



ulm university universität
uulm

Stochastik I

Vorlesungsskript

Prof. Dr. Evgeny Spodarev

ULM
2016

Vorwort

Dieses Skript entstand aus dem Zyklus der Vorlesungen über Statistik, die ich in den Jahren 2005–2012 an der Universität Ulm gehalten habe. Dabei handelt es sich um die erste Einführung in die Statistik, die durch die aufbauende Vorlesung Stochastik III ergänzt wird.

Dieses Skript gibt eine Übersicht über die typischen Fragestellungen und Methoden der mathematischen Statistik. Es stellt einen Versuch dar, einen Mittelweg zwischen praktisch orientierten (aber mathematisch oft sehr dürftigen) Statistik-Monographien einerseits und trockenen Büchern über die mathematische Statistik andererseits einzuschlagen. Ob es mir gelungen ist, soll der Leser beurteilen.

Ich möchte gerne meinen Kollegen aus dem Institut für Stochastik, Herrn Prof. Volker Schmidt und Herrn Dipl.-Math. Malte Spiess, für ihre Unterstützung und anregenden Diskussionen während der Entstehung des Skriptes danken. Herr Tobias Brosch hat eine hervorragende Arbeit beim Tippen des Skriptes und bei der Erstellung zahlreicher Abbildungen, die den Text begleiten, geleistet. Dafür gilt ihm mein herzlicher Dank.

Ulm, den 19.04.2012

Evgeny Spodarev

Inhaltsverzeichnis

1	Einführung	1
1.1	Typische Fragestellungen, Aufgaben und Ziele der Statistik	1
1.2	Statistische Merkmale und ihre Typen	4
1.3	Statistische Daten und Stichproben	5
1.4	Stichprobenfunktionen	6
2	Beschreibende Statistik	7
2.1	Verteilungen und ihre Darstellungen	7
2.1.1	Häufigkeiten und Diagramme	7
2.1.2	Empirische Verteilungsfunktion	10
2.2	Beschreibung von Verteilungen	12
2.2.1	Lagemaße	12
2.2.2	Streuungsmaße	16
2.2.3	Konzentrationsmaße	17
2.2.4	Maße für Schiefe und Wölbung	20
2.3	Quantilplots (Quantil-Grafiken)	22
2.4	Dichteschätzung	25
2.5	Beschreibung und Exploration von bivariaten Datensätzen	27
2.5.1	Grafische Darstellung von bivariaten Datensätzen	27
2.5.2	Zusammenhangsmaße	29
2.5.3	Einfache lineare Regression	33
3	Punktschätzer	41
3.1	Parametrisches Modell	41
3.2	Parametrische Familien von statistischen Prüfverteilungen	42
3.2.1	Gamma-Verteilung	42
3.2.2	Student-Verteilung (t-Verteilung)	45
3.2.3	Fisher-Snedecor-Verteilung (F-Verteilung)	48
3.3	Punktschätzer und ihre Grundeigenschaften	50
3.3.1	Eigenschaften von Punktschätzern	50
3.3.2	Schätzer des Erwartungswertes und empirische Momente	52
3.3.3	Schätzer der Varianz	54
3.3.4	Eigenschaften der Ordnungsstatistiken	62
3.3.5	Empirische Verteilungsfunktion	64
3.4	Methoden zur Gewinnung von Punktschätzern	72
3.4.1	Momentenschätzer	72
3.4.2	Maximum-Likelihood-Schätzer	74
3.4.3	Bayes-Schätzer	84
3.4.4	Resampling-Methoden zur Gewinnung von Punktschätzern	86

3.5	Weitere Güteeigenschaften von Punktschätzern	90
3.5.1	Ungleichung von Cramér-Rao	90
3.5.2	Bedingte Erwartung	95
3.5.3	Suffizienz	97
3.5.4	Vollständigkeit	102
3.5.5	Bester erwartungstreuer Schätzer	103
4	Konfidenzintervalle	106
4.1	Einführung	106
4.2	Ein-Stichproben-Probleme	108
4.2.1	Normalverteilung	108
4.2.2	Konfidenzintervalle aus stochastischen Ungleichungen	110
4.2.3	Asymptotische Konfidenzintervalle	111
4.3	Zwei-Stichproben-Probleme	113
4.3.1	Normalverteilte Stichproben	113
4.3.2	Poissonverteilte Stichproben	115
5	Tests statistischer Hypothesen	118
5.1	Allgemeine Philosophie des Testens	118
5.2	Nichtrandomisierte Tests	126
5.2.1	Parametrische Signifikanztests	126
5.3	Randomisierte Tests	131
5.3.1	Grundlagen	131
5.3.2	Neyman-Pearson-Tests bei einfachen Hypothesen	132
5.3.3	Einseitige Neyman-Pearson-Tests	137
5.3.4	Unverfälschte zweiseitige Tests	142
5.4	Anpassungstests	148
5.4.1	χ^2 -Anpassungstest	148
5.4.2	χ^2 -Anpassungstest von Pearson-Fisher	153
5.4.3	Anpassungstest von Shapiro	159
5.5	Weitere, nicht parametrische Tests	161
5.5.1	Binomialtest	161
5.5.2	Iterationstests auf Zufälligkeit	162
	Literatur	166
	Index	168

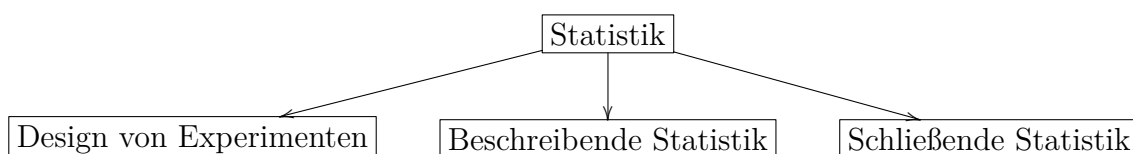
1 Einführung

1.1 Typische Fragestellungen, Aufgaben und Ziele der Statistik

Im alltäglichen Sprachgebrauch versteht man unter „Statistik“ eine Darstellung von Ergebnissen des Zusammenzählens von Daten und Fakten jeglicher Art, wie z.B. ökonomischen Kenngrößen, politischen Umfragen, Daten der Marktforschung, klinischen Studien in der Biologie und Medizin, usw.

Die *mathematische Statistik* jedoch kann viel mehr. Sie arbeitet mit *Daten-Stichproben*, die nach einem bestimmten Zufallsmechanismus aus der *Grundgesamtheit* aller Daten, die in Folge von Beobachtung, Experimenten (reale Daten) oder Computersimulation (synthetische Daten) erhoben wurden. Dabei beschäftigt sich die mathematische Statistik mit folgenden Fragestellungen:

1. Wie sollen die Daten gewonnen werden? (Design von Experimenten)
2. Wie sollen (insbesondere riesengroße) Datensätze beschrieben werden, um die Gesetzmäßigkeiten und Strukturen in ihnen entdecken zu können? (Beschreibende (deskriptive) und explorative Statistik)
3. Welche Schlüsse kann man aus den Daten ziehen? (Schließende oder induktive Statistik)



In dieser einführenden Vorlesung werden wir Teile der beschreibenden und schließenden Statistik kennenlernen, wobei die Datenerhebung aus Platzgründen ausgelassen wird. Die *Arbeitsweise eines Statistikers* sieht folgendermaßen aus:

1. *Datenerhebung*
2. *Visualisierung und beschreibende Datenanalyse*
3. *Datenbereinigung* (z.B. Erkennung fehlerhafter Messungen, Ausreißern, usw.)
4. *Explorative Datenanalyse* (Suche nach Gesetzmäßigkeiten)
5. *Modellierung der Daten* mit Methoden der Stochastik
6. *Modellanpassung* (Schätzung der Modellparameter)
7. *Modellvalidierung* (wie gut war die Modellanpassung?)

Pflanze	1	2	3	4	5	6	7	8	9	10
rund	45	27	24	19	32	26	88	22	28	25
kantig	12	8	7	10	11	6	24	10	6	7
Verhältnis ... : 1	3,8	3,4	3,4	1,9	2,9	4,3	3,7	2,2	4,7	3,6

Tab. 1.1: Ergebnisse für die 10 Pflanzen des ersten Versuchs von Mendel

8. Schließende Datenanalyse:

- Konstruktion von *Vertrauensintervallen* (Konfidenzintervallen) für Modellparameter und deren Funktionen,
- Tests statistischer Hypothesen,
- Vorhersage von Zielgrößen (z.B. auf Basis modellbezogener Computersimulation).

Uns werden in diesem Vorlesungsskript vor allem die Arbeitspunkte 2), 4)–6) und 8) beschäftigen.

Beispiel 1.1.1

Nachfolgend geben wir einige typische Fragestellungen der Statistik an Beispielen von Datensätzen:

1. Statistische Herleitung von Grundsätzen der biologischen Evolution (Mendel, 1865):

Es wurden Nachkommen von zwei Erbsensorten, die sich in der Samenform unterscheiden, gezüchtet: die erste Sorte hat runde, die zweite kantige Erbsen. Johann Gregor Mendel hat festgestellt, dass sich runde Samen dominant vererben. Dabei werden bei einer Bestäubung von Pflanzen der einen Sorte mit Pollen der anderen alle Nachkommen runde Samen zeigen, die genetisch heterozygot sind, d.h., beide Allele aufweisen. Kreuzt man diese hybriden Pflanzen, so zeigen sie runde und kantige Samen im Verhältnis 3 : 1 (Spaltungs- und Dominanzregeln von Mendel). Bei der statistischen Überprüfung seiner Vermutungen erhielt Mendel 5475 runde und 1850 kantige Samen, die somit im Verhältnis 2,96 : 1 stehen. In der Tabelle 1.1 sind Ergebnisse für die ersten 10 Pflanzen gezeigt. Man sieht, dass das oben genannte Verhältnis zufällig um 3 : 1 schwankt. Durch die Bildung des Mittels über das Gesamtkollektiv der Daten wird die Gesetzmäßigkeit 3 : 1 gefunden (explorative Statistik).

2. Kreditwürdigkeit bei Kreditvergabe

Die Banken sind offensichtlich daran interessiert, Bankkredite an Kunden zu vergeben, die in der Zukunft solvent bleiben, also die Kreditraten regelmäßig zurückzahlen können. Um die Kreditwürdigkeit zu überprüfen, werden Umfragen gemacht, wobei die Antworten unter anderem in folgenden Variablen kodiert werden:

- X_1 Laufendes Konto bei der Bank (1 = nein, 2 = ja und durchschnittlich geführt, 3 = ja und gut geführt)
- X_2 Laufzeit des Kredits in Monaten
- X_3 Kredithöhe in €
- X_4 Rückzahlung früherer Kredite (gut/ schlecht)

X_1 : laufendes Konto	Y	
	1	0
nein	45,0	19,9
gut	15,3	49,7
mittel	39,7	30,4
X_3 : Kredithöhe in €	1	0
$0 < \dots \leq 500$	1,00	2,14
$500 < \dots \leq 1000$	11,33	9,14
$1000 < \dots \leq 1500$	17,00	19,86
$1500 < \dots \leq 2500$	19,67	24,57
$2500 < \dots \leq 5000$	25,00	28,57
$5000 < \dots \leq 7500$	11,33	9,71
$7500 < \dots \leq 10000$	6,67	3,71
$10000 < \dots \leq 15000$	7,00	2,00
$15000 < \dots \leq 20000$	1,00	0,29
X_4 : Frühere Kredite	1	0
gut	82,33	94,85
schlecht	17,66	5,15
X_5 : Verwendungszweck	1	0
privat	57,53	69,29
beruflich	42,47	30,71

Tab. 1.2: Lernstichprobe zur Vergabe von Krediten

- X_5 Verwendungszweck (privat / geschäftlich)
- X_6 Geschlecht (weiblich / männlich)

Um an Hand eines ausgefüllten Fragebogens wie diesem eine Entscheidung über die Vergabe des Kredits treffen zu können, werden *Lernstichproben* herangezogen, bei denen das Ergebnis Y der erfolgten Kreditvergabe bekannt ist. Dabei bedeutet $Y = 0$ gut und $Y = 1$ schlecht. Betrachten wir eine solche Stichprobe einer süddeutschen Bank, die 1000 Umfragebögen umfasst. Dabei sind 700 kreditwürdig und 300 davon nicht kreditwürdig gewesen. Die Tabelle 1.2 zeigt Prozentzahlen dieses Datensatzes für ausgewählte Merkmale X_i . Dabei ist es möglich, mit Hilfe statistischer Methoden (Regression) eine Kreditentscheidung bei einem Kunden an Hand dieser Lernprobe automatisch treffen zu können. Dieser Vorgang wird manchmal auch „statistisches Lernen“ genannt. Fragestellungen wie diese werden erst in Stochastik III (verallgemeinerte lineare Modelle) behandelt.

3. Korrosion von Legierungen

In diesem Beispiel wurde der Korrosionsgrad einer Kupfer-Nickel-Legierung in Abhängigkeit ihres Eisengehalts untersucht. Dazu wurden 13 verschiedene Räder mit dieser Legierung beschichtet und 60 Tage lang in Meerwasser gedreht. Danach wurde der Gewichtsverlust in mg pro dm^2 und Tag bestimmt. Aus dem Bild 1.1 ist zu sehen, dass die Korrosion in Abhängigkeit vom Eisengehalt linear abnimmt. Mit statistischen Methoden (einfache lineare Regression) kann die Geschwindigkeit dieser Abnahme geschätzt werden.

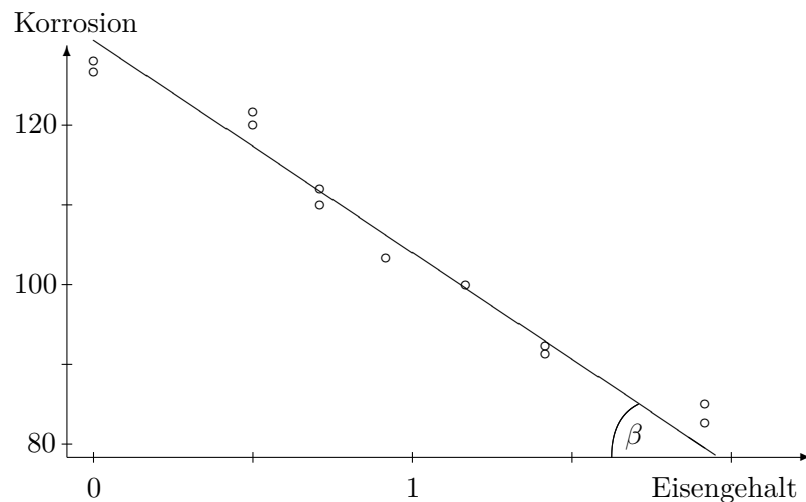
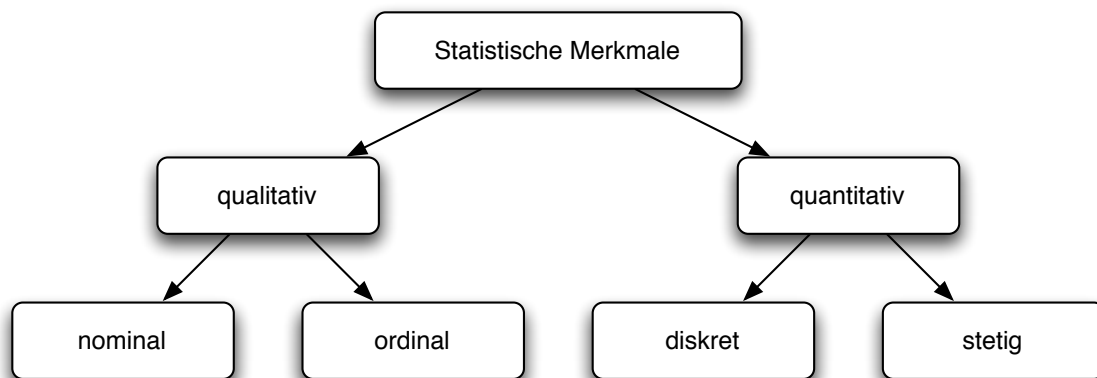


Abb. 1.1: Korrosion von Kupfer-Nickel-Legierung

1.2 Statistische Merkmale und ihre Typen

Die Daten, die zur statistischen Analyse vorliegen, können eine oder mehrere interessierende Größen (die auch *Variablen* oder *Merkmale* genannt werden) umfassen. Ihre Werte werden *Merkmalsausprägungen* genannt. In dem nachfolgenden Diagramm werden mögliche Typen der statistischen Merkmale gegeben.



Diese Typen entstehen in Folge der Klassifikation von Wertebereichen (Skalen) der Merkmale. Dennoch ist diese Einteilung nicht vollständig und kann bei Bedarf erweitert werden. Man unterscheidet *qualitative* und *quantitative* Merkmale. *Quantitative Merkmale* lassen sich inhaltlich gut durch Zahlen darstellen (z.B. Kredithöhe in €, Körpergewicht und Körpergröße, Blutdruck usw.). Sie können *diskrete* oder *stetige* Wertebereiche haben, wobei diskrete Merkmale isolierte Werte annehmen können (z.B. Anzahl der Schäden eines Versicherers pro Jahr). Stetige Wertebereiche hingegen sind überabzählbar. Dennoch liegen in der Praxis stetige Merkmale in gerundeter Form vor (z.B. Körpergröße auf cm gerundet, Geldbeträge auf € gerundet usw.).

Im Gegensatz zu den quantitativen Merkmalen sind die Inhalte der *qualitativen Merkmale*, wie z.B. Blutgruppe (0, A, B und AB) oder Familienstand (ledig, verheiratet, verwitwet),

nicht sinnvoll durch Zahlen darzustellen. Sie können zwar formell mit Zahlen kodiert werden (z.B. bei Blutgruppen $0 = 0$, $A = 1$, $B = 2$, $AB = 3$), aber solche Kodierungen stellen keinen inhaltlichen Zusammenhang zwischen Ausprägungen und Zahlen-Codes dar sondern dienen lediglich der besseren Identifikation der Merkmale auf einem Rechner. Es ist insbesondere unsinnig, Mittelwerte und ähnliches von solchen Codes zu bilden.

Ein qualitatives Merkmal mit nur 2 Ausprägungen (z.B. männlich / weiblich, Raucher / Nichtraucher) heißt *alternativ*. Ein qualitatives Merkmal kann *ordinal* (wenn sich eine natürliche lineare Ordnung in den Merkmalsausprägungen finden lässt, wie z.B. gut / mittel / schlecht bei Qualitätsbewertung in Umfragen oder sehr gut / gut / befriedigend / ausreichend / mangelhaft / ungenügend bei Schulnoten) oder *nominal* (wenn eine solche Ordnung nicht vorhanden ist) sein. Beispiele von nominalen Merkmalen sind Fahrzeugmarken in der KFZ-Versicherung (z.B. BMW, Peugeot, Volvo, usw.) oder Führerscheinklassen ($A, B, C, \dots$). Datenmerkmale können auch mehrdimensionale Ausprägungen haben. In dieser Vorlesung behandeln wir jedoch hauptsächlich eindimensionale Merkmale.

1.3 Statistische Daten und Stichproben

Aus den obigen Beispielen wird klar, dass ein Statistiker mit Datensätzen der Form $(x_1, \dots, x_n)$ arbeitet, wobei die Einzeleinträge x_i aus einer Grundgesamtheit $G \subset \mathbb{R}^k$ stammen, die hypothetisch unendlich groß ist. Der vorliegende Datensatz $(x_1, \dots, x_n)$ wird auch (*konkrete*) *Stichprobe* von Umfang n genannt. Die Menge B aller potentiell möglichen Stichproben bezeichnen wir als *Stichprobenraum* und setzen zur Vereinfachung der Notation $B = \mathbb{R}^{kn}$. In diesem Skript werden wir meistens die univariate statistische Analyse (also $k = 1$, ein eindimensionales Merkmal) betreiben. In der beschreibenden Statistik arbeitet man mit Stichproben $(x_1, \dots, x_n)$ und ihren Funktionen, um diese Daten visualisieren zu können. Für die Aufgabe der schließenden Statistik jedoch reicht diese Datenebene nicht mehr aus. Daher wird die zweite Ebene der Betrachtung eingeführt, die sogenannte *Modellebene*. Dabei wird angenommen, dass die konkrete Stichprobe $(x_1, \dots, x_n)$ eine *Realisierung* eines stochastischen Modells $(X_1, \dots, X_n)$ darstellt, wobei $X_1, \dots, X_n$ (meistens unabhängige identisch verteilte) Zufallsvariablen auf einem (nicht näher spezifizierten) Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P})$ sind. Diese Zufallsvariablen X_i , $i = 1, \dots, n$ können als konsequente Beobachtungen eines Merkmals interpretiert werden. In Bsp. 1.1.1, 1) z.B. die Erbsenform mit

$$X_i = \begin{cases} 0, & \text{falls Erbse } i \text{ rund,} \\ 1, & \text{falls Erbse } i \text{ eckig,} \end{cases} \quad i = 1, \dots, n.$$

Der Vektor $(X_1, \dots, X_n)$ wird dabei *Zufallsstichprobe* genannt. Man setzt weiter voraus, dass $EX_i^2 < \infty \forall i = 1, \dots, n$, damit man von der Varianz $\text{Var } X_i$ der Einzeleinträge sprechen kann. Es wird außerdem angenommen, dass ein $\omega \in \Omega$ existiert, sodass $X_i(\omega) = x_i \quad \forall i = 1, \dots, n$. Sei F die Verteilungsfunktion der Zufallsvariablen X_i . Eine der wichtigsten Aufgaben der Statistik ist die Bestimmung von F (man sagt, „Schätzung von F “) aus den konkreten Daten $(x_1, \dots, x_n)$. Dabei können auch Momente von F und ihre Funktionen (Erwartungswert, Varianz, Schiefe, usw.) von Interesse sein.

1.4 Stichprobenfunktionen

Um die obigen Aufgaben erfüllen zu können, braucht man gewisse Funktionen $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$, $m \in \mathbb{N}$ auf dem Stichprobenraum, die diese Stichprobe bewerten.

