



Statistik-Lehrbuch

Methoden der Statistik im
wirtschaftswissenschaftlichen
Bachelor-Studium

von

Univ.-Prof. Dr. Horst Degen

Heinrich-Heine-Universität Düsseldorf

apl. Prof. Dr. Peter Lorscheid

Heinrich-Heine-Universität Düsseldorf

4., korrigierte Auflage

Oldenbourg Verlag München

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.

© 2012 Oldenbourg Wissenschaftsverlag GmbH
Rosenheimer Straße 145, D-81671 München
Telefon: (089) 45051-0
www.oldenbourg-verlag.de

Das Werk einschließlich aller Abbildungen ist urheberrechtlich geschützt. Jede Verwertung außerhalb der Grenzen des Urheberrechtsgesetzes ist ohne Zustimmung des Verlages unzulässig und strafbar. Das gilt insbesondere für Vervielfältigungen, Übersetzungen, Mikroverfilmungen und die Einspeicherung und Bearbeitung in elektronischen Systemen.

Lektorat: Dr. Stefan Giesen
Herstellung: Constanze Müller
Einbandgestaltung: hauser lacour
Gesamtherstellung: Grafik & Druck GmbH, München

Dieses Papier ist alterungsbeständig nach DIN/ISO 9706.

ISBN 978-3-486-71420-3
eISBN 978-3-486-71588-0

Inhaltsverzeichnis

0	Einführung: Statistik – historische Entwicklung und heutige Arbeitsgebiete	1
0.1	Der Begriff ‚Statistik‘	1
0.2	Geschichte der Statistik	3
0.3	Arbeitsgebiete der Statistik	6
0.4	Phasen statistischen Arbeitens	7
Teil A:	Beschreibende Statistik	11
1	Grundbegriffe der Statistik	12
1.1	Statistische Massen	12
1.2	Statistische Merkmale	14
1.3	Skalierung von Merkmalen	16
1.4	Skalentransformationen und Klassenbildung	18
2	Häufigkeiten und ihre Darstellung in Tabellen und Grafiken	23
2.1	Absolute und relative Häufigkeiten	23
2.2	Kumulierte Häufigkeiten	25
2.3	Quantile von Häufigkeitsverteilungen	25
2.4	Tabellarische Darstellung von Häufigkeiten	27
2.5	Grafische Darstellung von Häufigkeiten	29
2.6	Typen von Häufigkeitsverteilungen	37
3	Statistische Maßzahlen für eindimensionale Häufigkeitsverteilungen	39
3.1	Vorbemerkungen	39
3.2	Mittelwerte	40
3.3	Steuungsmaße	46
3.4	Formmaße	53
3.5	Grafische Darstellung und statistische Messung der Konzentration	56
4	Beschreibung zweidimensionaler Häufigkeitsverteilungen	64
4.1	Zur Bedeutung mehrdimensionaler Häufigkeiten	64
4.2	Zweidimensionale und bedingte Häufigkeitsverteilungen	64
4.3	Grafische Darstellung zweidimensionaler Häufigkeiten	68
4.4	Korrelationsanalyse	71
4.5	Regressionsanalyse	83

5	Zeitreihenanalyse	89
5.1	Zum Begriff der Zeitreihe und ihrer Komponenten	89
5.2	Methoden der Komponentenbestimmung	91
5.3	Einfache Prognosetechniken	97
6	Maßzahlen des statistischen Vergleichs	100
6.1	Verhältniszahlen	100
6.2	Veränderungszahlen	105
6.3	Indexzahlen	109
Teil B: Wahrscheinlichkeitsrechnung		115
7	Wahrscheinlichkeiten	116
7.1	Zufallseignisse	116
7.2	Die klassische Definition der Wahrscheinlichkeit	117
7.3	Die Häufigkeitsdefinition der Wahrscheinlichkeit	123
7.4	Die axiomatische Definition der Wahrscheinlichkeit	125
7.5	Folgerungen aus den Axiomen	127
7.6	Abhängigkeit von Ereignissen und bedingte Wahrscheinlichkeiten	129
8	Zufallsvariablen und Verteilungen	134
8.1	Begriff der Zufallsvariablen	134
8.2	Wichtige Verteilungstypen	141
8.3	Verteilungen mehrdimensionaler Zufallsvariablen	153
8.4	Bedingte Verteilungen und Unabhängigkeit	159
8.5	Funktionen von Zufallsvariablen	161
9	Kennzahlen für Verteilungen	164
9.1	Kennzahlen für eindimensionale Verteilungen	164
9.2	Kennzahlen für zweidimensionale Verteilungen	168
9.3	Kennzahlen für Funktionen von Zufallsvariablen	169
9.4	Die Ungleichung von TSCHEBYSCHEFF	173
10	Approximation von Verteilungen	175
10.1	Das Gesetz der großen Zahl und der zentrale Grenzwertsatz	175
10.2	Faustregeln zur Approximation von Verteilungen	178
Teil C: Schließende Statistik		183
11	Grundeigenschaften von Stichproben	184
11.1	Grundbegriffe der Stichprobentheorie	184
11.2	Vor- und Nachteile von Stichprobenuntersuchungen	186

11.3	Einfache Zufallsstichproben	187
11.4	Stichprobenfunktionen und ihre Verteilungen	190
12	Punktschätzungen	194
12.1	Aufgabenstellung des Parameterschätzens	194
12.2	Qualitätseigenschaften von Schätzfunktionen	195
12.3	Gebräuchliche Schätzfunktionen und ihre Eigenschaften	198
13	Intervallschätzungen	203
13.1	Grundidee des Konfidenzintervalls	203
13.2	Spezielle Methoden der Intervallschätzung	204
13.3	Bestimmung des notwendigen Stichprobenumfangs	211
14	Signifikanztests	214
14.1	Testentscheidungen und Fehlerarten	214
14.2	Allgemeine Vorgehensweise bei Signifikanztests	216
14.3	Signifikanztests für eine einfache Zufallsstichprobe	219
14.4	Signifikanztests für verbundene Stichproben	226
14.5	Signifikanztests für mehrere unabhängige Stichproben	235
14.6	Gütefunktion und notwendiger Stichprobenumfang	246
15	Komplexere Stichprobenverfahren	249
15.1	Überblick	249
15.2	Geschichtete Stichproben	250
15.3	Klumpenstichproben	255
15.4	Hochrechnungsverfahren	258
Anhang		263
A.1	Verteilungstabellen	263
A.2	Übersichten zu Wahrscheinlichkeitsrechnung und schließender Statistik	273
A.3	Literaturverzeichnis	280
A.4	Stichwortverzeichnis	282

Vorwort zur vierten Auflage

Das vorliegende Lehrbuch richtet sich an Studierende in grundlegenden wirtschaftswissenschaftlichen Studiengängen und ist aus einem Skript zur statistischen Ausbildung im Diplom-Studiengang Betriebswirtschaftslehre an der Heinrich-Heine-Universität Düsseldorf hervorgegangen. Die zugehörigen Lehrveranstaltungen hatten zunächst einen Umfang von acht Semesterwochenstunden (SWS) Vorlesungen plus zwei SWS Übungen und wurden ergänzt um zwei freiwillige SWS Praxis am Computer.

