	0,0664	2	0,0332	9,0
6	10–15	427	0,0854	5	0,0171	12,5
7	15–30	252	0,0504	15	0,0034	22,5
Σ	x	5 000	1,0000	x	x	x

■

Die grafische Darstellung einer klassierten Häufigkeitsverteilung erfolgt mittels eines **Histogramms**. Es besteht aus soviel Rechtecken wie Klassen, deren Grundlinien auf der waagerechten Achse (der Merkmalsachse) die Klasseneinteilung wiedergeben. Die Flächeninhalte sollen dabei den relativen Häufigkeiten f_j proportional sein. Deshalb muss die Höhe von Rechteck j gleich der **Klassendichte** f_j^* sein, sonst würden Klassen mit großer Breite überproportional gewichtet, und es entstünde ein falscher optischer Eindruck. Bild 2.4 gibt ein ehrliches Histogramm zum vorigen Beispiel an in dem Sinne, dass die Flächeninhalte die relativen Klassenhäufigkeiten wiedergeben.

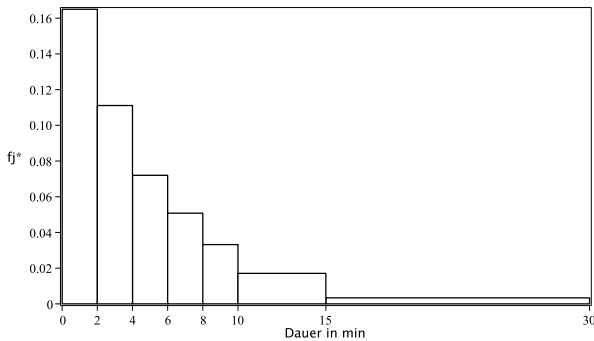


Bild 2.4: Histogramm zu den Telefongesprächen