Definition 1.4.1

Eine Borel-messbare Abbildung $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^m$ heißt *Stichprobenfunktion*. Wenn man auf der Modellebene mit einer Zufallsstichprobe $(X_1, \dots, X_n)$ arbeitet, so heißt die Zufallsvariable

$$\varphi(X_1, \dots, X_n)$$

eine *Statistik*. In der Schätztheorie spricht man dabei von *Schätzern* und bei statistischen Tests wird $\varphi(X_1, \dots, X_n)$ *Teststatistik* genannt.

Beispiele für Stichprobenfunktionen sind unter anderen das *Stichprobenmittel*

$$\bar{x}_n = \frac{1}{n} \sum_{i=1}^n x_i,$$

die *Stichprobenvarianz*

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

und die *Ordnungsstatistiken*

$$x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)},$$

die entstehen, wenn man eine Stichprobe, die aus quantitativen Merkmalen besteht, linear ordnet ($x_{(1)} = \min_{i=1, \dots, n} x_i, \dots, x_{(n)} = \max_{i=1 \dots n} x_i$). Weitere Beispiele und ihre Charakteristiken werden in Kapitel 2 gegeben.

2 Beschreibende Statistik

Sei eine konkrete Stichprobe $(x_1, \dots, x_n)$, $x_i \in \mathbb{R}$ gegeben, wobei die x_i als Realisierungen der Zufallsvariablen $X_i \stackrel{d}{=} X$ mit Verteilungsfunktion F interpretiert werden können.

2.1 Verteilungen und ihre Darstellungen

In diesem Abschnitt werden wir Methoden zur statistischen Beschreibung und grafischen Darstellung der (unbekannten) Verteilung F betrachten.

2.1.1 Häufigkeiten und Diagramme

Falls das quantitative Merkmal X eine endliche Anzahl von Ausprägungen $\{a_1, \dots, a_k\}$, $a_1 < a_2 < \dots < a_k$, besitzt, also

$$\mathbb{P}(X \in \{a_1, \dots, a_k\}) = 1,$$

dann kann eine Schätzung der Zähldichte $p_i = P(X = a_i)$ von X aus den Daten $(x_1, \dots, x_n)$ grafisch dargestellt werden. Ähnliche Darstellungen sind für die Dichte $f(x)$ von absolut stetigen Merkmalen X möglich, wobei ihr Wertebereich C sich in k Klassen aufteilen lässt: $(c_{i-1}, c_i]$, $i = 1, \dots, k$, wobei $c_0 = -\infty$, $c_1 < \dots < c_{k-1}$, $c_k = \infty$ ist. Dann kann die Zähldichte $p_i = \mathbb{P}(X \in (c_{i-1}, c_k])$ gegeben durch

$$p_i = \int_{c_{i-1}}^{c_i} f(x) dx, \quad i = 0, \dots, k$$

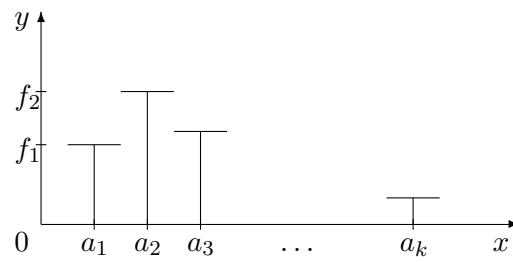
betrachtet werden.

Definition 2.1.1

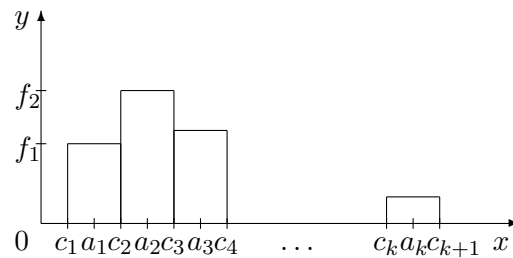
1. Die *absolute Häufigkeit* von Merkmalsausprägung a_i bzw. Klasse $(c_{i-1}, c_i]$, $i = 1, \dots, k$ ist $n_i = \#\{x_j, j = 1, \dots, n : x_j = a_i\}$ bzw. $n_i = \#\{x_j, j = 1, \dots, n : x_j \in (c_{i-1}, c_i]\}$.
2. Die *relative Häufigkeit* von Merkmalsausprägung a_i bzw. Klasse $(c_{i-1}, c_i]$ ist $f_i = n_i/n$, $i = 1, \dots, k$.

Es gilt offensichtlich $n = \sum_{i=1}^k n_i$, $0 \leq f_i \leq 1$, $\sum_{i=1}^k f_i = 1$. Die absoluten und relativen Häufigkeiten werden oft in Häufigkeitstabellen zusammengefasst. Zu ihrer Visualisierung dienen so genannte *Diagramme*. Es wird grundsätzlich zwischen *Histogrammen* und *Kreisdiagrammen* unterschieden.

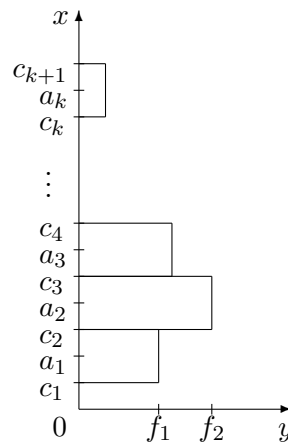
1. *Histogramme* werden gebildet, indem man die Paare (a_i, f_i) (bzw. $(1/2(c_1 + x_{(1)}), f_1)$, $(1/2(c_{i-1} + c_i), f_i)$, $i = 2, \dots, k-1$, $(1/2(c_{k-1} + x_{(n)}), f_k)$ im absolut stetigen Fall, wobei hier die Bezeichnung $a_i = 1/2(c_{i-1} + c_i)$ verwendet wird und $x_{(1)} < c_1$, $x_{(n)} > c_{k-1}$ angenommen wird.) auf der Koordinatenebene (x, y) folgendermaßen aufträgt:
 - *Stabdiagramm*: f_i wird als Höhe des senkrechten Strichs über a_i dargestellt:



- *Säulendiagramm*: genauso wie ein Stabdiagramm, nur werden Striche durch Säulen der Form $(c_{i-1}, c_i] \times f_i$ ersetzt, wobei im diskreten Fall die Aufteilung der reellen Achse $-\infty = c_0 < c_1 < c_2 < \dots < c_{k-1} < c_k = \infty$ in Intervalle beliebig vorgenommen werden kann.

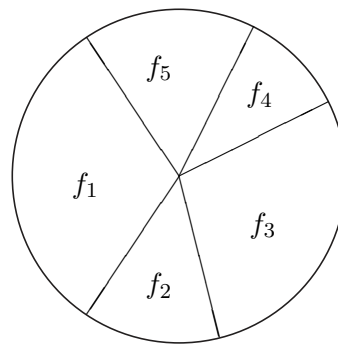


- *Balkendiagramm*: genauso wie Säulendiagramm, nur mit vertikalen statt horizontaler x -Achse.



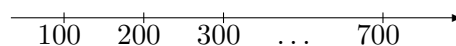
2. Kreisdiagramme (Tortendiagramme):

Ein Kreis wird in Segmente mit Öffnungswinkel α_i eingeteilt, die proportional zu f_i sind:
 $\alpha_i = 2\pi f_i$, $i = 1, \dots, n$.

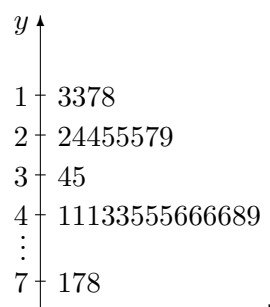


3. Stamm-Blatt-Diagramme (*stem-leaf display*):

Diese werden heutzutage relativ selten und nur für kleine Datensätze verwendet. Dabei arbeitet man mit Stichprobenwerten, die auf ganze Zahlen gerundet sind. Sei $(x_1, \dots, x_n)$ eine Stichprobe von solchen Werten, die Ausprägungen eines quantitativen Merkmals sind. Zunächst teilt man den Wertebereich $[x_{(1)}, x_{(n)}]$ in Klassen gleicher Breite 10^d , $d \in \mathbb{N}$, wobei jede Klasse mit den ersten Ziffern der dazugehörigen Beobachtungen markiert wird. Zum Beispiel, wenn die Klasseneinteilung so aussieht



werden die Klassen $[100(i-1), 100i)$ mit den Zahlen i markiert und auf der y -Achse wie folgt aufgetragen:



Auf diese Weise wird der Stamm des Baumes festgelegt. In jeder Klasse ordnet man Beobachtungen ihrer Größe nach und rundet sie auf die Stelle, die nach der gewählten Genauigkeit des Stammes folgt. Als Beispiel erhält man aus $127 \rightarrow 130$, aus $652 \rightarrow 650$ usw. und trägt diese Beobachtungen als Blätter des Baums horizontal ihrer Reihenfolge nach als 3 in Klasse 1 und 5 in Klasse 6 auf. Dabei darf man nicht vergessen, die Einheit zu notieren: $1/3 = 130$, um sich das Rückrechnen zu ermöglichen. Bei der Wahl der Klassenanzahl m hält man sich an die Faustregel $m \approx 10 \log_{10} n$, um einerseits den Dateiverlust durch das unnötige Runden zu minimieren und andererseits das Diagramm so übersichtlich wie möglich zu halten.

Bemerkung 2.1.1

Die in Abschnitt 2.1.1 betrachteten Methoden dienen der Visualisierung von (Zähl-) Dichten der Verteilung eines beobachteten Merkmals X . Aus dem Histogramm kann z.B. die Interpretation der Form der Dichte abgelesen werden:

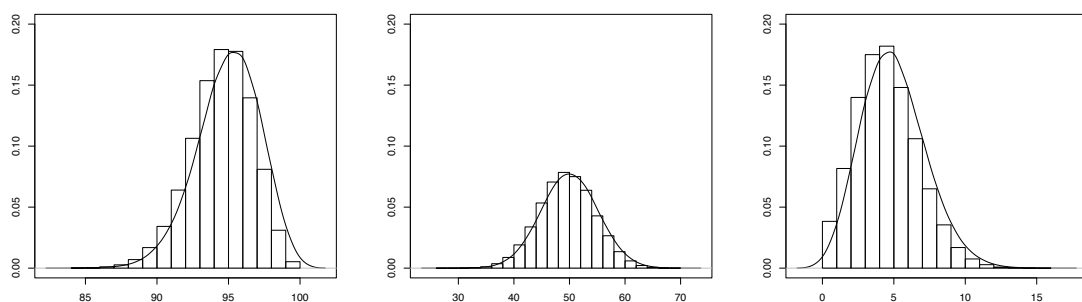


Abb. 2.1: Das Histogramm der Daten mit einer rechtssteilen (linksschiefen), symmetrischen und linkssteilen (rechtsschiefen) Verteilung und ihre Dichte.

Ist die zugrundeliegende Verteilung F_X symmetrisch bzw. linkssteil (rechtsschief) oder rechtssteil (linksschief) (vgl. Abb. 2.1) oder ist sie unimodal (d.h. eingipflig), bimodal (d.h. mit 2 Gipfeln) oder multimodal (also mit mehreren Gipfeln) (vgl. Abb. 2.2).

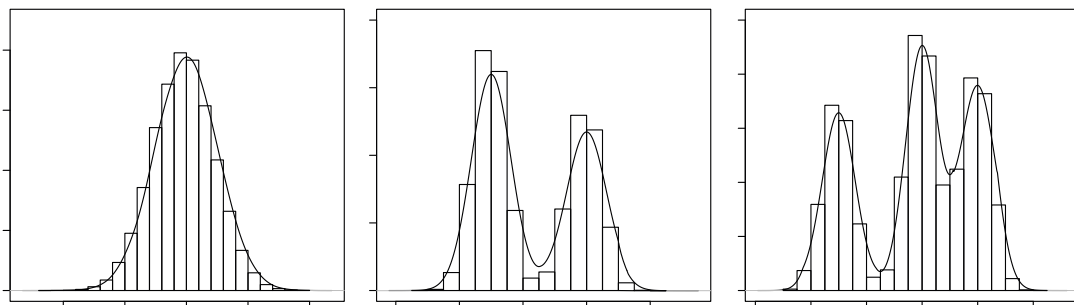


Abb. 2.2: Histogramm der Daten mit der Dichte einer unimodalen, bimodalen und multimodalen Verteilung

2.1.2 Empirische Verteilungsfunktion

Es sei eine konkrete Stichprobe $(x_1, \dots, x_n)$ gegeben, die eine Realisierung des statistischen Modells $(X_1, \dots, X_n)$ ist, wobei $X_1, \dots, X_n$ unabhängige identisch verteilte Zufallsvariablen mit Verteilungsfunktion $F_X : X_i \stackrel{d}{=} X \sim F_X$ sind. Wie kann die unbekannte Verteilungsfunktion F_X aus den Daten $(x_1, \dots, x_n)$ rekonstruiert (die Statistiker sagen „geschätzt“) werden? Dies ist mit Hilfe der sogenannten empirischen Verteilungsfunktion möglich:

Definition 2.1.2

1. Die Funktion $\hat{F}_n(x) = \#\{x_i : x_i \leq x, i = 1, \dots, n\}/n$, $\forall x \in \mathbb{R}$ heißt *empirische Verteilungsfunktion der konkreten Stichprobe* $(x_1, \dots, x_n)$. Dabei gilt $\hat{F}_n : \mathbb{R}^{n+1} \rightarrow [0, 1]$, weil $\hat{F}_n(x) = \varphi(x_1, \dots, x_n, x)$.
2. Die mit $x \in \mathbb{R}$ indizierte Zufallsvariable $\hat{F}_n : \Omega \times \mathbb{R} \rightarrow [0, 1]$ heißt *empirische Verteilungsfunktion der Zufallsstichprobe* $(X_1, \dots, X_n)$, wenn

$$\hat{F}_n(x, \omega) = \hat{F}_n(x) = \frac{1}{n} \#\{X_i, i = 1, \dots, n : X_i(\omega) \leq x\}, \quad x \in \mathbb{R}.$$

Äquivalent zur Definition 2.1.2 kann man

$$\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(x_i \leq x), \quad x \in \mathbb{R}$$

schreiben, wobei

$$\mathbb{I}(x \in A) = \begin{cases} 1, & x \in A \\ 0, & \text{sonst.} \end{cases}$$

Es gilt

$$\hat{F}_n(x) = \begin{cases} 1, & x \geq x_{(n)}, \\ \frac{i}{n}, & x_{(i)} \leq x < x_{(i+1)}, \quad i = 1, \dots, n-1, \\ 0, & x < x_{(1)}. \end{cases}$$

für $x_{(1)} < x_{(2)} < \dots < x_{(n)}$.

Dabei ist die Höhe des Sprungs an Stelle $x_{(i)}$ gleich der relativen Häufigkeit f_i des Wertes $x_{(i)}$. Falls $x_{(i)} = x_{(i+1)}$ für ein $i \in \{1, \dots, n\}$, so tritt der Wert i/n nicht auf. In Abbildung 2.3 sieht man, dass $\hat{F}_n(x)$ eine rechtsstetige monoton nichtfallende Treppenfunktion ist, für die

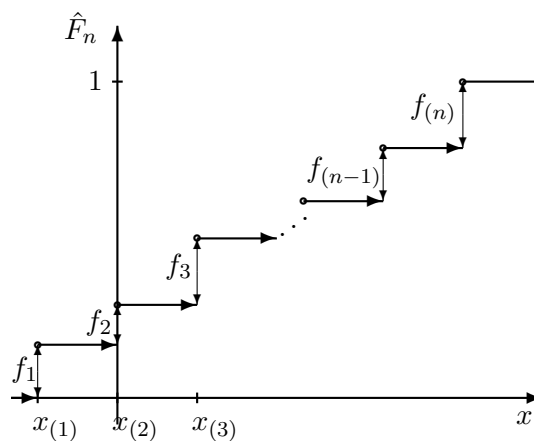


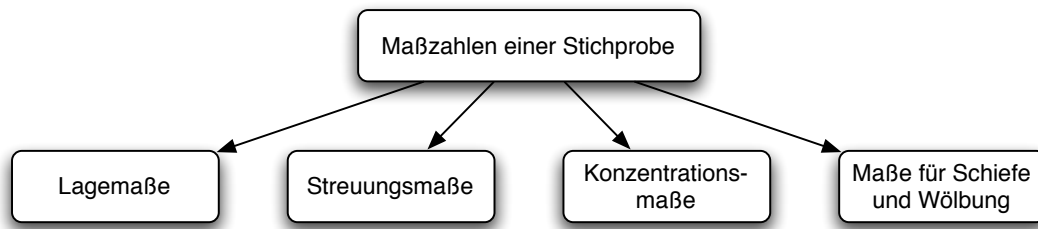
Abb. 2.3: Eine typische empirische Verteilungsfunktion

$\hat{F}_n(x) \xrightarrow{x \rightarrow -\infty} 0$, $\hat{F}_n(x) \xrightarrow{x \rightarrow \infty} 1$ gilt.

Übungsaufgabe 2.1.1

Zeigen Sie, dass $\hat{F}_n(x)$ eine Verteilungsfunktion ist.

2.2 Beschreibung von Verteilungen



Es sei eine konkrete Stichprobe $(x_1, \dots, x_n)$ gegeben. Im Folgenden werden Kennzahlen (die sogenannten Maße) dieser Stichprobe betrachtet, welche die wesentlichen Aspekte der der Stichprobe zugrundeliegenden Verteilung wiedergeben:

1. Wo liegen die Werte x_i (Mittel, Ordnungsstatistiken, Quantile)? $\implies$ Lagemaße
2. Wie stark streuen die Werte x_i (Varianz) $\implies$ Streuungsmaße
3. Wie stark sind die Werte x_i in gewissen Bereichen von $\mathbb{R}$ konzentriert $\implies$ Konzentrationsmaße
4. Wie schief bzw. gewölbt ist die Verteilung von X $\implies$ Maße für Schiefe und Wölbung

2.2.1 Lagemaße

Man unterscheidet folgende wichtige Lagemaße:

- Mittelwerte: Stichprobenmittel (arithmetisch), geometrisches und harmonisches Mittel, gewichtetes Mittel, getrimmtes Mittel
- Ordnungsstatistiken und Quantile, insbesondere Median und Quartile
- Modus

Betrachten wir sie der Reihe nach:

1. *Mittelwertbildung*: Seit der Antike kennt man mindestens 3 Arten der *Mittelberechnung* von n Zahlen $(x_1, \dots, x_n)$:

- *arithmetisch*: $\bar{x}_n = 1/n \sum_{i=1}^n x_i$, $\forall x_1, \dots, x_n \in \mathbb{R}$,
- *geometrisch*: $x_n^g = \sqrt[n]{x_1 \cdot \dots \cdot x_n}$, $x_1, \dots, x_n > 0$,
- *harmonisch*: $x_n^h = \left(1/n \sum_{i=1}^n x_i^{-1}\right)^{-1}$, $x_1, \dots, x_n \neq 0$.

- a) Das *arithmetische Mittel* wird in der Statistik am meisten benutzt, weil es keine Voraussetzungen über den Wertebereich von $x_1, \dots, x_n$ braucht. Es wird auch *Stichprobenmittel* genannt. Offensichtlich ist $\bar{x}_n$ ein Spezialfall des sogenannten gewichteten Mittels $x_n^w = \sum_{i=1}^n w_i x_i$, wobei für die Gewichte $w_i \geq 0 \quad \forall i = 1, \dots, n$ und $\sum_{i=1}^n w_i = 1$ gilt. Als eine natürliche Gewichtewahl kommt $w_i = 1/n$, $\forall i = 1, \dots, n$ bei einer konkreten Stichprobe $(x_1, \dots, x_n)$ in Frage. Die Summe aller Abweichungen von $\bar{x}_n$ ist Null, denn $\sum_{i=1}^n (x_i - \bar{x}_n) = n\bar{x}_n - n\bar{x}_n = 0$, d.h. $\bar{x}_n$ stellt geometrisch

den Schwerpunkt der Werte x_i dar, falls jedem Punkt eine Einheitsmasse zugeordnet wird. Wenn es in der Stichprobe große Ausreißer gibt, so beeinflussen sie das Stichprobenmittel entscheidend und erschweren so die objektive Datenanalyse. Deshalb verwendet man oft die robuste Version des arithmetischen Mittels, das sogenannte *getrimmte Mittel*:

$$\tilde{x}_n^{(k)} = \frac{1}{n-2k} \sum_{i=k+1}^{n-k} x_{(i)},$$

bei dessen Berechnung die k kleinsten und k größten Ausreißer ausgelassen werden, wobei $k \ll n/2$.

- b) Das *geometrische Mittel* wird hauptsächlich bei der Beobachtung von Wachstums- und Zinsfaktoren verwendet. Sei $x_i = B_i/B_{i-1}$, $i = 1, \dots, n$ der Wachstumsfaktor des Merkmals B_i , das in den Jahren $i = 1, \dots, n$ beobachtet wurde (z.B. Inflationsfaktor). Dann ist $B_n = B_0 \cdot x_1 \cdot \dots \cdot x_n$ und somit wäre der Zins im Jahre n

$$B_n^g = B_0 \cdot x_1 \cdot \dots \cdot x_n = B_0 \cdot (x_n^g)^n.$$

Für das geometrische Mittel gilt

$$\log x_n^g = \frac{1}{n} \sum_{i=1}^n \log x_i \leq \log \left(\frac{1}{n} \sum_{i=1}^n x_i \right)$$

wegen der Konkavität des Logarithmus, d.h. $\log x_n^g = \overline{\log x_n} \leq \log \bar{x}_n$ und somit $x_n^g \leq \bar{x}_n$, wobei $x_n^g = \bar{x}_n$ genau dann, wenn $x_1 = \dots = x_n$.

- c) Das *harmonische Mittel* wird bei der Ermittlung von z.B. durchschnittlicher Geschwindigkeiten gebraucht.

Beispiel 2.2.1

Seien x_i Geschwindigkeiten mit denen Bauteile eine Produktionslinie der Länge l durchlaufen. Die gesamte Bearbeitungszeit ist $l/x_1 + \dots + l/x_n$ und die Durchschnittslaufgeschwindigkeit

$$\frac{l + \dots + l}{l/x_1 + \dots + l/x_n} = x_n^h.$$

Es gilt $x_{(1)} \leq x_n^h \leq x_n^g \leq \bar{x}_n \leq x_{(n)}$ für $x_i > 0$, $i = 1, \dots, n$.

Übungsaufgabe 2.2.1

Beweisen Sie diese Relation per Induktion bzgl. n .