Durch die Umstellung auf konsekutive Bachelor- und Masterstudiengänge ist ab der dritten Auflage eine Änderung der ursprünglichen Konzeption notwendig geworden. Im Gegensatz zu den meisten Lehrbüchern zur Statistik für Wirtschaftswissenschaftler enthielt das Buch in den vorigen Auflagen auch einen substanzwissenschaftlichen Teil, nämlich die Wirtschafts- und Bevölkerungsstatistik, der in die statistische Methodenlehre eingebettet war. Obwohl wir auch weiterhin der Meinung sind, dass der Bereich der Wirtschafts- und Bevölkerungsstatistik unbedingt zum statistischen Basiswissen für angehende Wirtschaftswissenschaftler gehört, mussten wir der erforderlichen Kürzung des Statistikangebotes im Bachelorstudiengang Rechnung tragen: Das Pflicht-Lehrprogramm im Fach Statistik umfasst nur noch sechs SWS Plenum und zwei SWS Gruppenarbeit. Um dem statistischen Standard-Curriculum in wirtschaftswissenschaftlichen Bachelorstudiengängen zu entsprechen, trennten wir uns schweren Herzens von der Wirtschafts- und Bevölkerungsstatistik im Pflichtprogramm und verlagerten die diesbezüglichen Lehrveranstaltungen in den Wahlpflichtbereich ab dem dritten Fachsemester.

Dementsprechend haben wir auch das vorliegende Statistik-Lehrbuch ab der dritten Auflage um den bisherigen Teil zur Wirtschafts- und Bevölkerungsstatistik gekürzt. Die drei verbleibenden Themenschwerpunkte des Lehrbuchs orientieren sich weiterhin an unserer 1994 erstmals erschienen Statistik-Aufgabensammlung¹, die – neben der Wirtschafts- und Bevölkerungsstatistik – zur statistischen Methodenlehre die Themenbereiche beschreibende (deskriptive) Statistik, Wahrscheinlichkeitsrechnung und schließende Statistik (Stichprobentheorie) abdeckt. In allen drei Themenbereichen des Lehrbuchs werden Grundkenntnisse in Wirtschaftsmathematik vorausgesetzt.

¹ H. DEGEN & P. LORSCHIED: *Statistik-Aufgabensammlung*.

Da die Aufgabensammlung parallel zum Lehrbuch verwendet werden soll, haben wir im Lehrbuch die Zahl der Beispiele stark beschränkt und keine Übungsaufgaben aufgenommen. Zur Erläuterung der Methoden benutzen wir – soweit es die Verfahren erlauben – ein durchgehendes ökonomisches Beispiel eines fiktiven Unternehmens, wobei die Berechnungen häufig auf den zuvor ermittelten Ergebnissen basieren. Wir hoffen, durch die aufeinander aufbauenden Anwendungen nicht nur den Zusammenhang zwischen den Methoden aufzuzeigen, sondern auch eine gewisse Motivation zur Vor- und Rückschau zu vermitteln. Denn detailliertes Formelwissen ist für das Verständnis der Statistik weit weniger entscheidend als das Erkennen von Zusammenhängen, Unterschieden und Gemeinsamkeiten bei den Verfahren, die sich erst aus einer Gesamtschau der Methoden ergeben.

Neu aufgenommen bzw. übertragen aus der Statistik-Aufgabensammlung haben wir am Ende dieses Buches Übersichten zur Vorgehensweise bei Konfidenzintervallschätzungen und bei Signifikanztests. Erhalten geblieben ist ein knapper Tabellenanhang mit den Werten der wichtigsten benötigten Verteilungsfunktionen.

Die Texte des Lehrbuchs sind möglichst knapp gehalten, andererseits so ausführlich, dass sie in sich verständlich sein sollten. Formeln werden grundsätzlich erläutert – meist unter Verzicht auf mathematisch-statistische Herleitungen. Dadurch ist das Lehrbuch sowohl zum Selbststudium als auch als Begleittext beim Besuch von Lehrveranstaltungen geeignet.

Nach der dritten Auflage war bereits nach kurzer Zeit eine Neuauflage erforderlich. Inhaltlich ist die vierte gegenüber der dritten Auflage unverändert. Wir haben aber die Gelegenheit genutzt, einige Fehler, insbesondere in den Querweisen innerhalb des Buches, zu beseitigen.²

Horst Degen
Peter Lorscheid

² Soweit auch in dieser vierten Auflage noch Fehler im Text auftreten sollten, können uns diese unter der e-mail-Adresse ‚peter.lorscheid@uni-duesseldorf.de‘ mitgeteilt werden.

0 Statistik – historische Entwicklung und heutige Arbeitsgebiete

0.1 Der Begriff ‚Statistik‘

Ursprünglich wurde der Begriff ‚Statistik‘ zur Beschreibung der Verhältnisse im staatlichen Gemeinwesen benutzt; Gegenstand der Statistik war somit der ‚Zustand des Staates‘. Etymologisch konnte nicht eindeutig geklärt werden, ob das Wort direkt auf das lateinische ‚status‘ (Stand, Zustand) zurückgeht oder auf die hieraus entlehnte Bezeichnung ‚Staat‘.¹

Im Laufe der Zeit gesellten sich zu dieser ursprünglichen Bedeutung des Wortes ‚Statistik‘ andere hinzu:

- *materiell*: Eine ‚Statistik‘ ist eine tabellarische oder grafische Darstellung von zahlenmäßig erhobenen Daten oder von Ergebnissen statistischer Untersuchungen bestimmter Sachverhalte.

- *instrumental*: ‚Die Statistik‘ ist die Zusammenfassung von Methoden, die zur zahlenmäßigen Untersuchung (Beschreibung, Analyse) von Massenerscheinungen dienen.

- *institutionell*: Begriffe wie ‚Arbeitsmarktstatistik‘ oder ‚Bankenstatistik‘ bezeichnen die an der Durchführung bestimmter statistischer Erhebungen beteiligten Bereiche oder Institutionen.

- *speziell*: ‚statistic‘ ist der englische Ausdruck für eine Stichprobenfunktion, der zum Teil auch im deutschen Sprachraum verwendet wird.

Definitionen von Statistik

Diese Vielzahl von Bedeutungen des Begriffes ‚Statistik‘ führte dazu, dass es kaum möglich ist, eine allgemein gültige Definition hierfür anzugeben. Jede Definition von Statistik ist darauf angewiesen, innerhalb dieser vielschichtigen Bedeutungen Schwerpunkte zu setzen. Eine Auswahl von unterschiedlichen Statistik-Definitionen bekannter Statistiker soll die verschiedenen Schwerpunkte verdeutlichen:

¹ Dies liegt nicht zuletzt daran, dass im Englischen und Französischen bis heute für beide Bedeutungen das gleiche Wort verwendet wird.