2. Ordnungsstatistiken und Quantile

Definition 2.2.1

Die *Ordnungsstatistiken* $x_{(i)}$, $i = 1, \dots, n$ der Stichprobe $(x_1, \dots, x_n)$ sind durch die messbare Permutation $\varphi(x_1, \dots, x_n)$ gegeben, so dass

$$x_{(i)} = \min \{x_j : \#\{k : x_k \leq x_j\} \geq i\}, \quad \forall i = 1, \dots, n.$$

Somit gilt $x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(n)}$. Dieselbe Definition kann auch auf der Modellebene gegeben werden.

Definition 2.2.2

- a) Sei nun X die Zufallsvariable, die das Merkmal modelliert. Sei F_X ihre Verteilungsfunktion. Die verallgemeinerte Inverse von F_X , definiert durch

$$F_X^{-1}(y) = \inf \{x : F_X(x) \geq y\}, \quad y \in [0, 1],$$

heißt *Quantilfunktion* von F_X bzw. X . Es gilt $F_X^{-1} : [0, 1] \rightarrow \mathbb{R} \cup \{\pm\infty\}$. Die Zahl $F_X^{-1}(\alpha)$, $\alpha \in [0, 1]$ wird α -*Quantil* von F_X genannt.

- b) • $F_X^{-1}(0, 25)$ heißt *unteres Quartil*,
 • $F_X^{-1}(0, 75)$ heißt *oberes Quartil*,
 • $F_X^{-1}(0, 5)$ heißt der *Median* der Verteilung von X .

Zwischen Ordnungsstatistiken und Quantilen besteht ein enger Zusammenhang. So bedeutet $F_X^{-1}(\alpha)$, $\alpha \in (0, 1)$, dass ca. $\alpha \cdot 100\%$ aller Merkmalsausprägungen in der Stichprobe $(x_1, \dots, x_n)$ unter $F_X^{-1}(\alpha)$ und ca. $(1 - \alpha) \cdot 100\%$ über $F_X^{-1}(\alpha)$ liegen (im absolut stetigen Fall). Insbesondere gilt $F_X^{-1}(\alpha) \approx x_{([n\alpha])}$, deshalb werden Ordnungsstatistiken auch *empirische Quantile* genannt. Dabei ist x_α definiert als

$$x_\alpha = \begin{cases} x_{([n\alpha]+1)}, & n\alpha \notin \mathbb{N} \\ 1/2(x_{([n\alpha])} + x_{([n\alpha]+1)}), & n\alpha \in \mathbb{N} \end{cases}.$$

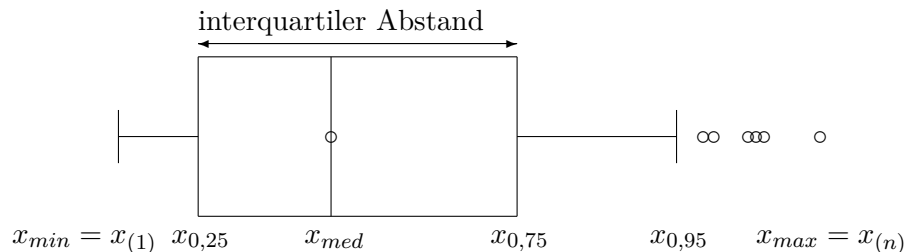
Dies ist die allgemeine Definition des *empirischen α -Quantils*.

Der *empirische Median* ist

$$x_{med} = \begin{cases} x_{(\frac{n+1}{2})}, & n \text{ ungerade} \\ \frac{1}{2} \left(x_{(\frac{n}{2})} + x_{(\frac{n}{2}+1)} \right), & n \text{ gerade.} \end{cases}$$

Somit sind mindestens 50% aller Stichprobenwerte kleiner gleich und 50% größer gleich x_{med} . Der Median ist ein Lagemaß, das ein robuster Ersatz für den Mittelwert darstellt, denn er ist bzgl. Ausreißern in der Stichprobe nicht sensibel.

Die oben genannten Statistiken werden in einem *Box-Plot* zusammengefasst und grafisch dargestellt:



Manchmal werden $x_{(1)}$ und $x_{(n)}$ durch $x_{0,05}$ und $x_{0,95}$ ersetzt. Die restlichen Werte werden darüber hinaus als Einzelpunkte auf der x -Achse abgebildet. Dann liegt ein sogenannter *modifizierter Box-Plot* vor.

3. *Modus*: Sei $(x_1, \dots, x_n)$ eine Stichprobe, die aus n unabhängigen Realisierungen des Merkmals X besteht. Sei $(p(x))$ $f(x)$ die (Zähl-) Dichte von X , wobei die Verteilung von X unimodal ist.

Definition 2.2.3

- a) Der Wert $x_{mod} = \arg \max f(x)$ ($\arg \max p(x)$) wird der *Modus der Verteilung von X* genannt (vgl. Abb. 2.4).
- b) Empirisch wird $\hat{x}_{mod}$ als $\frac{c_{m-1} + c_m}{2}$ für $m = \arg \max f_i$ definiert, also als die Mitte des Intervalls mit der größten Häufigkeit des Vorkommens in der Stichprobe, falls dieser eindeutig bestimmbar ist.

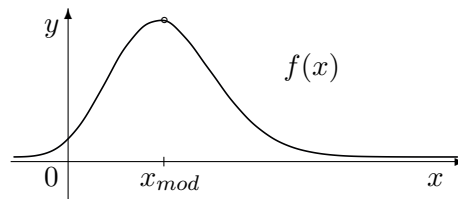


Abb. 2.4: Veranschaulichung des Modus

Den Mittelwert $\bar{x}_n$, Median x_{med} und Modus x_{mod} kann man auch wie folgt definieren:

$$\bar{x}_n = \arg \min_{x \in \mathbb{R}} \sum_{i=1}^n (x_i - x)^2$$

$$x_{med} = \arg \min_{x \in \mathbb{R}} \sum_{i=1}^n |x_i - x|$$

$$\hat{x}_{mod} = \frac{c_{m-1} + c_m}{2}, \quad \text{wobei } m = \arg \min_{j=1, \dots, n} \sum_{i=1}^n \mathbb{I}(x_i \notin (c_{j-1}, c_j])$$

Übungsaufgabe 2.2.2

Zeigen Sie die Äquivalenz der oben genannten Definitionen des Mittelwerts $\bar{x}_n$, Medians x_{med} und des Modus x_{mod} zu den bekannten Definitionen.

Die Größen $\bar{x}_n$, x_{med} und $\hat{x}_{mod}$ können auch zur Beschreibung der Symmetrie einer unimodalen Verteilung F_X von Daten $(x_1, \dots, x_n)$ verwendet werden, da

- bei symmetrischen Verteilung F_X gilt $\bar{x}_n \approx x_{med} \approx \hat{x}_{mod}$
- bei linkssteilen Verteilung F_X gilt $\hat{x}_{mod} < x_{med} < \bar{x}_n$
- bei rechtssteilen Verteilung F_X gilt $\bar{x}_n < x_{med} < \hat{x}_{mod}$.

2.2.2 Streuungsmaße

Bekannte Streuungsmaße einer konkreten Stichprobe $(x_1, \dots, x_n)$ sind die folgenden Größen:

- *Spannweite* $x_{(n)} - x_{(1)}$,
- *empirische Varianz* $\bar{s}_n^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^2$,
- *Stichprobenvarianz* $s_n^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}_n)^2 = \frac{n}{n-1} \bar{s}_n^2$,
- *empirische Standardabweichungen* $\bar{s}_n = \sqrt{\bar{s}_n^2}$, $s_n = \sqrt{s_n^2}$,
- *empirischer Variationskoeffizient* $\gamma_n = s_n / \bar{x}_n$, falls $\bar{x}_n > 0$.

Die Spannweite zeigt die *maximale Streuung* in den Daten, wobei sich die empirische Varianz mit der *mittleren quadratischen Abweichung* vom Stichprobenmittel auseinandersetzt. Hier sind einige Eigenschaften von $\bar{s}_n^2$ (bzw. s_n^2 , da sie sich nur durch einen Faktor unterscheiden):

Lemma 2.2.1

1. Für jedes $b \in \mathbb{R}$ gilt

$$\sum_{i=1}^n (x_i - b)^2 = \sum_{i=1}^n (x_i - \bar{x}_n)^2 + n(\bar{x}_n - b)^2$$

und somit für $b = 0$

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

2. *Transformationsregel:*

Falls die Daten $(x_1, \dots, x_n)$ linear transformiert werden, d.h. $y_i = ax_i + b$, $a \neq 0$, $b \in \mathbb{R}$, dann gilt

$$\bar{s}_{n,y}^2 = a^2 \bar{s}_{n,x}^2 \quad \text{bzw.} \quad \bar{s}_{n,y} = |a| \bar{s}_{n,x},$$

wobei

$$\bar{s}_{n,y}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \bar{y}_n)^2, \quad \bar{s}_{n,x}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^2.$$

Beweis 1. Es gilt:

$$\begin{aligned} \sum_{i=1}^n (x_i - b)^2 &= \sum_{i=1}^n (x_i - \bar{x}_n + \bar{x}_n - b)^2 \\ &= \sum_{i=1}^n (x_i - \bar{x}_n)^2 + 2 \sum_{i=1}^n (x_i - \bar{x}_n) \cdot (\bar{x}_n - b) + \sum_{i=1}^n (\bar{x}_n - b)^2 \\ &= \sum_{i=1}^n (x_i - \bar{x}_n)^2 + 2(\bar{x}_n - b) \cdot \underbrace{\sum_{i=1}^n (x_i - \bar{x}_n)}_{=0} + n(\bar{x}_n - b)^2, \quad \forall b \in \mathbb{R}. \end{aligned}$$

2. Es gilt:

$$\bar{s}_{n,y}^2 = \frac{1}{n} \sum_{i=1}^n (ax_i + b - a\bar{x}_n - b)^2 = \frac{a^2}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^2 = a^2 \bar{s}_{n,x}^2.$$

□

Der Skalierungsunterschied zwischen $\bar{s}_n^2$ und s_n^2 ist den Eigenschaften der *Erwartungstreue* von s_n^2 zu verdanken, die später im Laufe dieser Vorlesung behandelt wird, und besagt, dass für eine Zufallsstichprobe $(X_1, \dots, X_n)$ mit X_i unabhängig identisch verteilt, $X_i \sim X$, $\text{Var } X = \sigma^2 \in (0, \infty)$ gilt $\mathbb{E}s_n^2 = \sigma^2$, wobei $\mathbb{E}\bar{s}_n^2 = \frac{n}{n-1}\sigma^2 \xrightarrow{n \rightarrow \infty} \sigma^2$. Das heißt, während bei der Verwendung von s_n^2 zur Schätzung von σ^2 kein Fehler „im Mittel“ gemacht wird, ist diese Aussage für $\bar{s}_n^2$ nur asymptotisch (für große Datenmengen n) richtig.

Aufgrund von $\sum_{i=1}^n (x_i - \bar{x}_n) = 0$ ist z.B. $x_n - \bar{x}_n$ durch $x_i - \bar{x}_n$, $i = 1, \dots, n-1$ bestimmt. Somit verringert sich die *Anzahl der Freiheitsgrade* in der Summe $\sum_{i=1}^n (x_i - \bar{x}_n)^2$ um 1 und somit scheint die Normierung $\frac{1}{n-1}$ plausibel zu sein.

Die *Standardabweichungen* $\bar{s}_n$ und s_n werden verwendet, damit man die selben Einheiten (und nicht ihre Quadrate, also z.B. Euro und nicht Euro²) erhält. Für normalverteilte Stichproben ($X \sim N(\mu, \sigma^2)$) liefert $\bar{s}_n$ auch die „*k*-Sigma-Regel“ (vgl. Vorlesung WR), die besagt, dass in den Intervallen

$$\begin{array}{ll} [\bar{x}_n - \bar{s}_n, \bar{x}_n + \bar{s}_n] & \text{ca. } 68\%, \\ [\bar{x}_n - 2\bar{s}_n, \bar{x}_n + 2\bar{s}_n] & \text{ca. } 95\%, \\ [\bar{x}_n - 3\bar{s}_n, \bar{x}_n + 3\bar{s}_n] & \text{ca. } 99\% \end{array}$$

aller Daten liegen.

Der Vorteil vom *empirischen Variationskoeffizienten* ist, dass er *maßstabsunabhängig* ist und somit den Vergleich von Streuungseigenschaften unterschiedlicher Stichproben zulässt.

2.2.3 Konzentrationsmaße

Insbesondere in den Wirtschaftswissenschaften interessiert man sich oft für die Konzentration von Merkmalsausprägungen in der Stichprobe, z.B. wie sich das Familieneinkommen einer demographischen Einheit auf unterschiedliche Einkommensbereiche (Vielverdiener, Mittelstand, Wenigverdiener) aufteilt, oder wie sich der Markt auf Marktanbieter aufteilt (Marktkonzentration). Dabei ist es wünschenswert, diese Relation mit Hilfe weniger Zahlen oder einer Grafik zum Ausdruck zu bringen. Dies ist mit Hilfe folgender Stichprobenfunktionen möglich:

- *Lorenzkurve* L ,
- *Gini-Koeffizient* G ,
- *Konzentrationsrate* CR_g ,
- *Herfindahl-Index* H .

1. Die Lorenzkurve wurde von M. Lorenz am Anfang des XX. Jahrhunderts für die Charakterisierung der Vermögenskonzentration benutzt. Sei $(x_1, \dots, x_n)$ eine Stichprobe, die in aufsteigender Reihenfolge geordnet werden muss: $(x_{(1)}, \dots, x_{(n)})$. Die *Lorenzkurve* verbindet Punkte

$$(0, 0), (u_1, v_1), \dots, (u_n, v_n), (1, 1)$$

durch Liniensegmente, wobei $u_j = j/n$ der Anteil der j kleinsten Merkmalsträger und $v_j = \sum_{i=1}^j x_{(i)} / \sum_{i=1}^n x_i$ die kumulierte relative Merkmalssumme ist. Der Grundgedanke

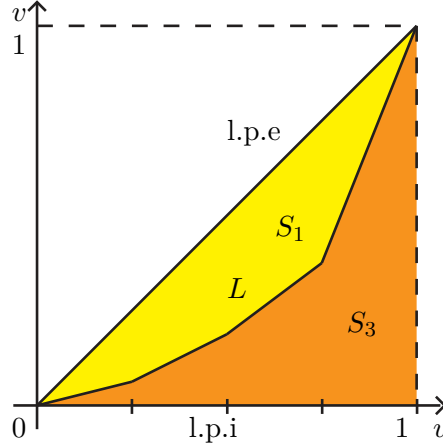


Abb. 2.5: Abbildung einer typischen Lorenzkurve

ist darzustellen, welcher Anteil des Merkmalsträgers auf welchen Anteil der Gesamtmerkmalssumme entfällt. Zum Beispiel lassen sich dadurch Aussagen wie etwa „Auf 20% aller Haushalte im Land entfällt 78% des Gesamteinkommens“ machen. Eine Interpretation der Lorenzkurve L ist nur an den Knoten (u_j, v_j) möglich: „Auf $u_j \cdot 100\%$ der kleinsten Merkmalsträger konzentrieren sich $v_j \cdot 100\%$ der Merkmalssumme“. Dabei liegt L auf $[0, 1]^2$ immer zwischen der „line of perfect equality“ (l.p.e.) $v_i = u_i \quad \forall i$ (Einkommen ist absolut gleichmäßig—also „gerecht“—verteilt) und „line of perfect inequality“ (l.p.i.) $v = 0, u \in [0, 1)$ und $(1, 1)$ (das Gesamteinkommen besitzt nur die reichste Familie) und ist immer monoton und konvex. Auf Modellebene gibt es ein Analogon der Lorenzkurve. Dieses ist

$$L = \left\{ (u, v) \in [0, 1]^2 : \quad v = \frac{\int_0^u F_X^{-1}(t) dt}{\int_0^1 F_X^{-1}(t) dt}, \quad u \in [0, 1] \right\},$$

wobei

$$\mathbb{E}X = \int_0^1 F_X^{-1}(t) dt$$

(vgl. WR Satz 4.3.2). Dementsprechend können die Knoten (u_i, v_j) der oben eingeführten empirischen Lorenzkurve als

$$v_j = \frac{\sum_{i=1}^j \frac{x_{(i)}}{n}}{\bar{x}_n}$$

interpretiert werden.

2. Der *Gini-Koeffizient* G ist gegeben durch $G = S_1/S_2$, wobei S_1 die Fläche zwischen der Lorenzkurve L und der Diagonalen $v = u$, S_2 die Fläche zwischen der Diagonalen und der u -Achse ($= 1/2|[0, 1]^2| = 1/2$) ist.

Satz 2.2.1 (Darstellung des Gini-Koeffizienten):

Es gilt

$$G = 2S_1 = \frac{2 \sum_{i=1}^n ix_{(i)}}{n \sum_{i=1}^n x_i} - \frac{n+1}{n}.$$

Beweis Beginnen wir damit, die Darstellung $G = (n+1)/n - 2\bar{v}_n$ zu zeigen. Nach Definition ist

$$G = \frac{S_1}{S_2} = \frac{S_2 - S_3}{S_2} = 1 - \frac{S_3}{S_2} = 1 - 2S_3,$$

wobei S_3 die Fläche zwischen der Lorenzkurve und der x -Achse ist (vgl. Abb. 2.5). Berechnen wir S_3 :

$S_3 = \sum_{j=1}^n F_j$, wobei $F_j = 1/n \cdot v_{j-1} + \frac{1}{2} \frac{1}{n} \cdot (v_j - v_{j-1}) = \frac{1}{2n}(v_j + v_{j-1})$ die Fläche unter einem Liniensegment der Lorenzkurve ist (vgl. Abb. 2.6). Es gilt

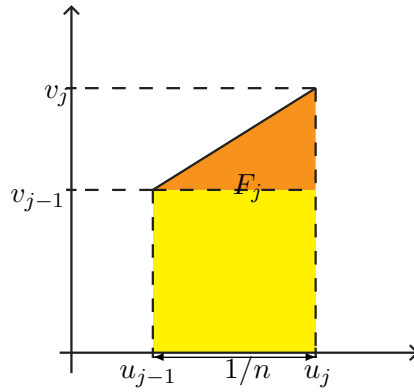


Abb. 2.6: Liniensegment der Lorenzkurve

$$S_3 = \frac{1}{2n} \sum_{j=1}^n (v_j + v_{j-1}) = \frac{1}{2n} \left(2 \sum_{j=1}^n v_j - 1 \right) = \bar{v}_n - \frac{1}{2n},$$

somit

$$G = 1 - 2\bar{v}_n + \frac{1}{n} = \frac{n+1}{n} - 2\bar{v}_n.$$

Beweisen wir jetzt, dass

$$G = \frac{2 \sum_{i=1}^n ix_{(i)}}{n \sum_{i=1}^n x_i} - \frac{n+1}{n}$$

ist. Sei $w = \sum_{i=1}^n ix_{(i)}$. Aufgrund der Definition von v_j gilt $s_j = \sum_{i=1}^j x_{(i)} = s_n \cdot v_j$, $\forall j = 1, \dots, n$ und $x_{(i)} = s_i - s_{i-1}$, $s_0 = 0$. Daher erhalten wir

$$\begin{aligned} w &= \sum_{i=1}^n i(s_i - s_{i-1}) = \sum_{i=1}^n is_i - \sum_{i=0}^{n-1} (i+1)s_i = ns_n - \sum_{i=0}^{n-1} s_i \\ &= (n+1)s_n - \sum_{i=1}^n s_i = (n+1)s_n - s_n \cdot \sum_{i=1}^n v_i = (n+1)s_n - s_n \cdot n\bar{v}_n \end{aligned}$$

und somit

$$\frac{2\omega}{ns_n} - \frac{n+1}{n} = \frac{2w - (n+1)s_n}{ns_n} = \frac{2(n+1)s_n - 2s_n n\bar{v}_n - (n+1)s_n}{ns_n} = \frac{n+1}{n} - 2\bar{v}_n = G.$$

□

Es gilt $G \in [0, (n-1)/n]$, wobei

$$\begin{aligned} G_{\min} &= 0 & \text{bei } x_1 = x_2 = \dots = x_n & \quad \text{„perfect equality“,} \\ G_{\max} &= \frac{n-1}{n} & \text{bei } x_1 = \dots = x_{n-1} = 0, x_n \neq 0 & \quad \text{„perfect inequality“}. \end{aligned}$$

Somit hängt $G_{\max}$ vom Datenumfang ab. Um dies zu vermeiden, betrachtet man oft den normierten Gini-Koeffizienten

$$G^* = \frac{G}{G_{\max}} = \frac{n}{n-1}G \in [0, 1]$$

(Lorenz-Münzner-Koeffizient).

3. Konzentrationsrate CR_g :

In den Punkten 1) und 2) betrachteten wir die *relative Konzentration*, wie etwa bei der Fragestellung „Wieviel % der Familien teilen sich wieviel % des Gesamteinkommens?“. Dabei beantwortet die Konzentrationsrate die Frage „Wieviele Familien haben wieviel Prozent des Gesamteinkommens?“ für die g reichsten Familien, somit wird auch die absolute Anzahl aller Familien berücksichtigt.

Sei $g \in \{1, \dots, n\}$ und seien $x_{(1)} \leq \dots \leq x_{(n)}$ die Ordnungsstatistiken der Stichprobe $(x_1, \dots, x_n)$. Für $i \in \{1, \dots, n\}$ sei

$$p_i = \frac{x_{(i)}}{\sum_{j=1}^n x_j} = \frac{x_{(i)}}{n\bar{x}_n} \quad (2.2.1)$$

der Merkmalsanteil der i -ten Einheit.

Dann gibt die *Konzentrationsrate* $CR_g = \sum_{i=n-g+1}^n p_i$ wieder, welcher Anteil des Gesamteinkommens von g reichsten Familien gehalten wird.

4. Der *Herfindahl-Index* ist definiert durch $H = \sum_{i=1}^n p_i^2$, wobei der Merkmalsanteil p_i nach (2.2.1) definiert ist. Bei der gleichen Verteilung des Einkommens ($x_1 = x_2 = \dots = x_n$) gilt $H_{\min} = 1/n$, bei völlig ungerechter Verteilung ($x_1 = \dots = x_{n-1} = 0, x_n \neq 0$) $H_{\max} = 1$. Sonst gilt $H \in [H_{\min}, H_{\max}]$, also $1/n \leq H \leq 1$. H ist umso kleiner, je gerechter das Gesamteinkommen verteilt ist.

2.2.4 Maße für Schiefe und Wölbung

Im Vorlesungsskript WR, Abschnitt 4.5 S. 99 wurden folgende Maße für Schiefe bzw. Wölbung der Verteilung einer Zufallsvariable X eingeführt:

Schiefe oder Symmetriekoeffizient:

$$\gamma_1 = \frac{\mu'_3}{\sigma^3} = E(\tilde{X}^3),$$

wobei

$$\mu'_k = \mathbb{E}(X - \mathbb{E}X)^k, \quad \sigma^2 = \mu'_2 = \text{Var } X, \quad \tilde{X} = \frac{X - \mathbb{E}X}{\sigma}.$$

Wölbung (*Exzess*):

$$\gamma_2 = \frac{\mu'_4}{\sigma^4} - 3 = \mathbb{E}(\tilde{X}^4) - 3,$$

vorausgesetzt, dass $\mathbb{E}(X^4) < \infty$. Für ihre Bedeutung und Interpretation siehe die oben genannten Seiten des WR-Vorlesungsskriptes. Falls nun das Merkmal X statistisch in einer Stichprobe $(x_1, \dots, x_n)$ beobachtet wird, wie können γ_1 und γ_2 aus diesen Daten geschätzt und interpretiert werden?

Als Schätzer für das k -te zentrierte Moment $\mu'_k = \mathbb{E}(X - \mathbb{E}X)^k$, $k \in \mathbb{N}$ schlagen wir

$$\hat{\mu}'_k = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^k$$

vor, die Varianz σ^2 wird durch

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

geschätzt. Somit bekommt man den Momentenkoeffizient an der Schiefe (engl. „skewness“)

$$\hat{\gamma}_1 = \frac{\hat{\mu}'_3}{\bar{s}_n^3} = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^3}{\left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^2\right)^{3/2}}.$$

Falls die Verteilung von X linksschief ist, überwiegen negative Abweichungen im Zähler und somit gilt $\hat{\gamma}_1 < 0$ für linksschiefe Verteilungen. Analog gilt $\hat{\gamma}_1 \approx 0$ für symmetrische und $\hat{\gamma}_1 > 0$ für rechtsschiefe Verteilungen.

Das *Wölbungsmaß von Fisher* (engl. „kurtosis“) ist gegeben durch

$$\hat{\gamma}_2 = \frac{\hat{\mu}'_4}{\bar{s}_n^4} - 3 = \frac{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^4}{\left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x}_n)^2\right)^2} - 3.$$

Falls $\hat{\gamma}_2 > 0$ so ist die Verteilung von X steilgipflig, für $\hat{\gamma}_2 < 0$ ist sie flachgipflig. Falls $X \sim N(\mu, \sigma^2)$, so gilt $\hat{\gamma}_2 \approx 0$. Die Ursache dafür ist, dass die steilgipfligen Verteilungen schwerere Tails haben als die flachgipfligen. Als Maß dient dabei die Normalverteilung, für die $\gamma_1 = \gamma_2 = 0$ und somit $\hat{\gamma}_1 \approx 0$, $\hat{\gamma}_2 \approx 0$. So definiert, sind $\hat{\gamma}_1$ und $\hat{\gamma}_2$ nicht resistent gegenüber Ausreißern. Eine robuste Variante von $\hat{\gamma}_1$ ist beispielsweise durch den sogenannten *Quantilkoeffizienten der Schiefe*

$$\hat{\gamma}_q(\alpha) = \frac{(x_{1-\alpha} - x_{med}) - (x_{med} - x_\alpha)}{x_{1-\alpha} - x_\alpha}, \quad \alpha \in (0, 1/2)$$

gegeben.

Für $\alpha = 0,25$ erhält man den Quartilkoeffizienten. $\hat{\gamma}_q(\alpha)$ misst den Unterschied zwischen der Entfernung des α - und $(1 - \alpha)$ -Quantils zum Median. Bei linkssteilen (bzw. rechtssteilen) Verteilungen liegt das (untere) x_α -Quantil näher an (bzw. weiter entfernt von) dem Median. Somit gilt

- $\hat{\gamma}_q(\alpha) > 0$ für linkssteile Verteilungen,

- $\hat{\gamma}_q(\alpha) < 0$ für rechtssteile Verteilungen,
- $\hat{\gamma}_q(\alpha) = 0$ für symmetrische Verteilungen.

Durch das zusätzliche Normieren (Nenner) gilt $-1 \leq \hat{\gamma}_q(\alpha) \leq 1$.

2.3 Quantilplots (Quantil-Grafiken)

Nach der ersten beschreibenden Analyse eines Datensatzes $(x_1, \dots, x_n)$ soll überlegt werden, mit welcher Verteilung diese Stichprobe modelliert werden kann. Hier sind die sogenannten *Quantilplots* behilflich, da sie grafisch zeigen, wie gut die Daten $(x_1, \dots, x_n)$ mit dem Verteilungsgesetz G übereinstimmen, wobei G die Verteilungsfunktion einer hypothetischen Verteilung ist.

Sei X eine Zufallsvariable mit (unbekannter) Verteilungsfunktion F_X . Auf Basis der Daten $(X_1, \dots, X_n)$, X_i unabhängig identisch verteilt und $X_i \stackrel{d}{=} X$ möchte man prüfen, ob $F_X = G$ für eine bekannte Verteilungsfunktion G gilt. Die Methode der *Quantil-Grafiken* besteht darin, dass man die entsprechenden Quantil-Funktionen $\hat{F}_n^{-1}$ und G^{-1} von $\hat{F}_n$ und G grafisch vergleicht. Hierzu

- plote man $G^{-1}(k/n)$ gegen $\hat{F}_n^{-1}(k/n) = X_{(k)}$, $k = 1, \dots, n$.
- Falls die Punktwolke

$$\left\{ \left(G^{-1}(k/n), X_{(k)} \right), \quad k = 1, \dots, n \right\}$$

näherungsweise auf einer Geraden $y = ax + b$ liegt, so sagt man, dass $F_X(x) \approx G\left(\frac{x-a}{b}\right)$, $x \in \mathbb{R}$.

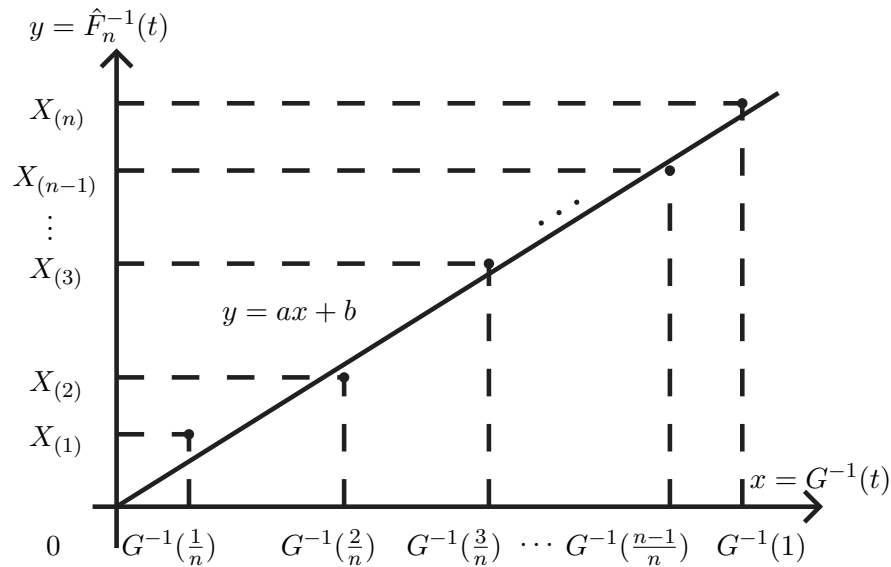


Abb. 2.7: Quantil-Grafik

Diese empirische Vergleichsmethode beruht auf folgenden Überlegungen:

- Man ersetzt die unbekannte Funktion F_X durch die aus den Daten berechenbare Funktion $\hat{F}_n$. Dabei macht man einen Fehler, der allerdings asymptotisch (für $n \rightarrow \infty$) klein ist. Dies folgt aus dem Satz 3.3.9 von Glivenko-Cantelli, der besagt, dass

$$\sup_{x \in \mathbb{R}} \left| \hat{F}_n(x) - F_X(x) \right| \xrightarrow{n \rightarrow \infty} 0.$$

Der Vergleich der entsprechenden Quantil-Funktionen wird durch folgendes Ergebnis bestärkt: Falls $\mathbb{E}X < \infty$, dann gilt

$$\sup_{t \in [0,1]} \left| \int_0^t \left(\hat{F}_n^{-1}(y) - F_X^{-1}(y) \right) dy \right| \xrightarrow[n \rightarrow \infty]{\text{f.s.}} 0.$$

Somit setzt man bei der Verwendung der Quantil-Grafiken voraus, dass der Stichprobenumfang n ausreichend groß ist, um $\hat{F}_n^{-1} \approx F_X^{-1}$ zu gewährleisten.

- Man setzt zusätzlich voraus, dass die Gleichungen

$$\begin{aligned} y &= ax + b, \\ y &= F_X^{-1}(t), \\ x &= G^{-1}(t) \end{aligned}$$

für alle t (und nicht nur näherungsweise für $t = k/n$, $k = 1, \dots, n$) gelten. Daraus folgt, dass $G(x) = t = F_X(y) = F_X(ax + b)$ für alle x , oder $F_X(y) = G\left(\frac{y-b}{a}\right)$ für alle y , weil $x = \frac{y-b}{a}$ ist.

Aus praktischer Sicht ist es besser, Paare $\left(G^{-1}\left(\frac{k}{n+1}\right), X_{(k)}\right)$, $k = 1, \dots, n$ zu plotten. Dadurch wird vermieden, dass $G^{-1}(n/n) = G^{-1}(1) = \infty$ vorkommt, wie es zum Beispiel bei einer Verteilung G der Fall ist, bei der $F(x) < 1$ gilt für alle $x \in \mathbb{R}$. Tatsächlich gilt für $k = n$, dass $\frac{n}{n+1} < 1$ und somit $G^{-1}\left(\frac{n}{n+1}\right) < \infty$.

Beispiel 2.3.1 (Exponential-Verteilung, $G(x) = (1 - e^{-\lambda x}) \cdot \mathbb{I}(x \geq 0)$):

Es gilt $G^{-1}(y) = -1/\lambda \log(1 - y)$, $y \in (0, 1)$. So wird man beim Quantil-Plot Paare

$$\left(-\frac{1}{\lambda} \log\left(1 - \frac{k}{n+1}\right), X_{(k)}\right), \quad k = 1, \dots, n$$

zeichnen, wobei der Faktor $1/\lambda$ für die Linearität unwesentlich ist und weggelassen werden kann.

Beispiel 2.3.2 (Normalverteilung, $G(x) = \Phi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-t^2/2} dt$, $x \in \mathbb{R}$):

Leider ist die analytische Berechnung von Φ^{-1} mit einer geschlossenen Formel nicht möglich. Aus diesem Grund wird $\Phi^{-1}\left(\frac{k}{n+1}\right)$ numerisch berechnet und in Tabellen oder statistischen Software-Paketen (wie z.B. R) abgelegt. Um die empirische Verteilung der Daten mit der Normalverteilung zu vergleichen, trägt man Punkte mit Koordinaten

$$\left(\Phi^{-1}\left(\frac{k}{n+1}\right), X_{(k)}\right), \quad k = 1, \dots, n$$

auf der Ebene auf und prüft, ob sie eine Gerade bilden (vgl. Abb. 2.8).

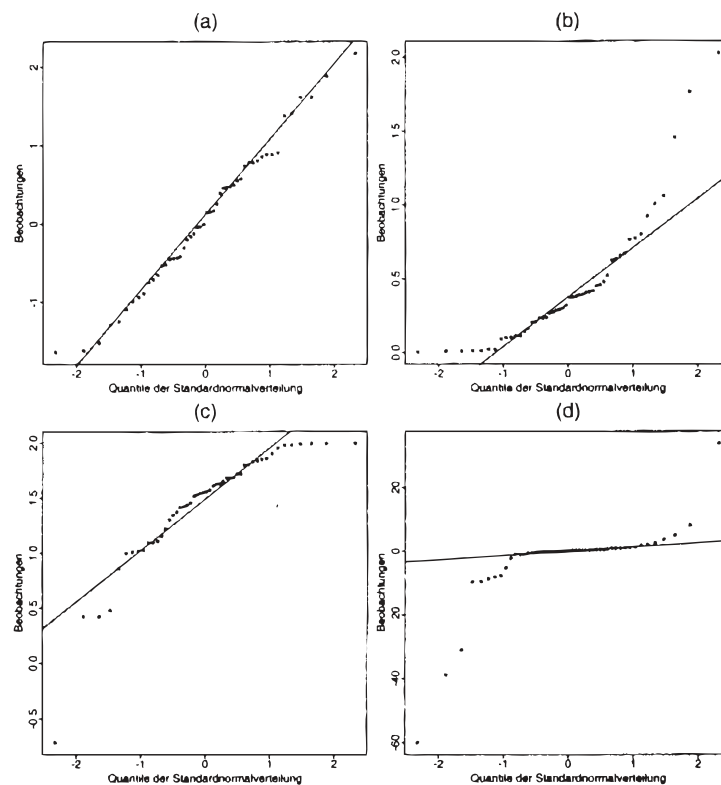


Abb. 2.8: QQ-Plot einer Normalverteilung (a), einer linkssteilen Verteilung (b), einer rechtssteilen Verteilung (c) und einer symmetrischen, aber stark gekrümmten Verteilung (d)

Übungsaufgabe 2.3.1

Entwerfen Sie die Quantil-Grafiken für den Vergleich der empirischen Verteilung mit der Lognormal und der Weibull-Verteilung.

Bemerkung 2.3.1

Falls $\bar{x}_n = 0$ und die Verteilung F_X linkssteil ist, so sind die Quantile von F_X kleiner als die von Φ . Somit ist der Normal-Quantilplot konvex. Falls $\bar{x}_n = 0$ und F_X rechtssteil ist, so wird der Normal-Quantilplot konkav sein.

Beispiel 2.3.3 (Haftpflichtversicherung (Belgien, 1992)):

In Abbildung 2.9 sind Ordnungsstatistiken der Stichprobe von $n = 227$ Schadenhöhen der Industrie-Unfälle in Belgien im Jahr 1992 (Haftpflichtversicherung) gegen Quantile von Exponential-, Pareto-, Standardnormal- und Weibull-Verteilungen geplottet. Im Bereich von Kleinschäden zeigen die Exponential- und Pareto-Verteilungen eine gute Übereinstimmung mit den Daten. Die Verteilung von mittelgroßen Schäden kann am besten durch die Lognormal- und Weibull-Verteilungen modelliert werden. Für Großschäden erweist sich die Weibull-Verteilung als geeignet.

Beispiel 2.3.4 (Rendite der BMW-Aktie):

In Abbildung 2.10 ist der Quantilplot für Renditen der BMW-Aktie beispielhaft zu sehen.

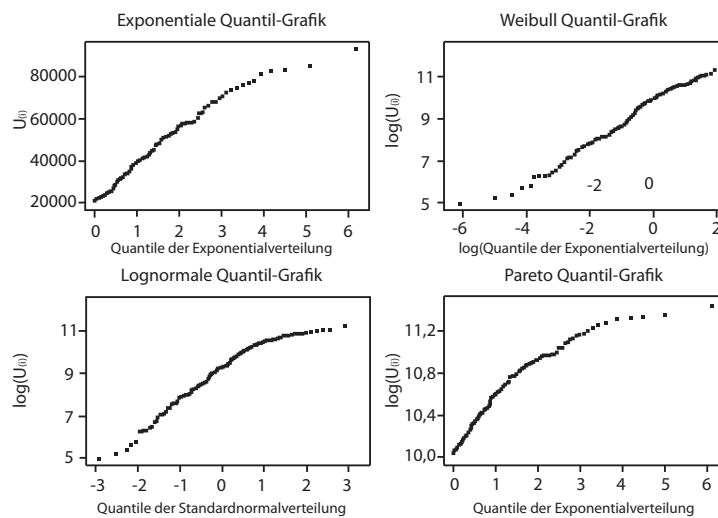


Abb. 2.9: Ordnungsstatistiken einer Stichprobe von Schadenhöhen der Industrie-Unfälle in Belgien im Jahr 1992

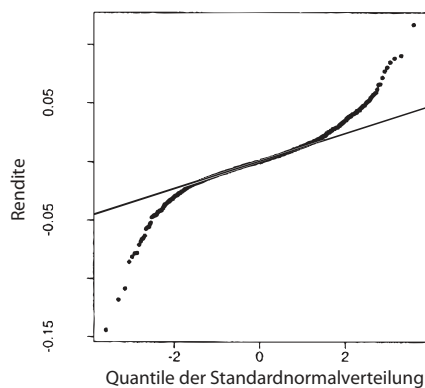


Abb. 2.10: Quantilplot der Rendite der BMW-Aktie

2.4 Dichteschätzung

Sei eine Stichprobe $(x_1, \dots, x_n)$ von unabhängigen Realisierungen eines absolut stetig verteilten Merkmals X mit Dichte f_X gegeben. Mit Hilfe der in Abschnitt 2.1.1 eingeführten Histogramme lässt sich f_X grafisch durch eine Treppenfunktion $\hat{f}_X$ darstellen. Dabei gibt es zwei entscheidende Nachteile der Histogrammdarstellung:

1. Willkür in der Wahl der Klasseneinteilung $[c_{i-1}, c_i]$,
2. Eine (möglicherweise) stetige Funktion f_X wird durch eine Treppenfunktion $\hat{f}_X$ ersetzt.

In diesem Abschnitt werden wir versuchen, diese Nachteile zu beseitigen, indem wir eine Klasse von Kerndichtenschätzern einführen, die (je nach Wahl des Kerns) auch zu stetigen Schätzern $\hat{f}_X$ führen.

Definition 2.4.1

Der Kern $K(x)$ wird definiert als eine nicht-negative messbare Funktion auf $\mathbb{R}$ mit der Eigenschaft $\int_{\mathbb{R}} K(x) dx = 1$.

Definition 2.4.2

Der *Kerndichteschätzer* der Dichte f_X aus den Daten $(x_1, \dots, x_n)$ mit Kernfunktion $K(x)$ ist gegeben durch

$$\hat{f}_X(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right), \quad x \in \mathbb{R},$$

wobei $h > 0$ die sogenannte *Bandbreite* ist.

Beispiele für Kerne:

1. *Rechteckskern:*

$$K(x) = 1/2 \cdot \mathbb{I}(x \in [-1, 1)).$$

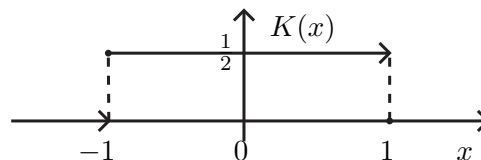
Dabei ist

$$\frac{1}{h} K\left(\frac{x - x_i}{h}\right) = \begin{cases} 1/(2h), & x_i - h \leq x < x_i + h, \\ 0, & \text{sonst,} \end{cases}$$

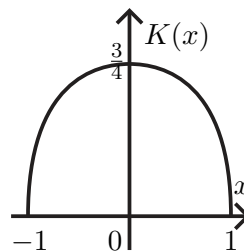
und somit

$$\hat{f}_X(x) = \frac{1}{nh} \sum_{i=1}^k K\left(\frac{x - x_i}{h}\right) = \frac{\#\{x_i \in [x - h, x + h)\}}{2nh},$$

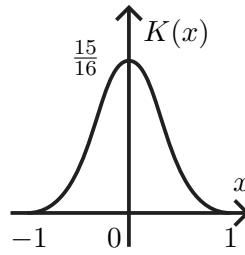
das auch *gleitendes Histogramm* genannt wird. Dieser Dichteschätzer ist (noch) nicht stetig, was durch die (besonders einfache rechteckige unstetige) Form des Kerns erklärt wird.

2. *Epanechnikov-Kern:*

$$K(x) = \begin{cases} 3/4(1 - x^2), & x \in [-1, 1) \\ 0, & \text{sonst.} \end{cases}$$

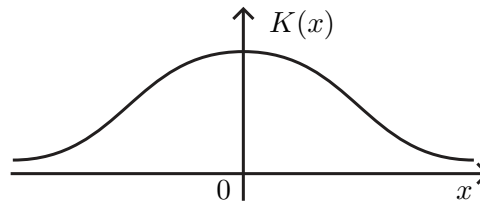
3. *Bisquare-Kern:*

$$K(x) = \frac{15}{16} \left((1 - x^2)^2 \cdot \mathbb{I}(x \in [-1, 1)) \right).$$



4. Gauss-Kern:

$$K(x) = \frac{1}{\sqrt{2\pi}} e^{-x^2/2}, \quad x \in \mathbb{R}.$$



Dabei ist die Wahl der Bandbreite h entscheidend für die Qualität der Schätzung. Je größer $h > 0$, desto glatter wird $\hat{f}_X$ sein und desto mehr „Details“ werden „herausgemittelt“. Für kleinere h wird $\hat{f}_X$ rauer. Dabei können aber auch Details auftreten, die rein stochastischer Natur sind und keine Gesetzmäßigkeiten zeigen. Mit der adäquaten Wahl von h beschäftigen sich viele wissenschaftliche Arbeiten, die empirische Faustregeln, aber auch kompliziertere Optimierungsmethoden dafür vorschlagen. Insgesamt ist das Problem der optimalen Dichteschätzung in der Statistik immer noch offen.