- HORST RINNE: „Statistik als wissenschaftliche Disziplin ist die Lehre von den Methoden zum Umgang mit quantitativen Informationen.“²
- PETER BOHLEY: „Statistik entsteht und besteht aus dem Zählen und Messen und dem Aufbereiten von Dingen und Phänomenen, die wiederholt oder meist sogar massenhaft auftreten.“³
- JOSEF BLEYMÜLLER, GÜNTHER GEHLERT & HERBERT GÜLICHER: „Heute wird das Wort ‚Statistik‘ im doppelten Sinn gebraucht: Einmal versteht man darunter quantitative Informationen über bestimmte Tatbestände schlechthin, wie z. B. die ‚Bevölkerungsstatistik‘ oder die ‚Umsatzstatistik‘, zum anderen aber eine formale Wissenschaft, die sich mit den Methoden der Erhebung, Aufbereitung und Analyse numerischer Daten beschäftigt.“⁴
- GÜNTER BAMBERG & FRANZ BAUR: „Zum einen beschreibt er die tabellarische oder grafische Darstellung eines konkret vorliegenden Datenmaterials, [...]. Zum anderen beschreibt der Begriff „Statistik“ die Gesamtheit der Methoden, die für die Gewinnung und Verarbeitung empirischer Informationen relevant sind.“⁵
- HEINZ-JÜRGEN PINNEKAMP & FRANK SIEGMANN: „Aufgabe der statistischen Methodenlehre ist es daher, allgemeine Grundsätze und Regeln zu formulieren, die es den jeweiligen Fachvertretern (Technikern, Medizinem, Soziologen, Ökonomen etc.) erlauben, Datensätze so zu komprimieren und darzustellen, daß sie überschaubar werden.“⁶
- WERNER NEUBAUER: „Statistik ist eine Methodik empirischer Erkenntnis von Massenerscheinungen, die quantifizierende, numerische Urteile über kategorial wohldefinierte Phänomene produziert.“⁷
- JOSEF SCHIRA: „Statistik ist die Wissenschaft vom Sammeln, Aufbereiten, Darstellen, Analysieren und Interpretieren von Fakten und Zahlen.“⁸
- JOCHEN SCHWARZE: „In unserer Sprache hat das Wort ‚Statistik‘ zwei, allerdings miteinander verwandte Bedeutungen. Einmal versteht man unter Statistik eine Zusammenstellung von Zahlen oder Daten, die bestimmte Zustände, Entwicklungen oder Phänomene beschreiben. [...] Die andere Bedeutung des Begriffs ‚Statistik‘ umfasst die Gesamtheit aller Methoden zur Untersuchung von Massenerscheinungen.“⁹

² H. RINNE: *Taschenbuch der Statistik*, S. 1.

³ P. BOHLEY: *Statistik*, S. 1.

⁴ J. BLEYMÜLLER, G. GEHLERT & H. GÜLICHER: *Statistik für Wirtschaftswissenschaftler*, S. 1.

⁵ G. BAMBERG, F. BAUR & M. KRAPP: *Statistik*, S. 1.

⁶ H.-J. PINNEKAMP & F. SIEGMANN: *Deskriptive Statistik*, S. 1.

⁷ W. NEUBAUER: *Statistische Methoden*, S. 6.

⁸ J. SCHIRA: *Statistische Methoden der VWL und BWL*, S. 13

⁹ J. SCHWARZE: *Grundlagen der Statistik I*, S. 11.

WERNER VOSS: „Die Statistik dient dazu, Entscheidungen in Fällen von Ungewißheit zu treffen.“¹⁰

Volksmund: „Es gibt drei Arten von Lügen: einfache Lügen, Notlügen und Statistiken.“¹¹

Angesichts dieser Anzahl und Vielfalt von Definitionen fällt es nicht leicht, sich für eine zu entscheiden. Alle angeführten Begriffsbestimmungen sind angemessen und heben dem jeweiligen Autor besonders wichtige Bestandteile des Begriffs ‚Statistik‘ hervor. Eine knappe Formulierung, die alle im Rahmen dieses Buches wesentlichen Elemente enthält, haben wir in einem der großen Standard-Lexika gefunden.¹²

BROCKHAUS-Lexikon: Statistik ist eine „methodische Hilfswissenschaft zur zahlenmäßigen Untersuchung von Massenerscheinungen.“

Jedes einzelne Wort dieser Definition muss man sorgfältig lesen, um den Sinngehalt der Formulierung zu erkennen: ‚Methodisch‘ bedeutet, dass die Vorgehensweise der Statistik in der planmäßigen Anwendung von Verfahren zur Lösung von Aufgaben besteht. ‚Hilfswissenschaft‘ betont, dass statistisches Arbeiten kein Selbstzweck ist, sondern stets innerhalb einer bestimmten Fachdisziplin erfolgt. ‚Zahlenmäßige Untersuchung‘ heißt, dass es bei der statistischen Arbeit vor allem um die quantitative Analyse von durch Zahlen geprägten Sachverhalten geht. ‚Massenerscheinungen‘ schließlich weist darauf hin, dass Statistik sich grundsätzlich nicht mit Einzelfällen näher befasst, sondern stets die Bearbeitung (Beschreibung, Analyse, Interpretation) von großen Datenmengen zum Ziel hat.

0.2 Geschichte der Statistik

Als (Hilfs-)Wissenschaft ist die Statistik erst im 20. Jahrhundert entstanden. Statistisches Arbeiten und Denken gibt es jedoch schon seit etwa 4000 Jahren in Form von agrar- und bevölkerungstatistischen Erhebungen. Bei der geschicht-

¹⁰ W. VOSS: *Statistische Methoden und PC-Einsatz*, S. 16.

¹¹ Dies zeugt von dem großen Misstrauen gegenüber der Statistik, das in weiten Teilen der Bevölkerung herrscht, was nicht zuletzt darauf zurückzuführen ist, dass schon bei der Definition des Begriffes große Unsicherheiten auftreten.

¹² *Der Große Brockhaus*, 16. Aufl. 1957, Bd. 11, S. 178.

lichen Entwicklung der Statistik bis ins 20. Jahrhundert hinein unterscheidet man üblicherweise vier große Phasen, die sich zum Teil zeitlich überschneiden:

Praktische Statistik

Sie ist ausgerichtet auf staatliche Phänomene und Bedürfnisse (gemeinsame Großbaustellen, Versorgungssicherung, Kriegsführung, fiskalische Zwecke). Volkszählungen zur Erfassung der arbeitsfähigen, waffentragenden und/oder steuerpflichtigen Männer waren in Ägypten seit 2500 v. Chr., in China seit 2200 v. Chr. und in Persien seit 500 v. Chr. üblich. Seit 435 v. Chr. wurden im römischen Reich Einkommens- und Vermögenslisten („Zensus“) eingeführt, die in einem Rhythmus von 5 Jahren (später 15 Jahren) erstellt wurden. Das Mittelalter war dagegen eine wenig statistikfreundliche Zeit. Es gab nur vereinzelt Zählungen (England 1086, Italien 1154, Dänemark 1231).

Universitäts- und Kathederstatistik

Es entstand eine erste Schule der Statistik (Conring 1606-1681) und damit eine Ausweitung und Systematisierung umfassender Staatsbeschreibungen in Italien, Holland und Deutschland zwischen dem 16. und 18. Jahrhundert (meist auf privater, d. h. universitärer Ebene). Statistik wird als „Lehre von den Staatsmerkwürdigkeiten“ betrachtet. Im Mittelpunkt stehen rein beschreibende Darstellungen ohne Zahlenangaben. Da er erstmals den Begriff „Statistik“ verwendete, wurde der deutsche Historiker und Jurist GOTTFRIED ACHENWALL (1719-1772) auch „Vater der Statistik“ genannt und als Mitbegründer der wissenschaftlichen Statistik in die Lexika aufgenommen. Bei PLAYFAIR (1786) findet man erste Versuche einer übersichtlichen Darstellung mittels Tabellen und Schaubildern. Ende des 19. Jahrhunderts kam es zum Zusammenbruch der Staatenkunde, als sich die Wissenschaften immer mehr spezialisierten und verselbstständigten (Nationalökonomie, Geographie usw.).

Erste staatliche statistische Ämter findet man in Schweden (1756), Frankreich (1796) und Bayern (1801). Bis 1870 kam es allgemein zur Institutionalisierung statistischer Ämter in allen europäischen Staaten. In Deutschland wurde im Jahre 1872 das Kaiserliche Statistische Amt gegründet. Seit 1880 erfolgte die regelmäßige Herausgabe des „Statistischen Jahrbuchs für das Deutsche Reich“. Umfassende Volks-, Berufs- und Betriebszählungen fanden bereits in den Jahren 1882, 1895 und 1907 statt.