2.5 Beschreibung und Exploration von bivariaten Datensätzen

Im Gegensatz zu der Datenlage in den Abschnitten 2.1 bis 2.4 betrachten wir im Folgenden Datensätze bestehend aus 2 Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$, die als Realisierungen von stochastischen Stichproben $(X_1, \dots, X_n)$ und $(Y_1, \dots, Y_n)$ aufgefasst werden, wobei $X_1, \dots, X_n$ unabhängige identisch verteilte Zufallsvariablen mit $X_i \stackrel{d}{=} X \sim F_X$, $Y_1, \dots, Y_n$ unabhängige identisch verteilte Zufallsvariablen mit $Y_i \stackrel{d}{=} Y \sim F_Y$ sind. Wir betrachten hier ausschließlich quantitative Merkmale X und Y . Es wird ein Zusammenhang zwischen X und Y vermutet, der an Hand von (konkreten) Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ näher untersucht werden soll. Mit anderen Worten, wir interessieren uns für die Eigenschaften der bivariaten Verteilung $F_{X,Y}(x, y) = P(X \leq x, Y \leq y)$ des Zufallsvektors $(X, Y)^T$.

2.5.1 Grafische Darstellung von bivariaten Datensätzen

Um die Verteilung von $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ zu visualisieren, betrachten wir drei Möglichkeiten:

1. Streudiagramme
2. Zweidimensionale Histogramme

3. *Kerndichteschätzer* (im Falle eines absolut stetig verteilten Zufallsvektors $(X, Y)^T$)

1. *Streudiagramme* sind die erste sehr einfache und intuitive Visualisierungsmöglichkeit von bivariaten Daten. Um ein Streudiagramm zu erstellen, plottet man die „Punktwolke“ $(x_i, y_i)_{i=1, \dots, n}$ auf einer Koordinatenebene im $\mathbb{R}^2$. Dabei zeigt die Form der Punktwolke, ob ein linearer ($y = ax + b$) bzw. polynomialer ($y = P_d(x)$) Zusammenhang in den Daten zu erwarten ist. Später werden solche Zusammenhänge im Rahmen der Regressionstheorie untersucht (vgl. Abschnitt 2.5.3 für die einfache lineare Regression).

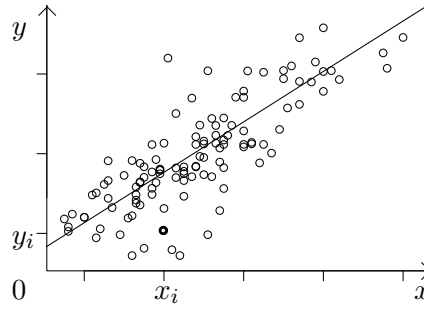


Abb. 2.11: Punktwolke

2. *Zweidimensionale Histogramme* dienen der Darstellung der bivariaten Zähldichte $p(x, y)$ des Zufallsvektors (X, Y) , falls er diskret verteilt ist, bzw. seiner Dichte $f(x, y)$ im Falle einer absolut stetigen Verteilung von (X, Y) aus den Daten $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$. Dabei teilt man den Wertebereich von X in Intervalle

$$[c_{i-1}, c_i), \quad i = 1, \dots, k, \quad -\infty = c_0 < c_1 < \dots < c_k = +\infty$$

und den Wertebereich von Y in Intervalle

$$[e_{i-1}, e_i), \quad i = 1, \dots, m, \quad -\infty = e_0 < e_1 < \dots < e_m = +\infty.$$

Bezeichnen wir

$$h_{ij} = \#\{(x_k, y_k), k = 1, \dots, n : x_k \in [c_{i-1}, c_i), y_k \in [e_{j-1}, e_j)\}$$

als die absolute Häufigkeit von (X, Y) in $[c_{i-1}, c_i) \times [e_{j-1}, e_j)$, $f_{ij} = h_{ij}/n$ als die relative Häufigkeit. Das zweidimensionale Histogramm setzt sich aus den Säulen mit Grundriss $[c_{i-1}, c_i) \times [e_{j-1}, e_j)$ und Höhe

$$\frac{h_{ij}}{(c_i - c_{i-1})(e_j - e_{j-1})}$$

für das Histogramm absoluter Häufigkeiten bzw.

$$\frac{f_{ij}}{(c_i - c_{i-1})(e_j - e_{j-1})}$$

für das Histogramm relativer Häufigkeiten zusammen, damit das Volumen dieser Säulen h_{ij} bzw. f_{ij} ist. Dabei hat solch ein Histogramm dieselben Vor- bzw. Nachteile wie ein eindimensionales, wenn es um die grafische Darstellung einer bivariaten Dichte $f(x, y)$ geht. Deshalb benutzt man oft Kerndichteschätzer, um eine glatte Darstellung zu bekommen.

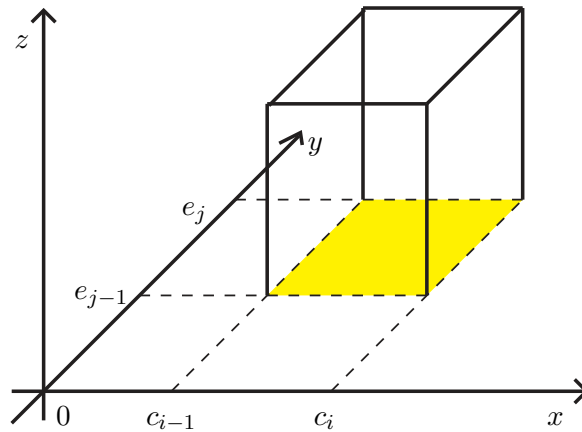


Abb. 2.12: Zweidimensionales Histogramm

3. Zweidimensionale Kerndichteschätzer haben die Form

$$\hat{f}(x, y) = \frac{1}{nh_1h_2} \sum_{i=1}^n K\left(\frac{x - x_i}{h_1}\right) K\left(\frac{y - y_i}{h_2}\right)$$

für die Bandbreiten $h_1, h_2 > 0$, die Glättungsparameter sind. Dabei ist $K(\cdot)$ eine Kernfunktion (vgl. Abschnitt 2.4). Seine Eigenschaften übertragen sich aus dem eindimensionalen Fall.

2.5.2 Zusammenhangsmaße

Jetzt wird uns die Frage beschäftigen, in welchem Maße die Merkmale X und Y voneinander abhängig sind. Um die $\text{Cov}(X, Y) = \mathbb{E}(X - \mathbb{E}X)(Y - \mathbb{E}Y)$ aus den Daten zu schätzen, setzt man die sogenannte *empirische Kovarianz*

$$S_{xy}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}_n)(y_i - \bar{y}_n)$$

ein. Dabei ist S_{xy}^2 jedoch von den Skalen von X und Y abhängig.

1. Um ein skaleninvariantes Zusammenhangsmaß zu bekommen, betrachtet man die empirische Variante des Korrelationskoeffizienten

$$\varrho(X, Y) = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var } X} \cdot \sqrt{\text{Var } Y}},$$

den sogenannten *Bravais-Pearson-Korrelationskoeffizienten*

$$\varrho_{xy} = \frac{S_{xy}^2}{\sqrt{S_{xx}^2 \cdot S_{yy}^2}},$$

wobei

$$S_{xx}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}_n)^2, \quad S_{yy}^2 = \frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y}_n)^2$$

die Stichprobenvarianzen der Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ sind. Dabei erbt ϱ_{xy} alle Eigenschaften des Korrelationskoeffizienten $\varrho(X, Y)$:

- a) $|\varrho_{xy}| \leq 1$
- b) $\varrho_{xy} = \pm 1$, falls ein linearer Zusammenhang in den Daten $(x_i, y_i)_{i=1, \dots, n}$ vorliegt, d.h. alle Punkte (x_i, y_i) , $i = 1, \dots, n$ liegen auf einer Gerade mit positivem (bei $\varrho_{xy} = 1$) bzw. negativem (bei $\varrho_{xy} = -1$) Anstieg.
- c) Wenn $|\varrho_{xy}|$ klein ist ($\varrho_{xy} \approx 0$), so sind die Datensätze unkorreliert. Dabei wird oft folgende grobe Einteilung vorgenommen:

Merkmale X und Y sind

- „*schwach korreliert*“, falls $|\varrho_{xy}| < 0.5$,
- „*stark korreliert*“, falls $|\varrho_{xy}| \geq 0.8$.

Ansonsten liegt ein mittlerer Zusammenhang zwischen X und Y vor.

Lemma 2.5.1

Für ϱ_{xy} gilt die alternative rechengünstige Darstellung

$$\varrho_{xy} = \frac{\sum_{i=1}^n x_i y_i - n \bar{x}_n \bar{y}_n}{\sqrt{(\sum_{i=1}^n x_i^2 - n \bar{x}_n^2) (\sum_{i=1}^n y_i^2 - n \bar{y}_n^2)}}. \quad (2.5.1)$$

Beweis Man muss lediglich zeigen, dass

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

Alles andere folgt daraus für $x_i = y_i$, $i = 1, \dots, n$. Es gilt

$$\begin{aligned} \sum_{i=1}^n (x_i - \bar{x}_n)(y_i - \bar{y}_n) &= \sum_{i=1}^n x_i y_i - \bar{x}_n \sum_{i=1}^n y_i - \bar{y}_n \sum_{i=1}^n x_i + n \bar{x}_n \bar{y}_n \\ &= \sum_{i=1}^n x_i y_i - n \bar{x}_n \bar{y}_n - n \bar{y}_n \bar{x}_n + n \bar{x}_n \bar{y}_n = \sum_{i=1}^n x_i y_i - n \bar{x}_n \bar{y}_n \end{aligned}$$

□

Falls die vorliegenden Daten $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ nur 2 Ausprägungen zeigen und somit *binär* kodiert werden können, d.h. $x_i, y_i \in \{0, 1\}$, dann gilt

$$\varrho_{xy} = \frac{h_{00}h_{11} - h_{01}h_{10}}{\sqrt{h_{0\cdot} \cdot h_{1\cdot} \cdot h_{\cdot 0} \cdot h_{\cdot 1}}} = \varphi$$

(der sogenannte Phi-Koeffizient), wobei

$$\begin{aligned} h_{00} &= \#\{(x_i, y_i) : x_i = y_i = 0\} \\ h_{11} &= \#\{(x_i, y_i) : x_i = y_i = 1\} \\ h_{01} &= \#\{(x_i, y_i) : x_i = 0, y_i = 1\} \\ h_{10} &= \#\{(x_i, y_i) : x_i = 1, y_i = 0\} \\ h_{0.} &= h_{00} + h_{01} \\ h_{.0} &= h_{00} + h_{10} \\ h_{1.} &= h_{10} + h_{11} \\ h_{.1} &= h_{01} + h_{11} \end{aligned}$$

Übungsaufgabe 2.5.1

Zeigen Sie diese Darstellungsform!

2. Spearmans Korrelationskoeffizient

Einen alternativen Korrelationskoeffizienten erhält man, wenn man die Stichprobenwerte x_i bzw. y_i in ϱ_{xy} durch ihre *Ränge* $\text{rg}(x_i)$ bzw. $\text{rg}(y_i)$ ersetzt, die als Position dieser Werte in den ansteigend geordneten Stichproben zu verstehen sind:

$\text{rg}(x_i) = j$, falls $x_i = x_{(j)}$ für ein $j \in \{1, \dots, n\}$, $\forall i = 1, \dots, n$. Es bedeutet, dass $\text{rg}(x_{(i)}) = i \forall i = 1, \dots, n$, falls $x_i \neq x_j$ für $i \neq j$.

Falls die Stichprobe $(x_1, \dots, x_n)$ k identische Werte x_i (die sogenannten *Bindungen*) enthält, so wird diesen Werten der sogenannte Durchschnittsrang $\text{rg}(x_i)$ zugewiesen, der als arithmetisches Mittel der k in Frage kommenden Ränge errechnet wird. Zum Beispiel findet folgende Zuordnung statt:

$$\frac{x_i}{\text{rg}(x_i)} \mid \begin{array}{l} (3, 1, 7, 5, 3, 3) \\ (a, 1, 6, 5, a, a) \end{array}$$

wobei der Durchschnittsrang a von Stichprobeneintrag 3 gleich $a = 1/3(2 + 3 + 4) = 3$ ist.

Somit wird der sogenannte *Spearmans Korrelationskoeffizient* (Rangkorrelationskoeffizient) der Stichproben

$$(x_1, \dots, x_n) \quad \text{und} \quad (y_1, \dots, y_n)$$

als der *Bravais-Pearson-Koeffizient* der Stichproben ihrer Ränge

$$(\text{rg}(x_1), \dots, \text{rg}(x_n)) \quad \text{und} \quad (\text{rg}(y_1), \dots, \text{rg}(y_n))$$

definiert:

$$\varrho_{sp} = \frac{\sum_{i=1}^n (\text{rg}(x_i) - \overline{\text{rg}_x})(\text{rg}(y_i) - \overline{\text{rg}_y})}{\sqrt{\sum_{i=1}^n (\text{rg}(x_i) - \overline{\text{rg}_x})^2 \sum_{i=1}^n (\text{rg}(y_i) - \overline{\text{rg}_y})^2}},$$

wobei

$$\begin{aligned} \overline{\text{rg}_x} &= \frac{1}{n} \sum_{i=1}^n \text{rg}(x_i) = \frac{1}{n} \sum_{i=1}^n \text{rg}(x_{(i)}) = \frac{1}{n} \sum_{i=1}^n i = \frac{n(n+1)}{2n} = \frac{n+1}{2}, \\ \overline{\text{rg}_y} &= \frac{1}{n} \sum_{i=1}^n \text{rg}(y_i) = \frac{n+1}{2}. \end{aligned}$$

Dieselbe Darstellung $\overline{\text{rg}}_y$ gilt auch, wenn Bindungen vorhanden sind.

Dieser Koeffizient misst monotone Zusammenhänge in den Daten. Aus den Eigenschaften der Bravais-Pearson-Koeffizienten folgt $|\varrho_{sp}| \leq 1$. Betrachten wir die Fälle $\varrho_{sp} = \pm 1$ gesondert:

- $\varrho_{sp} = 1$ bedeutet, dass die Punkte $(\text{rg}(x_i), \text{rg}(y_i))$, $i = 1, \dots, n$ auf einer Geraden mit positiver Steigung liegen. Da aber $\text{rg}(x_i), \text{rg}(y_i) \in \mathbb{N}$, kann diese Steigung nur 1 sein. Es bedeutet, dass dem kleinsten Wert in der Stichprobe $(x_1, \dots, x_n)$ der kleinste Wert in $(y_1, \dots, y_n)$ entspricht, usw., d.h., für wachsende x_i wachsen auch die y_i streng monoton: $x_i < x_j \implies y_i < y_j \quad \forall i \neq j$.
- Analog gilt dann für $\varrho_{sp} = -1$, dass $x_i < x_j \implies y_i > y_j \quad \forall i \neq j$.

Dies kann folgendermaßen zusammengefaßt werden:

- $\varrho_{sp} > 0$: gleichsinniger monotoner Zusammenhang (x_i groß $\iff y_i$ groß)
- $\varrho_{sp} < 0$: gegensinniger monotoner Zusammenhang (x_i groß $\iff y_i$ klein)
- $\varrho_{sp} \approx 0$: kein monotoner Zusammenhang.

Da der Spearmans Korrelationskoeffizient nur Ränge von x_i und y_i betrachtet, eignet er sich auch für ordinale (und nicht nur quantitative) Daten.

Lemma 2.5.2

Falls die Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ keine Bindung enthalten ($x_i \neq x_j, y_i \neq y_j \quad \forall i \neq j$), dann gilt

$$\varrho_{sp} = 1 - \frac{6}{(n^2 - 1)n} \sum_{i=1}^n d_i^2,$$

wobei $d_i = \text{rg}(x_i) - \text{rg}(y_i) \quad \forall i = 1, \dots, n$.

Beweis Als Übungsaufgabe. □

Satz 2.5.1 (Invarianzeigenschaften):

1. Wenn die Merkmale X und Y linear transformiert werden:

$$\begin{aligned} f(X) &= a_x X + b_x, & a_x &\neq 0, b_x \in \mathbb{R}, \\ g(Y) &= a_y Y + b_y, & a_y &\neq 0, b_y \in \mathbb{R}, \end{aligned}$$

dann gilt $\varrho_{f(x)g(y)} = \text{sgn}(a_x a_y) \cdot \varrho_{xy}$.

2. Falls Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}$ und $g: \mathbb{R} \rightarrow \mathbb{R}$ beide monoton wachsend oder beide monoton fallend sind, dann gilt

$$\varrho_{sp}(f(x), g(y)) = \varrho_{sp}(x, y).$$

Falls f monoton wachsend und g monoton fallend (oder umgekehrt) sind, dann gilt $\varrho_{sp}(f(x), g(y)) = -\varrho_{sp}(x, y)$.

Beweis Beweisen wir nur 1), weil 2) offensichtlich ist.

1.

$$\begin{aligned}\varrho_{f(x)g(y)} &= \frac{\sum_{i=1}^n ((a_x x_i + b_x) - (a_x \bar{x}_n + b_x))((a_y y_i + b_y) - (a_y \bar{y}_n + b_y))}{\sqrt{a_x^2 \sum_{i=1}^n (x_i - \bar{x}_n)^2 a_y^2 \sum_{i=1}^n (y_i - \bar{y}_n)^2}} \\ &= \frac{a_x a_y}{|a_x| |a_y|} \cdot \frac{\sum_{i=1}^n (x_i - \bar{x}_n)(y_i - \bar{y}_n)}{\sqrt{\sum_{i=1}^n (x_i - \bar{x}_n)^2 \sum_{i=1}^n (y_i - \bar{y}_n)^2}} = \operatorname{sgn}(a_x a_y) \cdot \varrho_{xy}.\end{aligned}$$

□

Bemerkung 2.5.1

1. Da lineare Transformationen monoton sind, gilt Aussage 1) auch für Spearmans Korrelationskoeffizienten ϱ_{sp} .
2. Der Koeffizient ϱ_{xy} erfasst lineare Zusammenhänge, während ϱ_{sp} monotone Zusammenhänge aufspürt.

2.5.3 Einfache lineare Regression

Wenn man den Zusammenhang von Merkmalen X und Y mit Hilfe von Streudiagrammen visualisiert, wird oft ein linearer Trend erkennbar, obwohl der Bravais-Pearson-Korrelationskoeffizient einen Wert kleiner als 1 liefert, z.B. $\varrho_{xy} \approx 0,6$ (vgl. Abb. 2.13). Dies ist der Fall, weil die Da-

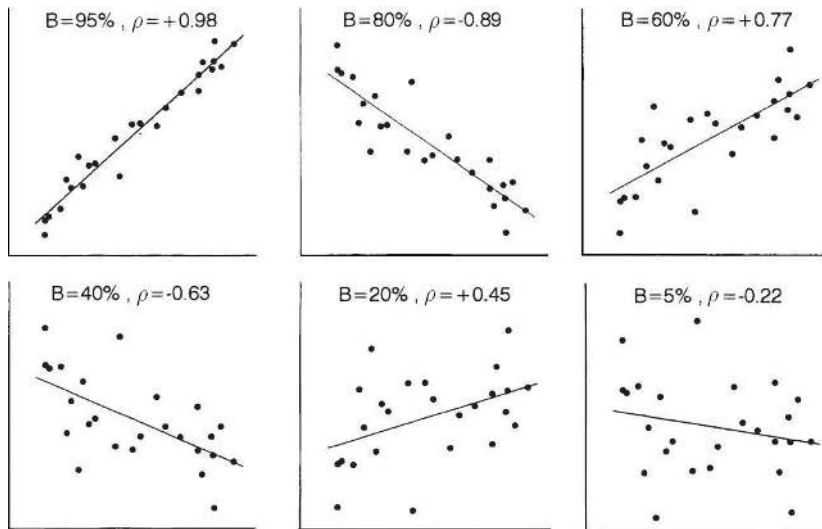


Abb. 2.13: Vergleich verschiedenwertiger Bestimmtheitsmaße. Es sind Regressionsgerade, Bestimmtheitsmaß B und Korrelationskoeffizient ρ verschiedener (fiktiver) Punktwolken vom Umfang $n = 25$ dargestellt. Die Beschriftung der Achsen ist weggelassen, weil sie hier ohne Bedeutung ist.

tenpunkte (x_i, y_i) , $i = 1, \dots, n$ oft um eine Gerade streuen und nicht exakt auf einer Geraden

X	Y
Geschwindigkeit	Länge des Bremswegs
Körpergröße des Vaters	Körpergröße des Sohnes
Produktionsfaktor	Qualität des Produktes
Spraydosen-Verbrauch	Ozongehalt der Atmosphäre
Noten im Bachelor-Studium	Noten im Master-Studium

Tab. 2.1: Beispiele möglicher Ausgangs- und Zielgrößen

liegen. Um solche Situationen stochastisch modellieren zu können, nimmt man den Zusammenhang der Form

$$Y = f(X) + \varepsilon$$

an, wobei ε die sogenannte Störgröße ist, die auf mehrere Ursachen wie z.B. Beobachtungsfehler (Messfehler, Berechnungsfehler, usw.) zurückzuführen sein kann. Dabei nennt man die Zufallsvariable Y *Zielgröße* oder *Regressand*, die Zufallsvariable X *Einflussfaktor*, *Regressor* oder *Ausgangsvariable*. Der Zusammenhang $Y = f(X) + \varepsilon$ wird *Regression* genannt, wobei man oft über ε voraussetzt, dass $\mathbb{E}\varepsilon = 0$ (kein systematischer Beobachtungsfehler). Wenn $f(x) = \alpha + \beta x$ eine lineare Funktion ist, so spricht man von der *einfachen linearen Regression*. Es sind aber durchaus andere Arten der Zusammenhänge denkbar, wie z.B.

$$f(x) = \sum_{i=0}^n \alpha_i x^i$$

(*polynomiale Regression*), usw. Beispiele für mögliche Ausgangs- bzw. Zielgrößen sind in Tabelle 2.1 zusammengefasst, einige Beispiele in Abbildung 2.14.