Politische Arithmetik

Sie entstand im 17. Jahrhundert unabhängig, aber auch als Gegenpol zur Universitätsstatistik und betonte die Erforschung der Bevölkerungsverhältnisse in Form einer mathematisch-statistischen Analyse, verbunden mit der Suche nach Regelmäßigkeiten im Wirtschafts- und Sozialleben. Als wichtige Neuerung ist hier z. B. die Aufstellung von Sterbetafeln durch GRAUNT (1662) und PETTY (1681) zu nennen. 1741 entwickelte der deutsche Statistiker und Nationalökonom JOHANN PETER SÜßMILCH statistische Modelle der allgemeinen Bevölkerungsentwicklung.

Wahrscheinlichkeitsrechnung

Parallel zu den beiden letztgenannten Entwicklungen kam es zu einer Übertragung von Methoden der Wahrscheinlichkeitsrechnung auf die Statistik. Statistik wird hier in erster Linie als Anwendungsgebiet der Stochastik (Zweig der angewandten Mathematik) betrachtet und damit die Grundlagen für die Entwicklung der analytischen Statistik gelegt. Dieser Zweig der Statistik hat vom 17. bis ins 20. Jahrhundert hinein die Geschichte der Statistik geprägt und in neuerer Zeit dominiert. Als wichtige Arbeiten in diesem Bereich gelten diejenigen von PASCAL (1623-1662), DE MOIVRE (1667-1754), J. BERNOULLI (1654-1705), D. BERNOULLI (1700-1782), BAYES (1702-1761), LAPLACE (1749-1827), GAUß (1777-1855), QUETELET (1796-1874), GALTON (1822-1911), PEARSON (1857-1936), FISHER (1890-1962), KOLMOGOROFF (1903-1987) und WALD (1903-1950).

Statistik in der zweiten Hälfte des 20. Jahrhunderts

Seit dem zweiten Weltkrieg hat sich mit dem Aufkommen des computergestützten Rechnens der Wissenschaftszweig Statistik außerordentlich schnell weiterentwickelt: Nichtparametrische Statistik [SIEGEL (1956), BICKEL und LEHMANN (1975/76)], Explorative Datenanalyse [TUKEY (1977)]; Computerabhängige Methoden, wie z. B. das Bootstrappen [EFRON (1979)]; Ausbau multivariater statistischer Verfahren, wie z. B. der LISREL-Ansatz [JÖRESKOG (1982)] – um nur einige herausragende Entwicklungen zu nennen. Im Zentrum der Diskussion steht die Polarisierung zwischen konfirmatorischem (induktivem) und explorativem (deskriptivem) statistischen Denken, d. h. zwischen der Nutzung von Wahrscheinlichkeitsmodellen zur statistischen Prüfung vorgegebener Hypothesen einerseits und der hypothesenbildenden statistischen Datenanalyse andererseits.

0.3 Arbeitsgebiete der Statistik

Die Arbeitsgebiete der Statistik gliedern sich in zwei Hauptbereiche: Einerseits handelt es sich um die Methodenlehre der Statistik, die statistische Verfahren und Techniken zur Analyse von Daten bereitstellt, ohne die Besonderheiten des fachwissenschaftlichen Ursprungs der zu analysierenden Daten zu beachten. Mit der Umsetzung der statistischen Methodik im Rahmen der jeweiligen Fachgebiete (insbesondere Wirtschafts- und Bevölkerungswissenschaften, Medizin, Soziologie, Physik, Psychologie, Rechtswissenschaft, Biologie, technische Wissenschaften usw.) beschäftigen sich andererseits die fachbezogenen Statistiken, wobei es neben der Anpassung der statistischen Methodik an die speziellen fachbezogenen Bedürfnisse vor allem auf die Operationalisierung fachlicher Begriffe für eine statistische Arbeit ankommt.

Statistische Methodenlehre

Unter diesem Sammelbegriff fasst man sowohl die Entwicklung und Bereitstellung statistischer Techniken (Verfahren, Methoden, Prozeduren) als auch die mathematischen Beweisführungen für ihre Richtigkeit zusammen. Ebenfalls zählen die Beschreibung typischer Anwendungsbereiche und die Erläuterung von Anwendungsgrenzen dazu. Seit etwa 1980 ist im Zusammenhang mit der Entwicklung statistischer Auswertungssysteme und geeigneter Software die Zusammenarbeit zwischen Informatikern und Statistikern sehr eng geworden. Man unterscheidet traditionell folgende Teilgebiete der statistischen Methodenlehre:

- *Beschreibende oder deskriptive Statistik*: Hierunter versteht man nicht nur das Erheben, Ordnen, Aufbereiten und Darstellen von Datenmengen in Tabellen und Schaubildern, sondern auch das Berechnen von charakteristischen Kenngrößen und deren Interpretation. Dabei beziehen sich Aussagen und Ergebnisse stets nur auf die untersuchten Daten. Seit den Arbeiten von J. W. TUKEY Ende der 70er Jahre ist zusammen mit der deskriptiven Statistik auch die so genannte explorative Datenanalyse (EDA) zu erwähnen.
- *Wahrscheinlichkeitsrechnung*: Sie ist ein Teilbereich der Mathematik, der als Bindeglied zwischen beschreibender und schließender Statistik dient. Die Untersuchung stochastischer Vorgänge mit Hilfe von Zufallsmodellen und darauf basierenden Zufallsvariablen steht im Mittelpunkt.
- *Schließende, induktive, inferentielle oder analytische Statistik (Stichprobenverfahren)*: Meist findet man die angeführten Bezeichnungen alternativ und gleichwertig nebeneinander. Allen ist gemeinsam, dass aufgrund der deskriptiven

Untersuchung eines Teils der Datenmenge (Stichprobe) von einer Teilmasse Rückschlüsse auf die zugrunde liegenden Gesetzmäßigkeiten in der Gesamtmasse gezogen werden sollen. Regeln aufzustellen für das stichprobenbasierte Schätzen unbekannter Parameter, Testen von Hypothesen und Fällen optimaler Entscheidungen sind die wichtigsten Aufgaben in diesem Bereich der Statistik.

Angewandte Statistiken

Hierunter versteht man die sogenannte *praktische Statistik* für spezielle Fachgebiete. Dies meint nicht nur die Anwendung der statistischen Methoden, sondern auch die Operationalisierung der verwendeten Begriffe. Zu diesem Zweck wird die praktische Statistik sehr tief gegliedert, z. B. die Wirtschaftsstatistik in *betriebliche und amtliche Statistik*, und dort weiter in institutionelle Ressorts (Agrarstatistik, Bankenstatistik, Industriestatistik, Verkehrsstatistik) und funktionelle Ressorts (Erwerbsstatistik, Preisstatistik, Außenhandelsstatistik, Einkommens- und Verbrauchsstatistik, Finanzstatistik, Statistik des produzierenden Gewerbes und volkswirtschaftliche Gesamtrechnung). Im Mittelpunkt steht die umfassende, kontinuierliche und aktuelle Information über wirtschaftliche, soziale und ökologische Zusammenhänge. Wegen der großen Bedeutung wird in den Industrieländern seit Beginn des 19. Jahrhunderts auch die *Bevölkerungsstatistik* in enger Verbindung mit der Wirtschaftsstatistik als Teil der amtlichen Statistik geführt.

0.4 Phasen des statistischen Arbeitens

Der Ablauf des statistischen Arbeitens umfasst sämtliche Phasen einer Studie, ausgehend von den ersten Vorüberlegungen über die eigentliche statistische Datenanalyse bis zur Interpretation der Ergebnisse. Im Einzelnen unterscheidet man folgende fünf Phasen:

- *Vorbereitung und Planung*: Zunächst muss eine Präzisierung der Ziele erfolgen, die mit der Untersuchung verfolgt werden. Hieraus ist eine sachliche, räumliche und zeitliche Abgrenzung des Problems und damit des Datenbedarfs abzuleiten. Insbesondere ist festzulegen, wie viele Objekte bzw. Fälle untersucht werden sollen. Dabei müssen auch die technisch-organisatorischen Abläufe der Untersuchung und ihre Kosten berücksichtigt werden.
- *Datenerfassung*: Hierunter versteht man die Gewinnung des statistischen Datenmaterials durch eine eigene Erhebung (*Primärstatistik*) oder durch den Rückgriff auf vorhandenes, für einen anderen Untersuchungszweck erhobenes

Datenmaterial (*Sekundärstatistik*). Primärstatistiken können genau auf das jeweilige Untersuchungsziel abgestellt werden, sind jedoch i. d. R. auch teurer als Sekundärstatistiken. Bei einer Primärstatistik unterscheidet man weiter zwischen *Voll- und Teilerhebung*. Vollerhebungen sind wegen der größeren Zahl von zu untersuchenden Einheiten kostspieliger und zeitraubender als Teilerhebungen. Demgegenüber sind Vollerhebungen nicht notwendigerweise genauer in den Ergebnissen als Teilerhebungen, was u. a. daran liegt, dass man angesichts der geringeren Anzahl von zu untersuchenden Einheiten in einer Teilerhebung bei der Bearbeitung der einzelnen Einheiten mehr Sorgfalt walten lassen kann. Die Techniken der Teilerhebungen sind vielfältig und werden im Rahmen der ‚Stichprobentheorie‘ ausführlich besprochen. Die *Erhebung* selbst erfolgt entweder durch Experiment, durch automatisches Erfassen, durch Beobachtung oder durch mündliche bzw. schriftliche Befragung.

● *Datenkontrolle und -aufbereitung*: Zunächst wird eine Kontrolle des Datenmaterials auf sachliche Richtigkeit (Plausibilitätskontrolle), Vollzähligkeit und Vollständigkeit durchgeführt. Das erklärte Ziel ist die Umwandlung des durch die Erhebung gewonnenen Urmaterials zu Aussagen über die zugrunde liegende Datenstruktur. Dazu werden die ungeordneten Daten verschlüsselt, sortiert und eventuell zu bestimmten Gruppen (Klassenbildung) zusammengefasst, um für die Anwendung statistischer Methoden vorbereitet zu sein. Gegebenenfalls sind auch Transformationen der Daten notwendig (z. B. Wechselkursumrechnungen, Veränderung der Skalenform).

● *Datenauswertung und -analyse*: Den Kern des statistischen Arbeitens bildet die Untersuchung des Datenmaterials mit Hilfe geeigneter statistischer Methoden. Hier kommen die Verfahren der deskriptiven und analytischen Statistik zum Einsatz. Die Arbeit in dieser Phase wird meist mit Hilfe von speziellen rechnergestützten Statistik-Programmpaketen durchgeführt. Gegenwärtig gibt es mehr als 200 kommerziell angebotene Programme, z. B. SPSS, SAS, Stata, STATISTICA, MINITAB.¹³

● *Datenpräsentation und Interpretation*: Die Darstellung der Ergebnisse in Tabellen und statistischen Schaubildern wird ebenfalls von rechnergestützten Programmen übernommen und dient als Ausgangspunkt für eine Interpretation der gefundenen Resultate. Vor allem bei der Interpretation sollte der Statistiker auf die Mithilfe von Experten der jeweiligen Fachdisziplin zurückgreifen.

¹³ Zu einer vergleichenden Übersicht vgl. G. BAMBERG, F. BAUR & M. KRAPP: *Statistik*, Kapitel 20.

// In allen fünf Phasen des statistischen Arbeitens können Fehler gemacht werden (von einer falschen Einschätzung der Problemlage bis zu fehlerhaften Interpretationen der Resultate).

Teil A

Beschreibende Statistik

Statistik beschäftigt sich, so die in Kapitel 0 gegebene Definition, mit der „zahlenmäßigen Untersuchung von Massenerscheinungen“. Eine solche Untersuchung bedingt offenbar, dass Daten in großen Mengen anfallen, die aufgrund ihres Umfangs einer unmittelbaren inhaltlichen Erfassung durch den Menschen nicht mehr zugänglich sind. Deshalb ist es unabdingbar, den Informationsgehalt der Daten in wenigen aussagekräftigen Zahlen und Grafiken zu verdichten: Maßzahlen für die Lage sollen die Größenordnung der erhobenen Werte wiedergeben, Streuungsmaße die Variationsbreite der Werte charakterisieren, Korrelationskoeffizienten den Zusammenhang der Beobachtungswerte für verschiedene Sachverhalte kennzeichnen.

Die sachgerechte Verdichtung der erhobenen Beobachtungswerte erfordert fundiertes Wissen über die Eigenschaften und Anwendungsvoraussetzungen der unterschiedlichen Methoden, welche die beschreibende Statistik zur Bewältigung dieser Aufgabe anbietet. Dieses Wissen stellt die Basis für eine angemessene Interpretation der statistischen Ergebnisse dar.

1 Grundbegriffe der Statistik

1.1 Statistische Massen

Statistische Untersuchungen haben häufig Massenerscheinungen zum Gegenstand, die jedoch sehr unpräzise beschrieben sein können, wie Wirtschaftswachstum, Arbeitslosigkeit oder Umweltverschmutzung. Wie misst oder zählt man z. B. ‚Umweltverschmutzung‘? Zunächst müssen geeignete statistische Massen definiert werden, die den zu untersuchenden Begriff den statistischen Methoden zugänglich machen. Bei der statistischen Untersuchung von ‚Umweltverschmutzung‘ müssen zähl- oder messbare statistische Massen gefunden werden: Gewässer verschiedener Wassergüteklassen, Schwefeldioxidmengen in der Luft, Investitionen für Lärmschutz, Primärenergieverbräuche der Einwohner. Welche Daten jeweils verwendet werden sollten, ist im Einzelfall zu klären.

Bevor man sich den Methoden zur Untersuchung von Daten zuwendet, ist zu klären, auf welche Weise Daten überhaupt entstehen, d. h. erfassbar und messbar werden. In diesem Zusammenhang sind einige Begriffe zu erläutern:

■ Statistische Einheit

Eine statistische Einheit ist ein Einzelobjekt einer statistischen Untersuchung, das Träger der Information(en) ist, für die man sich interessiert (Merkmalsträger).

■ Statistische Masse

Eine statistische Masse ist die Zusammenfassung aller in Verbindung mit dem Untersuchungsziel interessierenden statistischen Einheiten mit (mindestens) einer übereinstimmenden Eigenschaft, die durch *sachliche*, *räumliche* und *zeitliche Abgrenzung* exakt beschrieben und somit von den in diesem Zusammenhang nicht zu berücksichtigenden statistischen Einheiten deutlich zu unterscheiden sind.

Beispiel 1.1

Bei der statistischen Untersuchung eines Unternehmens kommen als statistische Massen beispielsweise der Personalbestand, der Materialbestand, der Auftragsbestand, aber auch die Gesamtheit der Zahlungsein- oder -ausgänge des Unternehmens in Betracht. Die einzelnen Beschäftigten, Materialien, Aufträge bzw. Zahlungsvorgänge bilden jeweils die einzelnen statistischen Einheiten.