Auf Modellebene ist damit folgende Fragestellung gegeben: Es gebe Zufallsstichproben von Ziel- bzw. Ausgangsvariablen $(Y_1, \dots, Y_n)$ und $(X_1, \dots, X_n)$, zwischen denen ein verrauschter linearer Zusammenhang $Y_i = \alpha + \beta X_i + \varepsilon_i$ besteht, wobei ε_i Störgrößen sind, die nicht direkt beobachtbar und uns somit unbekannt sind. Meistens nimmt man an, dass $\mathbb{E}\varepsilon_i = 0 \quad \forall i = 1, \dots, n$ und $\text{Cov}(\varepsilon_i, \varepsilon_j) = \sigma^2 \delta_{ij}$, d.h. $\varepsilon_1 \dots \varepsilon_n$ sind unkorreliert mit $\text{Var} \varepsilon_i = \sigma^2$. Wenn wir über die Eigenschaften der Schätzer für α , β und σ^2 reden, gehen wir davon aus, dass die X -Werte nicht zufällig sind, also $X_i = x_i \quad \forall i = 1, \dots, n$. Wenn man von einer konkreten Stichprobe $(y_1, \dots, y_n)$ für $(Y_1, \dots, Y_n)$ ausgeht, so sollen anhand von den Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ Regressionsparameter α (*Regressionskonstante*) und β (*Regressionskoeffizient*) sowie *Regressionsvarianz* σ^2 geschätzt werden. Dabei verwendet man die sogenannte *Methode der kleinsten Quadrate*, die den mittleren quadratischen Fehler von den Datenpunkten $(x_i, y_i)_{i=1, \dots, n}$ des Streudiagramms zur *Regressionsgeraden* $y = \alpha + \beta x$ minimiert:

$$(\alpha, \beta) = \arg \min_{\alpha, \beta \in \mathbb{R}} e(\alpha, \beta) \quad \text{mit} \quad e(\alpha, \beta) = \frac{1}{n} \sum_{i=1}^n (y_i - \alpha - \beta x_i)^2.$$

Diese Methode wurde 1809 von C.F. Gauß in seinem Werk „Theoria motus corporum coelestium“ verwendet, um die Laufbahnen der Himmelskörper an Hand von Beobachtungen zu bestimmen. Die Bezeichnung „Methode der kleinsten Quadrate“ stammt allerdings vom französischen Mathematiker A.M. Legendre (1752-1832), der sie unabhängig von Gauß entdeckt hat.

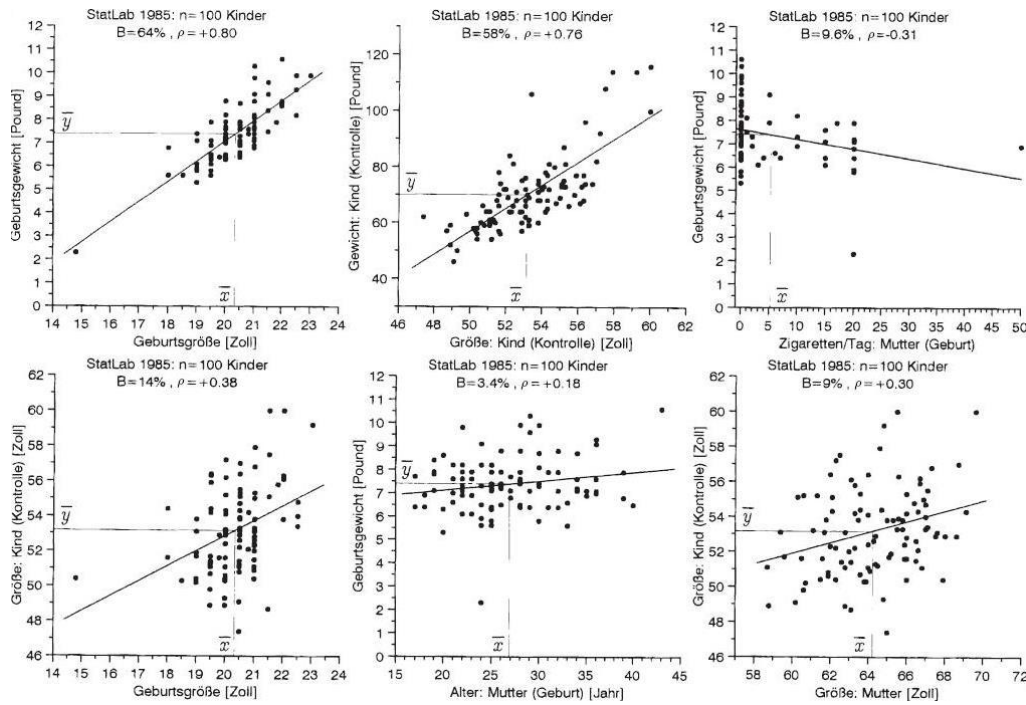


Abb. 2.14: Punktwolken verschiedener Merkmale der StatLab-Auswahl 1985 mit Regressionsgerade, Bestimmtheitsmaß B und Korrelationskoeffizient ρ .

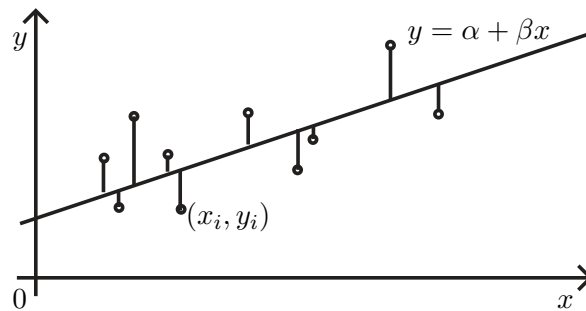


Abb. 2.15: Methode kleinster Quadrate

Da die Darstellung $y_i = \alpha + \beta x_i + \varepsilon_i$ gilt, kann man $e(\alpha, \beta) = 1/n \sum_{i=1}^n \varepsilon_i^2$ schreiben. Es ist der vertikale mittlere quadratische Abstand von den Datenpunkten (x_i, y_i) zur Geraden $y = \alpha + \beta x$ (vgl. Abb. 2.15). Das Minimierungsproblem $e(\alpha, \beta) \mapsto \min$ löst man durch das zweifache Differenzieren von $e(\alpha, \beta)$. Somit erhält man $\hat{\alpha} = \bar{y}_n - \hat{\beta} \bar{x}_n$, wobei

$$\hat{\beta} = \frac{S_{xy}^2}{S_{xx}^2}, \quad \bar{x}_n = \frac{1}{n} \sum_{i=1}^n x_i, \quad \bar{y}_n = \frac{1}{n} \sum_{i=1}^n y_i,$$

$$S_{xy}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}_n)(y_i - \bar{y}_n), \quad S_{xx}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}_n)^2.$$

Kind i	1	2	3	4	5	6	7	8	9
Fernsehzeit x_i	0,3	2,2	0,5	0,7	1,0	1,8	3,0	0,2	2,3
Tiefschlafdauer y_i	5,8	4,4	6,5	5,8	5,6	5,0	4,8	6,0	6,1

Tab. 2.2: Daten von Fernsehzeit und korrespondierender Tiefschlafdauer

Übungsaufgabe 2.5.2

Leiten Sie die Schätzer $\hat{\alpha}$ und $\hat{\beta}$ selbstständig her.

Die Varianz σ^2 schätzt man durch $\hat{\sigma}^2 = \frac{1}{n-2} \sum_{i=1}^n \hat{\varepsilon}_i^2$, wobei $\hat{\varepsilon}_i = y_i - \hat{\alpha} - \hat{\beta}x_i$, $i = 1, \dots, n$ die sogenannten *Residuen* sind. Die Gründe, warum $\hat{\sigma}^2$ diese Gestalt hat, können an dieser Stelle noch nicht angegeben werden, weil wir noch nicht die Maximum-Likelihood-Methode kennen. Zu gegebener Zeit (in der Vorlesung Stochastik III) wird jedoch klar, dass diese Art der Schätzung sehr natürlich ist.

Bemerkung 2.5.2

Die angegebenen Schätzer für α und β sind nicht symmetrisch bzgl. Variablen x_i und y_i . Wenn man also die *horizontalen* Abstände (statt vertikaler) zur Bildung des mittleren quadratischen Fehlers nimmt (was dem Rollentausch $x \leftrightarrow y$ entspricht), so bekommt man andere Schätzer für α und β , die mit $\hat{\alpha}$ und $\hat{\beta}$ nicht übereinstimmen müssen:

$$d_i = y_i - \alpha - \beta x_i \mapsto d'_i = x_i - \frac{(y_i - \alpha)}{\beta}.$$

Ein Ausweg aus dieser asymmetrischen Situation wäre es, die orthogonalen Abstände o_i von

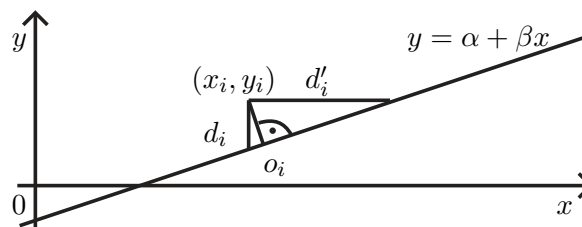


Abb. 2.16: Orthogonale Abstände

(x_i, y_i) zur Geraden $y = \alpha + \beta x$ zu betrachten (vgl. Abb. 2.16). Diese Art der Regression, die „errors-in-variables regression“ genannt wird, hat aber eine Reihe von Eigenschaften, die sie zur Prognose von Zielvariablen y_i durch die Ausgangsvariablen x_i unbrauchbar machen. Sie sollte zum Beispiel nur dann verwendet werden, wenn die Standardabweichungen für X und Y etwa gleich groß sind.

Beispiel 2.5.1

Ein Kinderpsychologe vermutet, dass sich häufiges Fernsehen negativ auf das Schlafverhalten von Kindern auswirkt. Um diese Hypothese zu überprüfen, wurden 9 Kinder im gleichen Alter befragt, wie lange sie pro Tag fernsehen dürfen, und zusätzlich die Dauer ihrer Tiefschlafphase gemessen. So ergibt sich der Datensatz in Tabelle 2.2 und die Regressionsgerade aus Abbildung 2.17.

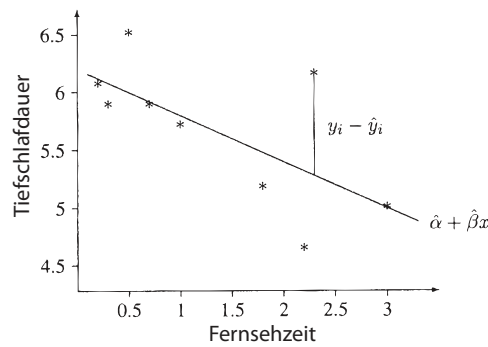


Abb. 2.17: Streudiagramm und Ausgleichsgerade zur Regression der Dauer des Tiefschlafs auf die Fernsehzeit

Es ergibt sich für die oben genannten Stichproben $(x_1, \dots, x_9)$ und $(y_1, \dots, y_9)$

$$\bar{x}_9 = 1,33, \quad \bar{y}_9 = 5,56, \quad \hat{\beta} = -0,45, \quad \hat{\alpha} = 6,16.$$

Somit ist

$$y = 6,16 - 0,45x$$

die Regressionsgerade, die eine negative Steigung hat, was die Vermutung des Kinderpsychologen bestätigt. Außerdem ist es mit Hilfe dieser Geraden möglich, Prognosen für die Dauer des Tiefschlafs für vorgegebene Fernsehzeiten anzugeben. So wäre z.B. für die Fernsehzeit von 1 Stunde der Tiefschlaf von $6,16 - 0,45 \cdot 1 = 5,71$ Stunden plausibel.

Bemerkung 2.5.3 (Eigenschaften der Regressionsgerade):

1. Es gilt $\text{sgn}(\hat{\beta}) = \text{sgn}(\rho_{xy})$, was aus $\hat{\beta} = s_{xy}^2 / s_{xx}^2$ folgt. Dies bedeutet (falls $s_{yy}^2 > 0$):
 - a) Die Regressionsgerade $y = \hat{\alpha} + \hat{\beta}x$ steigt an, falls die Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ positiv korreliert sind.
 - b) Die Regressionsgerade fällt ab, falls sie negativ korreliert sind.
 - c) Die Regressionsgerade ist konstant, falls die Stichproben unkorreliert sind.

Falls $s_{yy}^2 = 0$, dann ist die Regressionsgerade konstant ($y = \bar{y}_n$).

2. Die Regressionsgerade $y = \hat{\alpha} + \hat{\beta}x$ verläuft immer durch den Punkt $(\bar{x}_n, \bar{y}_n)$: $\hat{\alpha} + \hat{\beta}\bar{x}_n = \bar{y}_n$.
3. Seien $\hat{y}_i = \hat{\alpha} + \hat{\beta}x_i$, $i = 1, \dots, n$. Dann gilt

$$\overline{\hat{y}_n} = \frac{1}{n} \sum_{i=1}^n \hat{y}_i = \bar{y}_n \quad \text{und somit} \quad \sum_{i=1}^n \underbrace{(y_i - \hat{y}_i)}_{\hat{\varepsilon}_i} = 0.$$

Dabei sind $\hat{\varepsilon}_i$ die schon vorher eingeführten Residuen. Mit ihrer Hilfe ist es möglich, die Güte der Regressionsprognose zu beurteilen.

Residualanalyse und Bestimmtheitsmaß

Definition 2.5.1

Der relative Anteil der Streuungsreduktion an der Gesamtstreuung S_{yy}^2 heißt das *Bestimmtheitsmaß* der Regressionsgeraden:

$$R^2 = \frac{S_{yy}^2 - \frac{1}{n-1} \sum_{i=1}^n \hat{\varepsilon}_i^2}{S_{yy}^2} = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y}_n)^2}.$$

Es ist nur im Fall $S_{xx}^2 > 0$, $S_{yy}^2 > 0$ definiert, d.h., wenn nicht alle Werte x_i bzw. y_i übereinstimmen.

Warum R^2 in dieser Form eingeführt wird, zeigt folgende Überlegung, die *Streuungszerlegung* genannt wird:

Lemma 2.5.3

Die Gesamtstreuung („sum of squares total“) $\text{SQT} = (n-1)S_{yy}^2 = \sum_{i=1}^n (y_i - \bar{y}_n)^2$ lässt sich in die Summe der sogenannten erklärten Streuung „sum of squares explained“ $\text{SQE} = \sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2$ und der Residualstreuung „sum of squared residuals“ $\text{SQR} = \sum_{i=1}^n \hat{\varepsilon}_i^2 = \sum_{i=1}^n (y_i - \hat{y}_i)^2$ zerlegen:

$$\text{SQT} = \text{SQE} + \text{SQR}$$

bzw.

$$\sum_{i=1}^n (y_i - \bar{y}_n)^2 = \sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2 + \sum_{i=1}^n (y_i - \hat{y}_i)^2.$$

Beweis

$$\begin{aligned} \text{SQT} &= \sum_{i=1}^n (y_i - \bar{y}_n)^2 = \sum_{i=1}^n (y_i - \hat{y}_i + \hat{y}_i - \bar{y}_n)^2 \\ &= \underbrace{\sum_{i=1}^n (y_i - \hat{y}_i)^2}_{=\text{SQR}} + 2 \sum_{i=1}^n (y_i - \hat{y}_i)(\hat{y}_i - \bar{y}_n) + \underbrace{\sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2}_{=\text{SQE}} \\ &= \text{SQE} + \text{SQR} + 2 \sum_{i=1}^n \hat{y}_i (y_i - \hat{y}_i) - 2\bar{y}_n \underbrace{\sum_{i=1}^n (y_i - \hat{y}_i)}_{=0, \text{vgl. Eig. 3 S. 37}} = \text{SQE} + \text{SQR} + E, \end{aligned}$$

wobei noch zu zeigen ist, dass $E = 2 \sum_{i=1}^n \hat{y}_i (y_i - \hat{y}_i) = 0$, also

$$\begin{aligned} E &= 2 \sum_{i=1}^n (\hat{\alpha} + \hat{\beta}x_i)(y_i - \hat{\alpha} - \hat{\beta}x_i) = 2\hat{\alpha} \underbrace{\sum_{i=1}^n \hat{\varepsilon}_i}_{=0} + 2\hat{\beta} \sum_{i=1}^n x_i (y_i - \hat{\alpha} - \hat{\beta}x_i) \\ &= 2\hat{\beta} \left(\sum_{i=1}^n x_i y_i - \hat{\alpha} \sum_{i=1}^n x_i - \hat{\beta} \sum_{i=1}^n x_i^2 \right) \stackrel{\hat{\alpha} = \bar{y}_n - \bar{x}_n \hat{\beta}}{=} 2\hat{\beta} \left(\underbrace{\sum_{i=1}^n x_i y_i - n\bar{x}_n \bar{y}_n + \hat{\beta} n \bar{x}_n^2}_{=(n-1)S_{xy}^2} - \hat{\beta} \sum_{i=1}^n x_i^2 \right) \\ &= 2\hat{\beta} \left((n-1)S_{xy}^2 - \hat{\beta}(n-1)S_{xx}^2 \right) \stackrel{\hat{\beta} = \frac{S_{xy}^2}{S_{xx}^2}}{=} 2\hat{\beta}(n-1) \left(S_{xy}^2 - \frac{S_{xy}^2}{S_{xx}^2} \cdot S_{xx}^2 \right) = 0. \end{aligned}$$

□

Die erklärte Streuung gibt die Streuung der Regressionsgeradenwerte um $\bar{y}_n$ an. Sie stellt damit die auf den linearen Zusammenhang zwischen X und Y zurückzuführende Variation der y -Werte dar. Das oben eingeführte Bestimmtheitsmaß ist somit der Anteil dieser Streuung an der Gesamtstreuung:

$$R^2 = \frac{\text{SQE}}{\text{SQT}} = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2}{\sum_{i=1}^n (y_i - \bar{y}_n)^2} = \frac{\text{SQT} - \text{SQR}}{\text{SQT}} = 1 - \frac{\text{SQR}}{\text{SQT}}.$$

Es folgt aus dieser Darstellung, dass $R^2 \in [0, 1]$ ist.

1. $R^2 = 0$ bedeutet $\text{SQE} = \sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2 = 0$ und somit $\hat{y}_i = \bar{y}_n \ \forall i$. Dies weist darauf hin, dass das lineare Modell in diesem Fall schlecht ist, denn aus $\hat{y}_i = \hat{\alpha} + \hat{\beta}x_i = \bar{y}_n$ folgt $\hat{\beta} = \frac{S_{xy}^2}{S_{xx}^2} = 0$ und somit $S_{xy}^2 = 0$. Also sind die Merkmale X und Y unkorreliert.
2. $R^2 = 1$ bedingt $\text{SQR} = \sum_{i=1}^n \hat{\varepsilon}_i^2 = 0$. Somit liegen alle (x_i, y_i) perfekt auf der Regressionsgeraden. Dies bedeutet, dass die Daten x_i und y_i , $i = 1, \dots, n$ perfekt linear abhängig sind.

Faustregel zur Beurteilung der Güte der Anpassung eines linearen Modells an Hand von Bestimmtheitsmaß R^2 :

R^2 ist deutlich von Null verschieden (d.h. es besteht noch ein linearer Zusammenhang), falls $R^2 > \frac{4}{n+2}$, wobei n der Stichprobenumfang ist.

Allgemein gilt folgender Zusammenhang zwischen dem Bestimmtheitsmaß R^2 und dem Bravais-Pearson-Korrelationskoeffizienten ϱ_{xy} :

Lemma 2.5.4

$$R^2 = \varrho_{xy}^2$$

Beweis Aus der Eigenschaft 3 S. 37 folgt $\bar{y}_n = \overline{\hat{y}_n}$. Somit gilt

$$\text{SQE} = \sum_{i=1}^n (\hat{y}_i - \bar{y}_n)^2 = \sum_{i=1}^n (\hat{y}_i - \overline{\hat{y}_n})^2 = \sum_{i=1}^n (\hat{\alpha} + \hat{\beta}x_i - \hat{\alpha} - \hat{\beta}\bar{x}_n)^2 = \hat{\beta}^2 \sum_{i=1}^n (x_i - \bar{x}_n)^2$$

und damit

$$R^2 = \frac{\text{SQE}}{\text{SQT}} = \frac{\hat{\beta}^2 \sum_{i=1}^n (x_i - \bar{x}_n)^2}{\sum_{i=1}^n (y_i - \bar{y}_n)^2} = \frac{(S_{xy}^2)^2}{(S_{xx}^2)^2} \cdot \frac{(n-1)S_{xx}^2}{(n-1)S_{yy}^2} = \left(\frac{S_{xy}^2}{S_{yy}S_{xx}} \right)^2 = \varrho_{xy}^2$$

□

Folgerung 2.5.1

1. Der Wert von R^2 ändert sich bei einer Lineartransformation der Daten $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$ nicht. Grafisch kann man die Güte der Modellanpassung bei der linearen Regression folgendermaßen überprüfen:

Man zeichnet Punktepaaire $(\hat{y}_i, \hat{\varepsilon}_i)_{i=1, \dots, n}$ als Streudiagramm (der sogenannte *Residual-plot*). Falls diese Punktwolke gleichmäßig um Null streut, so ist das lineare Modell gut gewählt worden. Falls das Streudiagramm einen erkennbaren Trend aufweist, bedeutet das, dass die Annahme des linearen Modells für diese Daten ungeeignet sei (vgl. Abb. 2.18)

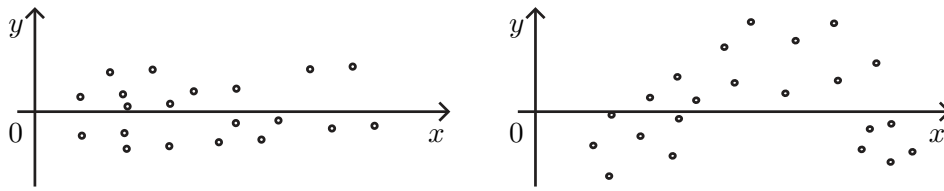


Abb. 2.18: Links: Gute, Rechts: Schlechte Übereinstimmung mit dem linearen Modell

2. Da $R^2 = \varrho_{xy}^2$, ist der Wert von R^2 symmetrisch bzgl. der Stichproben $(x_1, \dots, x_n)$ und $(y_1, \dots, y_n)$:

$$\varrho_{xy}^2 = R^2 = \varrho_{yx}^2 \quad \text{bzw.} \quad R_{xy}^2 = R_{yx}^2,$$

wobei R_{xy}^2 das Bestimmtheitsmaß bezeichnet, das sich aus der normalen Regression ergibt und R_{yx}^2 das mit vertauschten Achsen.

3 Punktschätzer

3.1 Parametrisches Modell

Sei $(x_1, \dots, x_n)$ eine konkrete Stichprobe. Es wird angenommen, dass $(x_1, \dots, x_n)$ eine Realisierung einer Zufallsstichprobe $(X_1, \dots, X_n)$ ist, wobei $X_1, \dots, X_n$ unabhängige identisch verteilte Zufallsvariablen mit der unbekannten Verteilungsfunktion F sind und F zu einer bekannten parametrischen Familie $\{F_\theta : \theta \in \Theta\}$ gehört. Hier ist $\theta = (\theta_1, \dots, \theta_m) \in \Theta$ der *m-dimensionale Parametervektor* der Verteilung F_θ und $\Theta \subset \mathbb{R}^m$ der sogenannte *Parameterraum* (eine Borel-Teilmenge von $\mathbb{R}^m$, die die Menge aller zugelassenen Parameterwerte darstellt). Es wird vorausgesetzt, dass die Parametrisierung $\theta \rightarrow F_\theta$ *identifizierbar* ist, indem $F_{\theta_1} \neq F_{\theta_2}$ für $\theta_1 \neq \theta_2$ gilt.

Eine wichtige Aufgabe der Statistik, die wir in diesem Kapitel betrachten werden, besteht in der Schätzung des Parametervektors θ (oder eines Teils von θ) an Hand von der konkreten Stichprobe $(x_1, \dots, x_n)$. In diesem Fall spricht man von einem *Punktschätzer* $\hat{\theta} : \mathbb{R}^n \rightarrow \mathbb{R}^m$, der eine gültige Stichprobenfunktion ist. Meistens wird angenommen, dass

$$\mathbb{P}(\hat{\theta}(X_1, \dots, X_n) \in \Theta) = 1,$$

wobei es zu dieser Regel auch Ausnahmen gibt. Bisher haben wir den Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P})$, auf dem unsere Zufallsstichprobe definiert ist, nicht näher spezifiziert. Dies kann man aber leicht tun, indem man den sogenannten *kanonischen Wahrscheinlichkeitsraum* angibt, wobei

$$\Omega = \mathbb{R}^\infty, \quad \mathcal{F} = \mathcal{B}_\mathbb{R}^\infty = \mathcal{B}_\mathbb{R} \times \mathcal{B}_\mathbb{R} \times \dots$$

und das Wahrscheinlichkeitsmaß $\mathbb{P}$ durch

$$\mathbb{P}(\{\omega = (\omega_1, \dots, \omega_n, \dots) \in \mathbb{R}^\infty : \omega_{i_1} \leq x_{i_1}, \dots, \omega_{i_k} \leq x_{i_k}\}) = F_\theta(x_{i_1}) \dots F_\theta(x_{i_k})$$

$\forall k \in \mathbb{N}, \quad 1 \leq i_1 < \dots < i_k$ gegeben sei. Um zu betonen, dass $\mathbb{P}$ vom Parameter θ abhängt, werden wir Bezeichnungen $\mathbb{P}_\theta, \mathbb{E}_\theta$ und Var_θ für das Maß $\mathbb{P}$, den Erwartungswert und die Varianz bzgl. $\mathbb{P}$ verwenden.

Auf dem kanonischen Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P}_\theta)$ gilt $X_i(\omega) = \omega_i$ (Projektion auf die Koordinate i), $i = 1, \dots, n$,

$$\mathbb{P}_\theta(X_i \leq x_i) = \mathbb{P}_\theta(\{\omega \in \Omega : \omega_i \leq x_i\}) = F_\theta(x_i), \quad i = 1, \dots, n, \quad x_i \in \mathbb{R}.$$

Beispiel 3.1.1

1. Sei X die Dauer des fehlerfreien Arbeitszyklus eines technischen Systems. Oft wird $X \sim \text{Exp}(\lambda)$ angenommen. Dann stellt $\{F_\theta : \theta \in \Theta\}$ mit $m = 1, \theta = \lambda, \Theta = \mathbb{R}_+$ und

$$F_\theta(x) = (1 - e^{-\theta x}) \cdot \mathbb{I}(x \geq 0)$$

ein parametrisches Modell dar, wobei der Parameterraum eindimensional ist. Später wird für λ der (Punkt-) Schätzer $\hat{\lambda}(x_1, \dots, x_n) = 1/\bar{x}_n$ vorgeschlagen.

2. In den Fragestellungen der statistischen Qualitätskontrolle werden n Erzeugnisse auf Mängel untersucht. Falls $p \in (0, 1)$ die unbekannte Wahrscheinlichkeit des Mangels ist, so wird mit $X \sim \text{Bin}(n, p)$ die Gesamtanzahl der mangelhaften Produkte beschrieben. Dabei wird folgendes parametrische Modell unterstellt:

$$\Theta = \{(n, p) : n \in \mathbb{N}, p \in (0, 1)\}, \quad \theta = (n, p), \quad m = 2,$$

$$F_\theta(x) = \mathbb{P}_\theta(X \leq x) = \begin{cases} 1, & x > n \\ \sum_{k=0}^{[x]} \binom{n}{k} p^k (1-p)^{n-k}, & x \in [0, n] \\ 0, & x < 0. \end{cases}$$

Falls n bekannt ist, kann die Wahrscheinlichkeit p des Ausschusses durch den Punktschätzer $\hat{p}(x_1, \dots, x_n) = \bar{x}_n$, $x_i \in \{0, 1\}$ näherungsweise berechnet werden.

3.2 Parametrische Familien von statistischen Prüfverteilungen

In der Vorlesung Wahrscheinlichkeitsrechnung wurden bereits einige parametrische Familien von Verteilungen eingeführt. Hier geben wir weitere Verteilungsfamilien an, die in der Statistik eine besondere Stellung einnehmen, weil sie als Referenzverteilungen in der Schätztheorie, statistischen Tests und Vertrauensintervallen ihre Anwendung finden.

3.2.1 Gamma-Verteilung

Als erstes führen wir zwei spezielle Funktionen aus der Analysis ein:

1. Die *Gamma-Funktion*:

$$\Gamma(p) = \int_0^\infty x^{p-1} e^{-x} dx \quad \text{für } p > 0.$$

Es gelten folgende Eigenschaften:

$$\begin{aligned} \Gamma(1) &= 1, & \Gamma(1/2) &= \sqrt{\pi} \\ \Gamma(p+1) &= p\Gamma(p) \quad \forall p > 0, & \Gamma(n+1) &= n!, \quad \forall n \in \mathbb{N}. \end{aligned}$$

2. Die *Beta-Funktion*:

$$B(p, q) = \int_0^1 t^{p-1} (1-t)^{q-1} dt, \quad p, q > 0.$$

Es gelten folgende Eigenschaften:

$$B(p, q) = B(q, p), \quad B(p, q) = \frac{\Gamma(p)\Gamma(q)}{\Gamma(p+q)}, \quad p, q > 0.$$

Definition 3.2.1

Die *Gamma-Verteilung* mit Parametern $\lambda > 0$ und $p > 0$ ist eine absolut stetige Verteilung mit der Dichte

$$f_X(x) = \begin{cases} \frac{\lambda^p x^{p-1}}{\Gamma(p)} e^{-\lambda x}, & x \geq 0, \\ 0, & x < 0. \end{cases} \quad (3.2.1)$$

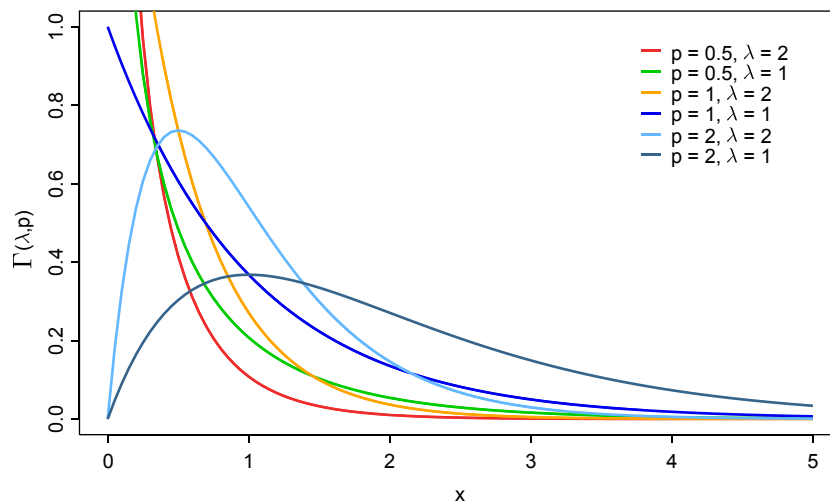


Abb. 3.1: Dichte der Gammaverteilung

Dabei verwenden wir die Bezeichnung $X \sim \Gamma(\lambda, p)$ für eine Zufallsvariable X , die Gamma-verteilt mit Parametern λ und p ist. Es gilt offensichtlich $X \geq 0$ fast sicher für $X \sim \Gamma(\lambda, p)$.

Übungsaufgabe 3.2.1

Zeigen Sie, dass (3.2.1) eine Dichte ist.

Beispiel 3.2.1

1. In der Kraftfahrzeugversicherung wird die Gamma-Verteilung oft zur Modellierung des Gesamtschadens verwendet.
2. Falls $p = 1$, dann ist $\Gamma(\lambda, 1) = \text{Exp}(\lambda)$.

Satz 3.2.1 (Momenterzeugende und charakteristische Funktion der Gammaverteilung):

Falls $X \sim \Gamma(\lambda, p)$, dann gilt Folgendes:

1. Die momenterzeugende Funktion der Gammaverteilung $\Psi_X(s)$ ist gegeben durch

$$\Psi_X(s) = \mathbb{E}e^{sX} = \frac{1}{(1 - s/\lambda)^p}, \quad s < \lambda.$$

Die charakteristische Funktion der Gammaverteilung $\varphi_X(s)$ ist gegeben durch

$$\varphi_X(s) = \mathbb{E}e^{isX} = \frac{1}{(1 - is/\lambda)^p}, \quad s \in \mathbb{R}.$$

2. k -te Momente:

$$\mathbb{E}X^k = \frac{p(p+1) \cdot \dots \cdot (p+k-1)}{\lambda^k}, \quad k \in \mathbb{N}.$$

Beweis 1. Betrachte

$$\begin{aligned}\Psi_X(s) &= \int_0^\infty e^{sx} f_X(x) dx = \frac{\lambda^p}{\Gamma(p)} \int_0^\infty x^{p-1} e^{\overbrace{(s-\lambda)x}^{<0}} dx \\ &= \frac{\lambda^p}{\Gamma(p)} \int_0^\infty \frac{y^{p-1}}{(-(s-\lambda))^p} e^{-y} dy = \frac{\lambda^p \Gamma(p)}{\Gamma(p)(\lambda-s)^p} \\ &= \left(\frac{\lambda}{\lambda-s} \right)^p = \frac{1}{(1-s/\lambda)^p}, \quad \lambda > s.\end{aligned}$$

Falls $s \in \mathbb{C}$, $\operatorname{Re}(s) < \lambda$, dann ist $\Psi_X(s)$ holomorph auf $D = \{s = x + iy \in \mathbb{C} : x < \lambda\}$. Es gilt

$$\Psi_X(s) = \varphi_X(-is), \quad s = it, t < \lambda$$

Daraus folgt

$$\Psi_X(s) = \varphi_X(-is), \quad s \in D \implies \varphi_X(s) = \frac{1}{(1-is/\lambda)^p}, \quad s \in \mathbb{R}.$$

2.

$$\mathbb{E}X^k = \Psi^{(k)}(0) \implies \mathbb{E}X^k = \frac{p \cdot (p+1) \cdot \dots \cdot (p+k-1)}{\lambda^k}, \quad k \in \mathbb{N}.$$

□

Folgerung 3.2.1 (Faltungsstabilität der Γ -Verteilung):

Falls $X \sim \Gamma(\lambda, p_1)$ und $Y \sim \Gamma(\lambda, p_2)$, X, Y unabhängig, dann ist $X + Y \sim \Gamma(\lambda, p_1 + p_2)$.

Beweis Es gilt

$$\varphi_{X+Y}(s) = \varphi_X(s) \cdot \varphi_Y(s) = \frac{1}{(1-is/\lambda)^{p_1}} \cdot \frac{1}{(1-is/\lambda)^{p_2}} = \left(\frac{1}{1-is/\lambda} \right)^{p_1+p_2} = \varphi_{\Gamma(\lambda, p_1+p_2)}(s).$$

Da die charakteristischen Funktionen die Verteilungen eindeutig bestimmen, folgt damit $X + Y \sim \Gamma(\lambda, p_1 + p_2)$. □

Beispiel 3.2.2

Seien $X_1, \dots, X_n \sim \operatorname{Exp}(\lambda)$ unabhängig. Nach der Folgerung 3.2.1 gilt $X = X_1 + \dots + X_n \sim \Gamma(\lambda, \underbrace{1 + \dots + 1}_n) = \Gamma(\lambda, n)$, denn $\operatorname{Exp}(\lambda) = \Gamma(\lambda, 1)$. Dabei heißt X *Erlang-verteilt* mit Parametern λ und n . Man schreibt $X \sim \operatorname{Erl}(\lambda, n)$.

$$\text{Zusammengefasst:} \quad \operatorname{Erl}(\lambda, n) = \Gamma(\lambda, n)$$

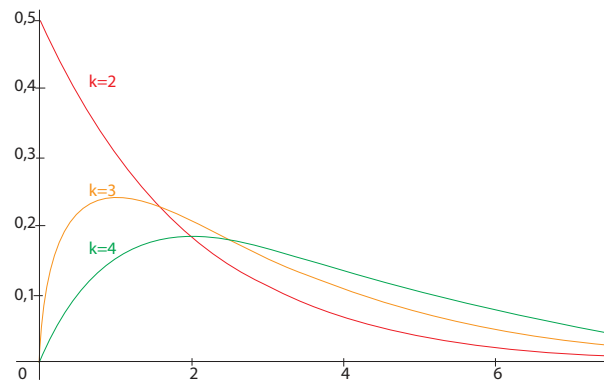
Interpretation: In der Risikotheorie z.B. sind X_i Zwischenankunftszeiten der Einzelschäden. Dann ist $X = \sum_{i=1}^n X_i$ die Ankunftszeit des n -ten Schadens, $X \sim \operatorname{Erl}(\lambda, n)$.

Definition 3.2.2 (χ^2 -Verteilung):

X ist eine χ^2 -verteilte Zufallsvariable mit k Freiheitsgraden ($X \sim \chi_k^2$), falls $X \stackrel{d}{=} X_1^2 + \dots + X_k^2$, wobei $X_1, \dots, X_k \sim N(0, 1)$ unabhängige identisch verteilte Zufallsvariablen sind.

Satz 3.2.2 (χ^2 -Verteilung: Spezialfall der Γ -Verteilung mit $\lambda = 1/2$, $p = k/2$):

Falls $X \sim \chi_k^2$, dann gilt:

Abb. 3.2: Dichte der χ^2 -Verteilung für $k = 2, 3, 4$

1. $X \sim \Gamma(1/2, k/2)$, d.h.

$$f_X(x) = \begin{cases} \frac{x^{k/2-1} e^{-x/2}}{2^{k/2} \Gamma(k/2)}, & x \geq 0 \\ 0, & x < 0 \end{cases}. \quad (3.2.2)$$

2. Insbesondere ist $\mathbb{E}X = k$, $\text{Var } X = 2k$.

Beweis 1. Sei $X = X_1^2 + \dots + X_k^2$ mit $X_i \sim N(0, 1)$ unabhängigen identisch verteilten Zufallsvariablen. Errechnen wir zunächst die Verteilung der X_i^2 :

$$\begin{aligned} P(X_1^2 \leq x) &= P(X_1 \in [-\sqrt{x}, \sqrt{x}]) = \int_{-\sqrt{x}}^{\sqrt{x}} \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} dy \\ &= \int_0^{\sqrt{x}} \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} dy + \int_{-\sqrt{x}}^0 \frac{1}{\sqrt{2\pi}} e^{-\frac{y^2}{2}} dy \\ &\stackrel{y^2=t}{=} \int_0^x \frac{1}{\sqrt{2\pi}} e^{-\frac{t}{2}} \frac{1}{2\sqrt{t}} dt + \int_x^0 \frac{1}{\sqrt{2\pi}} e^{-t/2} \frac{-1}{2\sqrt{t}} dt \\ &= \int_0^x \frac{(1/2)^{-1/2} t^{1/2-1}}{\Gamma(1/2)} e^{-t/2} dt, \quad x \geq 0. \end{aligned}$$

Somit folgt $X_1^2 \sim \Gamma(1/2, 1/2) \implies X \sim \Gamma(1/2, \underbrace{1/2 + \dots + 1/2}_k) = \Gamma(1/2, k/2)$ und daher gilt der Ausdruck (3.2.2) für die Dichte.

2. Wegen der Additivität des Erwartungswertes und der Unabhängigkeit von X_i gilt

$$\mathbb{E}X = k \cdot \mathbb{E}X_1^2, \quad \text{Var } X = k \text{Var } X_1^2, \quad \mathbb{E}(X_1^2) = \mathbb{E}(\Gamma(1/2, 1/2)).$$

Bitte zeigen Sie selbstständig, dass $\mathbb{E}X_1^2 = 1$, $\text{Var } X_1^2 = 2$.

□

3.2.2 Student¹-Verteilung (t-Verteilung)

Definition 3.2.3

Seien X, Y unabhängige Zufallsvariablen, wobei $X \sim N(0, 1)$ und $Y \sim \chi_r^2$. Dann heißt die

¹Genannt nach dem Entdecker William Sealy Gosset, der seine Arbeiten mit Student unterzeichnete.

Zufallsvariable

$$U \stackrel{d}{=} \frac{X}{\sqrt{Y/r}}$$

Student- oder t -verteilt mit r Freiheitsgraden. Wir schreiben $U \sim t_r$.

Satz 3.2.3 (Dichte der t -Verteilung):

Falls $X \sim t_r$, dann gilt:

1.

$$f_X(x) = \frac{1}{\sqrt{r} B\left(\frac{r}{2}, \frac{1}{2}\right)} \cdot \frac{1}{\left(1 + \frac{x^2}{r}\right)^{\frac{r+1}{2}}}, \quad x \in \mathbb{R}.$$

2. $\mathbb{E}X = 0$, $\text{Var } X = \frac{r}{r-2}$, $r \geq 3$.

Bemerkung 3.2.1

1. **Grafik von f_r :** Die t_r -Verteilung ist symmetrisch. Insbesondere gilt:

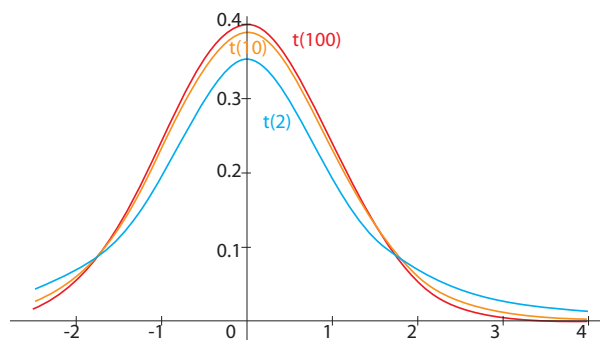


Abb. 3.3: Dichte $\hat{f}$ der t -Verteilung für $r = 2, 10, 100$

$$t_{r,\alpha} = -t_{r,1-\alpha}, \quad \alpha \in (0, 1),$$

wobei $t_{r,\alpha}$ das α -Quantil der Student-Verteilung mit r Freiheitsgraden ist.

2. Falls $r \rightarrow \infty$, dann $f_r(x) \rightarrow \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$, $x \in \mathbb{R}$. (Übungsaufgabe)
3. Für $r = 1$ gilt: $t_1 = \text{Cauchy}(0, 1)$ mit Dichte $f(x) = \frac{1}{\pi(1+x^2)}$. Der Erwartungswert von t_1 existiert nicht.

Beweis des Satzes 3.2.3:

1. Es gilt $X := \varphi(Y, Z)$, wobei $\varphi(x, y) = \frac{x}{\sqrt{y/r}}$ und $V = (Y, Z)$ ein zweidimensionaler Zufallsvektor ist, $Y \sim N(0, 1)$, $Z \sim \chi_r^2$, Y und Z unabhängig.

Wir wollen den sogenannten *Dichtetransformationssatz für Zufallsvektoren* verwenden, der besagt, dass unter bestimmten Voraussetzungen

$$f_{\varphi(V)}(x) = f_V(\varphi^{-1}(x))|J|$$

gilt, wobei $|J| = |\det J|$, $J = \left(\frac{\partial \varphi_i^{-1}(x)}{\partial x_j} \right)_{i,j=1}^n$, $\varphi = (\varphi_1, \dots, \varphi_n) : \mathbb{R}^n \rightarrow \mathbb{R}^n$. Berechnen wir hier φ^{-1} von $\varphi : (x, y) \mapsto (v, w)$, wobei $v = \frac{x}{\sqrt{y/r}}$, $w = y$:

$$\varphi^{-1} : v = \frac{x}{\sqrt{\frac{y}{r}}} \implies x = v \sqrt{\frac{y}{r}} = v \sqrt{\frac{w}{r}}. \quad \text{Somit } \varphi^{-1} : (v, w) \mapsto \left(v \sqrt{\frac{w}{r}}, w \right)$$

und die Jacobi-Matrix ist gleich

$$J = \begin{pmatrix} \frac{\partial \varphi_1^{-1}}{\partial v} & \frac{\partial \varphi_1^{-1}}{\partial w} \\ \frac{\partial \varphi_2^{-1}}{\partial v} & \frac{\partial \varphi_2^{-1}}{\partial w} \end{pmatrix} = \begin{pmatrix} \sqrt{\frac{w}{r}} & \frac{v}{2\sqrt{wr}} \\ 0 & 1 \end{pmatrix}.$$

Falls $V = (Y, Z)$, Y und Z unabhängig, dann

$$f_V(x, y) = f_Y(x) \cdot f_Z(y) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} \cdot \frac{y^{r/2-1} e^{-y/2}}{\Gamma(r/2) 2^{r/2}} = \frac{y^{r/2-1} e^{-\frac{y+x^2}{2}}}{2^{\frac{r+1}{2}} \Gamma(1/2) \Gamma(r/2)}, \quad x \in \mathbb{R}, y > 0,$$

und nach dem Dichtetransformationssatz gilt

$$\begin{aligned} f_X(v) &= \int_0^\infty f_{\varphi(V)}(u, w) dw = \int_0^\infty f_V(\varphi^{-1}(v, w)) |J| dw \\ &= \int_0^\infty \frac{e^{-(v^2 \frac{w}{r} + w)/2} w^{r/2-1}}{2^{\frac{r+1}{2}} \Gamma(1/2) \Gamma(r/2)} \sqrt{w/r} dw \\ &= \frac{1}{\sqrt{r} 2^{\frac{r+1}{2}} \Gamma(1/2) \Gamma(r/2)} \cdot \int_0^\infty w^{\frac{r-1}{2}} e^{-\overbrace{\frac{v^2}{r} + 1}^t \cdot w} dw \\ &\stackrel{w = \frac{2t}{v^2/r + 1}}{=} \frac{1}{\sqrt{r} 2^{\frac{r+1}{2}} \Gamma(1/2) \Gamma(r/2)} \cdot \int_0^\infty \frac{2^{\frac{r-1}{2} + 1} t^{\frac{r-1}{2}}}{(v^2/r + 1)^{\frac{r-1}{2} + 1}} e^{-t} dt \\ &= \frac{2^{\frac{r+1}{2}} \Gamma(\frac{r+1}{2})}{(v^2/r + 1)^{\frac{r+1}{2}} \sqrt{r} 2^{\frac{r+1}{2}} \Gamma(1/2) \Gamma(r/2)} = \frac{1}{\sqrt{r} B(r/2, 1/2) (1 + v^2/r)^{\frac{r+1}{2}}} \end{aligned}$$

2. Übungsaufgabe

□

Da im WR-Skript der Dichtetransformationssatz nur für Zufallsvariablen formuliert wurde, geben wir hier die notwendigen Begriffe und verallgemeinerten Sätze für Zufallsvektoren (ohne Beweis). Hierbei verwenden wir die folgende Notation:

Für Vektoren $x = (x_1, \dots, x_n)^T$ und $y = (y_1, \dots, y_n)^T$ schreiben wir $x \leq y$, falls $x_i \leq y_i$ für $i = 1, \dots, n$. Ferner sei für einen Zufallsvektor $X = (X_1, \dots, X_n)^T$ die Verteilungsfunktion definiert als $F(x) = \mathbb{P}(X \leq x)$ für $x = (x_1, \dots, x_n)^T$.