Die Abgrenzung der statistischen Massen muss in zeitlicher, räumlicher und sachlicher Hinsicht erfolgen, sodass für jede statistische Einheit unzweifelhaft ist, ob sie zur fraglichen statistischen Masse gehört oder nicht. Beim Personalbestand ist in zeitlicher Hinsicht ein Zeitpunkt (d. h. ein Stichtag) festzulegen, zu dem die statistische Masse erhoben werden soll. Die räumliche Abgrenzung besteht in der Definition des Unternehmens selbst; für die sachliche Abgrenzung des Unternehmens ist der Begriff ‚Personal‘ zu definieren: Gehören zum Personal z. B. auch die Teilzeitbeschäftigten, die Unternehmenseigner und die in Erziehungsurlaub befindlichen Personen?

Nach der Art der zeitlichen Abgrenzung lassen sich die statistischen Massen in zwei Grundtypen untergliedern:

● *Bestandsmassen (stock)*: Haben die Einheiten einer statistischen Masse nebeneinander Bestand, so sind sie von Dauer und können zu einem bestimmten Zeitpunkt (Punktmasse) durch eine statistische Erhebung erfasst werden.

In Beispiel 1.1 ist dies beim Personalbestand, Materialbestand und Auftragsbestand der Fall.

● *Bewegungsmassen (flow)*: Ereignisse können dagegen nicht stichzeitmäßig erfasst werden, da sie im Allgemeinen nicht in genügender Anzahl gleichzeitig auftreten. Diese Ereignismassen müssen innerhalb eines längeren Zeitabschnittes durch laufende Registrierung der betreffenden Ereignisse erfasst werden.

In Beispiel 1.1 handelt es sich bei den Zahlungsein- und -ausgängen um Bewegungsmassen. Veränderungen der Bestandsmassen im Zeitablauf ergeben sich ebenfalls aufgrund von Bewegungsmassen, beim Personalbestand beispielsweise handelt es sich dabei um Kündigungen, Neueinstellungen usw.

Zu jeder Bestandsmasse lassen sich Bewegungsmassen angeben, aufgrund derer sich der Bestand der Bestandsmassen im Zeitablauf ändert. Für die Bestandsmasse kann ausgehend von ihrem Anfangsbestand über die zugehörigen Bewegungsmassen mittels *Fortschreibung* der Bestand zum Endzeitpunkt ermittelt werden, auch wenn für diesen Endzeitpunkt keine erneute Erhebung vorliegt:

$$(1.1) \quad n_t = n_{t-1} + z_{t-1,t} - a_{t-1,t}.$$

Dabei bezeichnen n_t den Umfang der Bestandsmasse zum Zeitpunkt t , $z_{t-1,t}$ den Umfang der Bewegungsmasse der Zugänge zwischen $t-1$ und t und $a_{t-1,t}$ den Umfang der entsprechenden Bewegungsmasse der Abgänge. Bei der Fortschreibung ist jedoch zu berücksichtigen, dass die Ungenauigkeit der Ergebnisse, die man mittels Fortschreibung erhält, mit dem zeitlichen Abstand zur letzten Erhebung der Bestandsmasse ansteigt, da sich Fehler bei der Erhebung der Bewegungsmassen mit der Fortschreibung der Bestandsmasse kumulieren.

1.2 Statistische Merkmale

Nach der Abgrenzung einer statistischen Masse müssen nun noch die an ihr zu erhebenden Untersuchungstatbestände genau umrissen werden:

■ Merkmal

Die Eigenschaft einer statistischen Einheit (Merkmalsträger), für die man sich im Rahmen der statistischen Untersuchung interessiert, heißt Merkmal. An einer statistischen Masse können mehrere Merkmale gleichzeitig untersucht werden.

■ Merkmalsausprägung

Unter einer Merkmalsausprägung versteht man eine der grundsätzlich möglichen Ausformungen eines Merkmals bei einem Merkmalsträger. Dies betrifft sowohl die verschiedenen Zahlen, die ein quantitatives Merkmal annehmen kann, als auch die unterschiedlichen Kategorien, die bei einem qualitativen Merkmal auftreten können.

■ Beobachtungswert

Die bei der statistischen Untersuchung an einer bestimmten statistischen Einheit (Merkmalsträger) einer statistischen Masse hinsichtlich eines bestimmten Merkmals festgestellte Merkmalsausprägung heißt Beobachtungswert. Beobachtungswerte sind der Ausgangspunkt für sämtliche Anwendungen statistischer Methoden.

Bei einer statistischen Erhebung wird an sämtlichen n statistischen Einheiten einer statistischen Masse ein Merkmal X beobachtet, d. h., an jeder Einheit wird die Ausprägung dieses Merkmals festgestellt. Sind $a_1, a_2, \dots, a_k$ die möglichen Ausprägungen des Merkmals X , so wird der i -ten statistischen Einheit ($i = 1, 2, \dots, n$) seine Ausprägung a_j als Beobachtungswert x_i zugeordnet:

$$(1.2) \quad x_i = a_j(i).$$

Die insgesamt n Beobachtungswerte $x_1, x_2, \dots, x_n$ heißen *Urliste*. Die Urliste umfasst i. d. R. so viele Daten in ungeordneter Form, dass die Übersichtlichkeit nicht mehr gewährleistet ist. Eine gewisse Abhilfe kann hier die geordnete Urliste bringen, bei der sämtliche in der Urliste vorkommenden Merkmalsausprägungen der Größe nach sortiert sind:

$$(1.3) \quad x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}.$$

// Selbstverständlich können innerhalb einer statistischen Masse mehrere Merkmale betrachtet werden. Bis einschließlich Kapitel 3 werden allerdings zunächst nur Methoden betrachtet, mit denen jeweils nur ein Merkmal statistisch analysiert werden kann. Die in Kapitel 4 vorgestellten bivariaten Methoden der Korrelations- und Regressionsrechnung beschäftigen sich mit der statistischen Analyse des Zusammenhangs zweier Merkmale. Statistische Methoden zur gemeinsamen Analyse von mehr als zwei Merkmalen (sog. *multivariate Statistik*) können im Rahmen dieser Einführung in die Statistik nicht behandelt werden.¹

// Man beachte zudem, dass die Abgrenzung der Begriffe Umfang einer statistischen Masse und Beobachtungswert auch vom Betrachtungsstandpunkt der Untersuchung abhängt. So kann z. B. die Mitarbeiterzahl eines Unternehmens den Umfang der statistischen Masse der Mitarbeiter dieses Unternehmens kennzeichnen, zugleich aber auch die Merkmalsausprägung dieses Unternehmens bei einer Erhebung in der statistischen Masse der Unternehmen sein.

Merkmale lassen sich in verschiedene Typen einteilen, u. a. nach den Unterscheidungstatbeständen qualitativ / quantitativ, häufbar / nicht häufbar sowie nach dem Skalenniveau (vgl. Abschnitt 1.3).

● *Qualitative Merkmale* besitzen die Eigenschaft, dass ihre Ausprägungen verbal und nicht durch Zahlen beschrieben sind. Werden diesen Merkmalsausprägungen Zahlen zugeordnet, so kann (ohne Veränderung des Informationsgehaltes) ein qualitatives Merkmal in ein quantitatives überführt werden (sog. *Quantifizierung*), was z. B. bei EDV-gestützten statistischen Auswertungen von Vorteil sein kann. Vielfach ist bei qualitativen Merkmalen mit sehr vielen möglichen Ausprägungen deren *Systematisierung* erforderlich, um sie statistisch auswerten zu können (vgl. Abschnitt 1.4).

● *Quantitative Merkmale* sind solche Merkmale, deren Ausprägungen durch reelle Zahlen dargestellt werden. Sie werden üblicherweise folgendermaßen weiter untergliedert in

- *diskrete* quantitative Merkmale, die als Merkmalsausprägungen nur abzählbar viele, i. d. R. ganzzahlige Werte besitzen, sowie in
- *stetige* quantitative Merkmale, die (theoretisch) zwischen zwei Merkmalsausprägungen beliebig viele Zwischenwerte annehmen können, deren

¹ Vgl. hierzu z. B. J. HARTUNG & B. ELPET: *Multivariate Statistik*; K. BACKHAUS: *Multivariate Analysemethoden*.

Anzahl nur durch die technischen Möglichkeiten der Messgenauigkeit oder durch das Erreichen einer zufriedenstellenden Messgenauigkeit praktisch begrenzt wird.

Darüber hinaus kann man bei Merkmalen noch zwischen *häufbaren* und *nicht häufbaren* unterscheiden, je nachdem, ob mehrere Merkmalsausprägungen desselben Merkmals bei einem Merkmalsträger vorkommen können oder nicht. Der Tatbestand der Häufbarkeit eines Merkmals muss bei der statistischen Auswertung besonders beachtet werden.

Beispiel 1.2

Von den in Beispiel 1.1 betrachteten statistischen Massen sei exemplarisch der Personalbestand betrachtet. Mögliche interessierende Merkmale sind das Geschlecht, das Einkommen, die tarifliche Gehaltsgruppe, das Alter, die Kinderzahl oder der ausgeübte Beruf der Beschäftigten. Qualitative Merkmale sind hierbei Geschlecht, Gehaltsgruppe und ausgeübter Beruf, quantitative Merkmale das Alter und das Einkommen.

Geschlecht, Alter, Gehaltsgruppe und Einkommen sind nichthäufbare Merkmale, während der ausgeübte Beruf – sofern zugelassen wird, dass eine Person mehreren Beschäftigungen nachgeht – ein häufbares Merkmal ist.

1.3 Skalierung von Merkmalen

Geht man davon aus, dass – ggf. nach Quantifizierung der qualitativen Merkmalsausprägungen – quantitative Merkmale vorliegen, so ist es für die Auswahl geeigneter statistischer Methoden wesentlich, welche inhaltlichen Bedeutungen durch diese Merkmalsausprägungen ausgedrückt werden. Messen heißt, dass den Eigenschaften von statistischen Einheiten nach bestimmten Regeln Zahlen zugeordnet werden. Die Messlatte, die diesem Messvorgang zugrunde liegt, nennt man *Skala*.

Man unterscheidet die folgenden Typen von Skalen, die auch als *Skalenniveaus* bezeichnet werden:

Nominalskala

Die Ausprägungen des untersuchten Merkmals werden durch die zugeordneten Zahlen lediglich nach dem Kriterium ‚gleich‘ oder ‚verschieden‘ unterschieden. Die zugeordneten Zahlen haben reine Identifizierungsfunktion und sind deshalb auswechselbar, z. B. Kfz-Kennzeichen, Hausnummern, Postleitzahlen. Dieser Skalentyp findet bei Merkmalen Anwendung, deren Ausprägungen keine natürliche Rangfolge bilden. Man unterscheidet alternative nominale Merkmale

(z. B. Geschlecht) und mehrklassige nominale Merkmale (z. B. Augenfarbe, Familienstand). Um mit nominalskalierten Daten statistisch zu arbeiten, können ihnen mittels Verschlüsselung Zahlen zugeordnet werden.

Ordinal- oder Rangskala

Die Ausprägungen des untersuchten Merkmals unterscheiden sich nicht nur, sondern können auch in eine Rangordnung gebracht werden. Vielfach wird diese Rangordnung durch Zahlen zum Ausdruck gebracht, die den Merkmalsausprägungen bei der numerischen Verschlüsselung zugewiesen werden. Die Ordinalskala ist für die Sozial- und Verhaltenswissenschaft von großer Bedeutung, weil dort viele Problemstellungen von einer Größer-kleiner- bzw. Vor-nach-Relation ausgehen (z. B. Intelligenz, sozialer Status, Aggressivität). Ordinalskalen geben jedoch keine Auskunft über den Abstand zwischen je zwei aufeinanderfolgenden Rangplätzen (z. B. Schulzensuren, Medaillenränge), d. h., über den Abstand verschiedener Ausprägungen des Merkmals wird nichts ausgesagt.

Metrische Skala

Eine metrische Skala lässt neben der Verschiedenartigkeit und der Rangfolge der Merkmalsausprägungen auch mess- und quantifizierbare Unterschiede erkennen. Man untergliedert dabei weiter in:

- *Intervall- oder Einheitsskala*: Die Ausprägungen des untersuchten Merkmals können nicht nur in eine Rangordnung gebracht werden, sondern die Abstände zwischen den Merkmalsausprägungen können miteinander verglichen werden. Die Intervallskala besitzt jedoch keinen natürlichen Nullpunkt und keine natürliche Einheit. Man kann den Nullpunkt willkürlich festsetzen und verändern, indem zu jedem Beobachtungswert eine konstante Zahl addiert oder subtrahiert wird (z. B. aus Gründen der rechnerischen Vereinfachung). Bezüglich der Ausprägungen der Merkmale besitzen zwar die Differenzen, aber nicht die Quotienten Aussagekraft. Ein typisches Beispiel ist die Temperaturskala: Auf der Celsius-Skala ist der Skalennullpunkt willkürlich am Gefrierpunkt des Wassers festgemacht, sodass die Aussage, 20 Grad Celsius sei doppelt so warm wie 10 Grad Celsius, nicht sinnvoll ist. Verschiebt man den Nullpunkt stattdessen z. B. auf denjenigen der Fahrenheitskala, d. h. um 32 °C, erhielte man demgegenüber ein Verhältnis von etwa 4:3. Ökonomische Beispiele für diesen Skalentyp sind allerdings eher selten.

- *Ratio- oder Verhältnisskala*: Die Ausprägungen des untersuchten Merkmals besitzen einen natürlichen Nullpunkt. Dadurch wird der Quotient zweier Merkmalsausprägungen unabhängig von der gewählten Maßeinheit. Für die Merkmals-

ausprägungen sind daher sowohl Differenzen als auch Quotienten sinnvoll berechenbar. Beispiele sind Flächen, Gewichte, Temperaturen in Kelvin (natürlicher Nullpunkt), Entfernungen (unveränderte Entfernungsverhältnisse beim Übergang z. B. von Kilometer auf Meilen).

● *Absolutskala*: Die Ausprägungen des untersuchten Merkmals besitzen einen natürlichen Nullpunkt und eine natürliche Einheit. Daher sind Aussagen über Summen, Differenzen, Produkte und Quotienten von Merkmalsausprägungen sinnvoll. Dies ist die Skalenform, die den größtmöglichen Informationsgehalt beinhaltet. Die Merkmalsausprägung basiert auf einem Zählvorgang und stellt daher meist Stückzahlen oder Häufigkeiten dar.

Beispiel 1.3

Die in Beispiel 1.2 betrachteten Merkmale zur Beschreibung des Personalbestands lassen sich üblicherweise auf folgenden Skalenniveaus messen:

Nominalskaliert sind das Geschlecht und der ausgeübte Beruf. Ordinalskaliert ist die Gehaltsgruppe, in die der Beschäftigte eingruppiert ist. Metrisch skaliert sind die Merkmale Alter, Einkommen und Kinderzahl, wobei die beiden ersten Merkmale verhältnisskaliert sind, während das letzte auf einer Absolutskala zu messen ist.

Bei der *Festlegung der Skalenart* kommt es darauf an, welche Aussage der Benutzer als empirisch gerechtfertigt und sinnvoll akzeptiert. Der Anwender statistischer Methoden muss sich daher fragen, was er mit der Zuordnung von Merkmalsausprägungen bezweckt und welche Rechenoperationen bzw. welche Transformationen benötigt werden.

// Die Einsatzmöglichkeiten der einzelnen statistischen Verfahren hängen ganz wesentlich vom Skalenniveau der Untersuchungseinheiten ab. Die folgenden Beispiele mögen dies verdeutlichen:

- Nominalskala: Modus, Kontingenzkoeffizient;
- Ordinalskala: Median, Quartile, Rangkorrelationskoeffizient;
- Intervallskala: arithmetisches Mittel, Standardabweichung, Produktmomentkorrelationskoeffizient;
- Ratioskala: geometrisches Mittel, Variationskoeffizient.

1.4 Skalentransformationen und Klassenbildung

Bei der Erfassung und Aufbereitung statistischer Daten sind i. d. R. *Skalentransformationen* erforderlich: Dabei werden die Werte einer Skala in Werte einer anderen Skala übertragen. Beispielweise geschieht dies bei der Verschlüsselung

qualitativer Daten zur EDV-mäßigen Auswertung, indem die Werte ‚männlich‘ und ‚weiblich‘ einer Nominalskala zur Messung der Ausprägungen des Merkmals ‚Geschlecht‘ in die Werte ‚0‘ und ‚1‘ transformiert werden. Wichtigste Bedingung ist dabei, dass die Messeigenschaften der Skala erhalten bleiben, da andernfalls eine Veränderung des Informationsgehalts der Daten eintritt. Obwohl die Werte ‚0‘ und ‚1‘ für sich genommen auch Ausprägungen einer Absolutskala sein könnten, würde jede andere als eine nominalskalierte Interpretation dieser Werte diesen einen Informationsgehalt zuerkennen, der in den ursprünglich erhobenen Daten gar nicht vorhanden war.

Welche Skalentransformationen möglich sind, ohne dass Informationsverluste hinzunehmen sind, hängt vom Skalenniveau der Ursprungsskala ab. Hierbei gilt Folgendes:

- Nominalskala: Skalentransformationen müssen umkehrbar eindeutig sein.
- Ordinalskala: Skalentransformationen müssen umkehrbar eindeutig und streng monoton sein.
- Metrische Skala: Hier sind nur lineare Transformationen zulässig, und zwar in Abhängigkeit von der Art der metrischen Skala:
 - Intervallskala: $b_j = \alpha \cdot a_j + \beta$ mit $\alpha > 0$ und β beliebig,
 - Ratioskala: $b_j = \alpha \cdot a_j + \beta$ mit $\alpha > 0$ und $\beta = 0$,
 - Absolutskala: $b_j = \alpha \cdot a_j + \beta$ mit $\alpha = 1$ und $\beta = 0$.

Von der Nominal- bis zur Absolutskala nimmt der in den Merkmalsausprägungen steckende Informationsgehalt zu. Ist man bereit, auf den mit einer Skaleneigenschaft verbundenen Teil des Informationsgehalts zu verzichten, so kann man Daten eines höheren Skalenniveaus auch mit Methoden auswerten, die eigentlich für ein niedrigeres Skalenniveau geeignet sind. Beispielsweise kann der Median (vgl. Abschnitt 3.2) auch für metrisch skalierte Daten benutzt werden. Die Gründe für die Inkaufnahme eines solchen Informationsverlusts können z. B. die Einfachheit der Methoden niedrigeren Skalenniveaus oder die Notwendigkeit sein, bei mehreren Merkmalen ein gemeinsames Skalenniveau zu verwenden.

Eine besondere Art der Skalentransformation stellt die *Klassenbildung* für die Merkmalsausprägungen dar (sog. *Gruppierung der Daten*). Dabei werden jeweils mehrere Ausprägungen des Ursprungsmerkmals zusammengefasst und in eine Ausprägung (eine sogenannte *Klasse*) des transformierten Merkmals überführt. Diese Transformation, bei der i. d. R. zwar das Skalenniveau beibehalten wird, jedoch die Messgenauigkeit auf der verwendeten Skala reduziert wird, ist natürlich ebenfalls mit einem Informationsverlust verbunden.

Die Begründung für dieses Vorgehen ist darin zu sehen, dass – vor allem bei stetigen Merkmalen – die in der Urliste enthaltenen Beobachtungswerte oft alle verschieden sind. Die Ermittlung von absoluten oder relativen Häufigkeiten wäre nicht sinnvoll, da die meisten Merkmalsausprägungen keinmal, einmal oder höchstens zweimal vorkommen. Echte Häufungen der Merkmalsausprägungen lassen sich dann kaum feststellen; eine grafische Darstellung der Verteilung des Merkmals in der statistischen Masse ist praktisch unmöglich.

Bei der Klassenbildung zerlegt man das Gesamtintervall, in dem alle möglichen Merkmalsausprägungen liegen, in Teilintervalle I_k , die (im Fall metrischer Skalierung) eindeutig bestimmt sind durch ihre Klassenmitten a_k^* und Klassenbreiten d_k bzw. durch ihre unteren Klassengrenzen u_k und oberen Klassengrenzen o_k :

$$(1.4) \quad \begin{aligned} u_k &= a_k^* - \frac{1}{2} d_k, & o_k &= a_k^* + \frac{1}{2} d_k; \\ a_k^* &= \frac{1}{2}(u_k + o_k), & d_k &= o_k - u_k. \end{aligned}$$

Wenn nach der Klassenbildung nur noch die klassierten Daten vorliegen, besitzt man keine Informationen mehr über die Verteilung der Beobachtungswerte innerhalb der Klassen. Man unterstellt dann als naheliegende Annahme eine der beiden Arbeitshypothesen, dass alle Beobachtungswerte, die in eine Klasse fallen, sich entweder gleichmäßig über diese Klasse verteilen oder sich auf den Punkt der Klassenmitte konzentrieren.

Bei der Festlegung der Klassengrenzen sollte die Gesamtzahl K der zu bildenden Klassen wesentlich kleiner als n sein, andernfalls würde sich eine Klassenbildung erübrigen. Zudem sollten die Klassen (vor allem im Hinblick auf grafische Darstellungsmöglichkeiten) möglichst gleich breit gewählt werden.

// Konstante Klassenbreiten führen jedoch immer dann zu Verzerrungen, wenn hierdurch homogen besetzte Intervalle zerschnitten und heterogene Bereiche zusammengefasst werden. Auch wenn in bestimmten Bereichen sehr dünne Klassenbelegungen auftreten, sollten ungleiche Klassenbreiten, d. h. breitere Klassen bei Bereichen mit geringerer Anzahl von Merkmalsausprägungen, verwendet werden. So empfiehlt es sich beispielsweise bei der Klassenbildung für das Jahreseinkommen von Steuerpflichtigen im oberen Einkommensbereich breitere Klassen zu wählen.

Die Klassen werden links geschlossen und rechts offen gewählt, d. h., sie sind von der Form:

$$(1.5) \quad I_k = [a_k^* - \frac{1}{2} d_k ; a_k^* + \frac{1}{2} d_k) = [u_k ; o_k).$$