Definition 3.2.4

Die Zufallsvektoren $X_i : \Omega \rightarrow \mathbb{R}^{m_i}$, $i = 1, \dots, n$ sind *unabhängig*, falls

$$F_{(X_1, \dots, X_n)}(x_1, \dots, x_n) = P(X_1 \leq x_1, \dots, X_n \leq x_n) = \prod_{i=1}^n P(X_i \leq x_i) = \prod_{i=1}^n F_{X_i}(x_i),$$

$x_i \in \mathbb{R}^{m_i}$, $i = 1, \dots, n$.

Satz 3.2.4

Falls X_i absolut stetig verteilte und unabhängige Zufallsvektoren mit Dichten f_{X_i} , $i = 1, \dots, n$, sind, dann ist auch $(X_1, \dots, X_n)$ absolut stetig verteilt mit Dichte

$$f_{X_1, \dots, X_n}(x_1, \dots, x_n) = \prod_{i=1}^n f_{X_i}(x_i), \quad x_i \in \mathbb{R}^{m_i}, \quad i = 1, \dots, n.$$

Satz 3.2.5

Falls $X_i : \Omega \rightarrow \mathbb{R}^{m_i}$, $i = 1, \dots, n$ unabhängige Zufallsvektoren sind, und $\varphi_i : \mathbb{R}^{m_i} \rightarrow \mathbb{R}^{n_i}$, $\forall i = 1, \dots, n$ Borel-messbare Funktionen, dann sind Zufallsvektoren $\varphi_1(X_1), \dots, \varphi_n(X_n)$ unabhängig.

Satz 3.2.6 (Dichtetransformationssatz für Zufallsvektoren):

Sei $X = (X_1, \dots, X_m)^T : \Omega \rightarrow \mathbb{R}^m$ ein absolut stetig verteilter Zufallsvektor mit Dichte f_X . Sei $\varphi = (\varphi_1, \dots, \varphi_m)^T : \mathbb{R}^m \rightarrow \mathbb{R}^m$ eine Borel-messbare Abbildung, die innerhalb von einem Quader $B \subset \mathbb{R}^m$ stetig differenzierbar ist. Falls $\text{supp} f_X \subset B$ und $\det \left(\frac{\partial \varphi_i}{\partial x_j} \right)_{i,j=1, \dots, m} \neq 0$ auf B , dann $\exists \varphi^{-1} : \varphi(B) \rightarrow B$ stetig differenzierbar und

$$f_{\varphi(X)}(x) = \begin{cases} f_X(\varphi^{-1}(x)) \cdot |J|, & x \in \varphi(B), \\ 0, & x \notin \varphi(B), \end{cases}$$

wobei $J = \det \left(\frac{\partial \varphi_i^{-1}}{\partial x_j} \right)_{i,j=1, \dots, m}$

3.2.3 Fisher-Snedecor-Verteilung (F-Verteilung)**Definition 3.2.5**

Falls $X \stackrel{d}{=} \frac{U_r/r}{U_s/s}$, wobei $U_r \sim \chi_r^2$, $U_s \sim \chi_s^2$, $r, s \in \mathbb{N}$, U_r, U_s unabhängig, dann hat X eine F-Verteilung mit Freiheitsgraden r, s . Bezeichnung: $X \sim F_{r,s}$.

Lemma 3.2.1

Falls $X \sim F_{r,s}$, dann ist X absolut stetig verteilt mit Dichte

$$f_X(x) = \frac{x^{r/2-1}}{B(r/2, s/2)(r/s)^{-r/2}(1 + (r/s) \cdot x)^{\frac{r+s}{2}}} \cdot \mathbb{I}(x > 0).$$

Beweis Da $U_r \sim \chi_r^2$, gilt für ihre Dichte

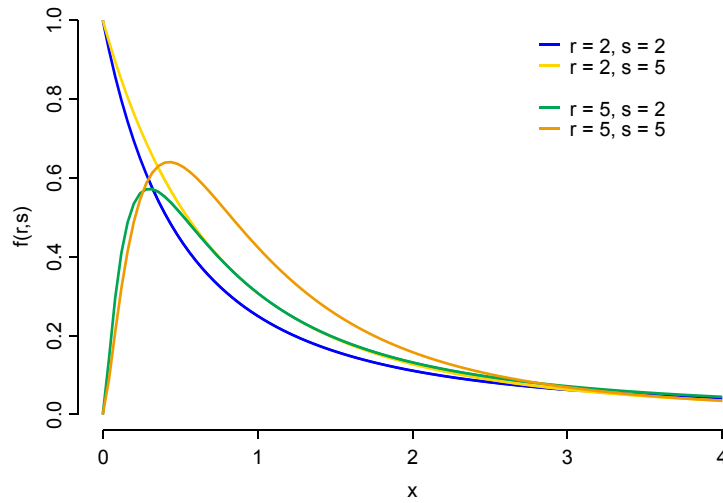
$$f_{U_r}(x) = \frac{x^{r/2-1} e^{-x/2}}{\Gamma(r/2) 2^{r/2}}, \quad x > 0, \quad r \in \mathbb{N}.$$

Somit

$$P(U_r/r \leq x) = P(U_r \leq rx) = F_{U_r}(rx)$$

und deshalb

$$\begin{aligned} f_{U_r/r}(x) &= (F_{U_r}(rx))' = r \cdot f_{U_r}(rx) = \frac{r(rx)^{r/2-1} e^{-\frac{rx}{2}}}{\Gamma(r/2) 2^{r/2}} \cdot \mathbb{I}(x > 0) \\ &= \frac{r^{r/2} x^{r/2-1} e^{-r/2 \cdot x}}{\Gamma(r/2) 2^{r/2}} \cdot \mathbb{I}(x > 0). \end{aligned}$$

Abb. 3.4: Dichte der F-Verteilung für verschiedene Parameter r und s .

Nach dem Dichtetransformationssatz für das Verhältnis von zwei Zufallsvariablen (vgl. Wahrscheinlichkeitsskript Satz 3.15) gilt

$$f_{\frac{U_r/r}{U_s/s}}(x) = \int_0^\infty t f_{U_r/r}(xt) \cdot f_{U_s/s}(t) dt \cdot \mathbb{I}(x > 0).$$

Somit

$$\begin{aligned} f_X(x) &= \int_0^\infty t \frac{r^{r/2} (tx)^{r/2-1} e^{-\frac{rtx}{2}}}{\Gamma(r/2) 2^{r/2}} \cdot \frac{s^{s/2} t^{s/2-1} e^{-st/2}}{\Gamma(s/2) 2^{s/2}} dt \\ &= \frac{r^{r/2} s^{s/2} x^{r/2-1}}{\Gamma(r/2) \Gamma(s/2) 2^{\frac{r+s}{2}}} \cdot \int_0^\infty t^{r/2+s/2-1} e^{-\frac{\overbrace{rx+s}^y}{2} t} dt \\ &= \frac{r^{r/2} s^{s/2} x^{r/2-1}}{\Gamma(r/2) \Gamma(s/2)} \cdot \int_0^\infty \frac{y^{\frac{r+s}{2}-1}}{(rx+s)^{\frac{r+s}{2}}} \cdot e^{-y} dy \\ &\stackrel{t=\frac{y}{rx+s}}{=} \frac{r^{r/2} s^{s/2} x^{r/2-1}}{\Gamma(r/2) \Gamma(s/2) s^{\frac{r+s}{2}} (1 + \frac{r}{s} \cdot x)^{\frac{r+s}{2}}} \cdot \Gamma\left(\frac{r+s}{2}\right) \\ &= \frac{(r/s)^{r/2} x^{r/2-1}}{B(r/2, s/2) (1 + \frac{r}{s} x)^{\frac{r+s}{2}}} \cdot \mathbb{I}(x > 0). \end{aligned}$$

□

Bemerkung 3.2.2

Sei $X \sim F_{r,s}$, $r, s \in \mathbb{N}$ mit Dichte f_X .

1. Einige Graphen der F-Verteilung sind in Abbildung 3.4 dargestellt.
2. Einige Eigenschaften der F-Verteilung:

Lemma 3.2.2

Es gilt:

a)
$$\mathbb{E}X = \frac{s}{s-2}, \quad s \geq 3.$$

b)
$$\text{Var } X = \frac{2s^2(r+s-2)}{r(s-4)(s-2)^2}, \quad s \geq 5.$$

c) Falls $F_{r,s,\alpha}$ das α -Quantil der $F_{r,s}$ -Verteilung ist, dann gilt

$$F_{r,s,\alpha} = \frac{1}{F_{s,r,1-\alpha}}, \quad \alpha \in (0,1).$$

Übungsaufgabe 3.2.2

Beweisen Sie Lemma 3.2.2!

3. Für Quantile $F_{r,s,\alpha}$ gilt folgende Näherungsformel (Abramowitz, Stegun (1972)):

$F_{r,s,\alpha} \approx e^\omega$, wobei

$$\begin{aligned} \omega &= 2 \left(\frac{\alpha(h+a)^{1/2}}{h} - \left(\frac{1}{r-1} - \frac{1}{s-1} \right) \cdot \left(a + \frac{5}{6} - \frac{2}{3h} \right) \right), \\ h &= 2 \left(\frac{1}{r-1} + \frac{1}{s-1} \right)^{-1}, \\ a &= \frac{z_\alpha^2 - 3}{6} \end{aligned}$$

und z_α das α -Quantil der $N(0,1)$ -Verteilung ist.

3.3 Punktschätzer und ihre Grundeigenschaften

Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe, definiert auf dem kanonischen Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P}_\theta)$. Seien X_i , $i = 1, \dots, n$ unabhängige identisch verteilte Zufallsvariablen mit Verteilungsfunktion $F \in \{F_\theta : \theta \in \Theta\}$, $\Theta \subset \mathbb{R}^m$. Finde einen Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für den Parameter θ mit vorgegebenen Eigenschaften.

Unser Ziel im nächsten Abschnitt ist es, zunächst grundlegende Eigenschaften der Schätzer kennenzulernen.

3.3.1 Eigenschaften von Punktschätzern**Definition 3.3.1 (Erwartungstreue):**

Ein Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für θ heißt *erwartungstreu* oder *unverzerrt*, falls

$$\mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n) = \theta, \quad \theta \in \Theta.$$

Dabei wird vorausgesetzt, dass

$$\mathbb{E}_\theta |\hat{\theta}(X_1, \dots, X_n)| < \infty, \quad \theta \in \Theta.$$

Der *Bias* (*Verzerrung*) eines Schätzers $\hat{\theta}(X_1, \dots, X_n)$ ist gegeben durch

$$\text{Bias}(\hat{\theta}) = \mathbb{E}_{\theta} \hat{\theta}(X_1, \dots, X_n) - \theta.$$

Falls $\hat{\theta}(X_1, \dots, X_n)$ erwartungstreu ist, dann gilt $\text{Bias}(\hat{\theta}) = 0$ (kein systematischer Schätzfehler).

Definition 3.3.2 (Asymptotische Erwartungstreue):

Der Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für θ heißt *asymptotisch erwartungstreu* (oder *asymptotisch unverzerrt*), falls (für große Datenmengen)

$$\mathbb{E}_{\theta} \hat{\theta}(X_1, \dots, X_n) \xrightarrow{n \rightarrow \infty} \theta, \quad \theta \in \Theta.$$

Definition 3.3.3 (Konsistenz):

Falls

$$\hat{\theta}(X_1, \dots, X_n) \xrightarrow{n \rightarrow \infty} \theta, \quad \theta \in \Theta$$

in L^2 , stochastisch bzw. fast sicher, dann heißt der Schätzer $\hat{\theta}(X_1, \dots, X_n)$ ein *konsistenter Schätzer* für θ im *mittleren quadratischen*, *schwachen* bzw. *starken Sinne*.

- $\hat{\theta}$ *L^2 -konsistent*: für $\mathbb{E}_{\theta} \hat{\theta}^2(X_1, \dots, X_n) < \infty$ gilt

$$\hat{\theta} \xrightarrow[n \rightarrow \infty]{L^2} \theta \iff \mathbb{E}_{\theta} |\hat{\theta}(X_1, \dots, X_n) - \theta|^2 \xrightarrow{n \rightarrow \infty} 0, \quad \theta \in \Theta.$$

- $\hat{\theta}$ *schwach konsistent*:

$$\hat{\theta} \xrightarrow[n \rightarrow \infty]{\mathbb{P}} \theta \iff P_{\theta}(|\hat{\theta}(X_1, \dots, X_n) - \theta| > \varepsilon) \xrightarrow{n \rightarrow \infty} 0, \quad \varepsilon > 0, \quad \theta \in \Theta.$$

- $\hat{\theta}$ *stark konsistent*:

$$\hat{\theta} \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \theta \iff P_{\theta} \left(\lim_{n \rightarrow \infty} \hat{\theta}(X_1, \dots, X_n) = \theta \right) = 1, \quad \theta \in \Theta.$$

Daraus ergibt sich folgendes Diagramm (vgl. Wahrscheinlichkeitsrechnungsskript, Kapitel 6).



Definition 3.3.4 (Mittlerer quadratischer Fehler (mean squared error)):

Der mittlere quadratische Fehler eines Schätzers $\hat{\theta}(X_1, \dots, X_n)$ für θ ist definiert als

$$MSE(\hat{\theta}) = \mathbb{E}_{\theta} |\hat{\theta}(X_1, \dots, X_n) - \theta|^2.$$

Lemma 3.3.1

Falls $m = 1$ und $\mathbb{E}_{\theta} \hat{\theta}^2(X_1, \dots, X_n) < \infty$, $\theta \in \Theta$, dann gilt

$$MSE(\hat{\theta}) = \text{Var}_{\theta} \hat{\theta} + (\text{Bias}(\hat{\theta}))^2.$$

Beweis

$$\begin{aligned} MSE(\hat{\theta}) &= \mathbb{E}_{\theta} (\hat{\theta} - \theta)^2 = \mathbb{E}_{\theta} (\hat{\theta} - \mathbb{E}_{\theta} \hat{\theta} + \mathbb{E}_{\theta} \hat{\theta} - \theta)^2 \\ &= \underbrace{\mathbb{E}_{\theta} (\hat{\theta} - \mathbb{E}_{\theta} \hat{\theta})^2}_{\text{Var}_{\theta} \hat{\theta}} + \underbrace{2 \mathbb{E}_{\theta} (\hat{\theta} - \mathbb{E}_{\theta} \hat{\theta}) (\mathbb{E}_{\theta} \hat{\theta} - \theta)}_{=0} + \underbrace{(\mathbb{E}_{\theta} \hat{\theta} - \theta)^2}_{= \text{Bias}(\hat{\theta})^2} \\ &= \text{Var}_{\theta} \hat{\theta} + (\text{Bias}(\hat{\theta}))^2. \end{aligned}$$

□

Bemerkung 3.3.1

Falls $\hat{\theta}$ erwartungstreu für θ ist, dann gilt $MSE(\hat{\theta}) = \text{Var}_\theta \hat{\theta}$.

Definition 3.3.5 (Vergleich von Schätzern):

Seien $\hat{\theta}_1(X_1, \dots, X_n)$ und $\hat{\theta}_2(X_1, \dots, X_n)$ zwei Schätzer für θ . Man sagt, dass $\hat{\theta}_1$ *besser* ist als $\hat{\theta}_2$, falls

$$MSE(\hat{\theta}_1) < MSE(\hat{\theta}_2), \quad \theta \in \Theta.$$

Falls $m = 1$ und die Schätzer $\hat{\theta}_1, \hat{\theta}_2$ erwartungstreu sind, so ist $\hat{\theta}_1$ *besser als* $\hat{\theta}_2$, falls $\hat{\theta}_1$ die kleinere Varianz besitzt. Dabei wird stets vorausgesetzt, dass $\mathbb{E}_\theta \hat{\theta}_i^2 < \infty$, $\theta \in \Theta$.

Definition 3.3.6 (Asymptotische Normalverteiltheit):

Sei $\hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ ($m = 1$). Falls $0 < \text{Var}_\theta \hat{\theta}(X_1, \dots, X_n) < \infty$, $\theta \in \Theta$ und

$$\frac{\hat{\theta}(X_1, \dots, X_n) - \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n)}{\sqrt{\text{Var}_\theta \hat{\theta}(X_1, \dots, X_n)}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1),$$

dann ist $\hat{\theta}(X_1, \dots, X_n)$ *asymptotisch normalverteilt*.

Definition 3.3.7 (Bester erwartungstreuer Schätzer):

Der Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für θ ist der *beste erwartungstreue Schätzer*, falls

$$\mathbb{E}_\theta \hat{\theta}^2(X_1, \dots, X_n) < \infty, \quad \theta \in \Theta, \quad \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n) = \theta, \quad \theta \in \Theta,$$

und $\hat{\theta}$ die minimale Varianz in der Klasse aller erwartungstreuen Schätzer für θ besitzt. Das heißt, dass für einen beliebigen erwartungstreuen Schätzer $\tilde{\theta}(X_1, \dots, X_n)$ mit

$$\mathbb{E}_\theta \tilde{\theta}^2(X_1, \dots, X_n) < \infty \quad \text{gilt} \quad \text{Var}_\theta \hat{\theta} \leq \text{Var}_\theta \tilde{\theta}, \quad \theta \in \Theta.$$

3.3.2 Schätzer des Erwartungswertes und empirische Momente

Sei $X \stackrel{d}{=} X_i$, $i = 1, \dots, n$ ein statistisches Merkmal. Sei weiter $\mathbb{E}|X_i|^k < \infty$ für ein $k \in \mathbb{N}$, $m = 1$ und der zu schätzende Parameter $\theta = \mu_k = \mathbb{E}X_i^k$. Insbesondere gilt im Fall $k = 1$, dass $\theta = \mu_1 = \mu$ der Erwartungswert ist.

Definition 3.3.8

Das k -te *empirische Moment* von X wird als

$$\hat{\mu}_k = \frac{1}{n} \sum_{i=1}^n X_i^k$$

definiert. Unter dieser Definition gilt, dass $\hat{\mu}_1 = \bar{X}_n$, also das erste empirische Moment gleich dem Stichprobenmittel ist.

Satz 3.3.1 (Eigenschaften der empirischen Momente):

Unter obigen Voraussetzungen gelten folgende Eigenschaften:

1. $\hat{\mu}_k$ ist erwartungstreu für μ_k (insbesondere $\bar{X}_n$).
2. $\hat{\mu}_k$ ist stark konsistent.
3. Falls $\mathbb{E}_\theta |X|^{2k} < \infty$, $\forall \theta \in \Theta$, dann ist $\hat{\mu}_k$ asymptotisch normalverteilt.

4. Es gilt $\text{Var } \bar{X}_n = \frac{\sigma^2}{n}$, wobei $\sigma^2 = \text{Var}_\theta X$. Falls $X_i \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$ (eine normalverteilte Stichprobe), dann gilt:

$$\bar{X}_n \sim N\left(\mu, \frac{\sigma^2}{n}\right).$$

Beweis

1.
$$\mathbb{E}_\theta \hat{\mu}_k = \frac{1}{n} \sum_{i=1}^n \mathbb{E}_\theta X_i^k = \frac{1}{n} \sum_{i=1}^n \mu_k = \frac{n\mu_k}{n} = \mu_k.$$

2. Aus dem starken Gesetz der großen Zahlen folgt

$$\frac{1}{n} \sum_{i=1}^n X_i^k \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta X_i^k = \mu_k.$$

3. Mit dem zentralen Grenzwertsatz gilt

$$\frac{\sum_{i=1}^n X_i^k - n \cdot \mathbb{E} X^k}{\sqrt{n \cdot \text{Var } X^k}} = \frac{\frac{1}{n} \sum_{i=1}^n X_i^k - \mu_k}{\frac{1}{\sqrt{n}} \sqrt{\text{Var } X^k}} = \sqrt{n} \frac{\hat{\mu}_k - \mu_k}{\sqrt{\text{Var } X^k}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1).$$

Insbesondere gilt für den Spezialfall $k = 1$

$$\sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1).$$

- 4.

$$\text{Var } \bar{X}_n = \text{Var} \left(\frac{1}{n} \sum_{i=1}^n X_i \right) \underset{X_i \text{ u.i.v.}}{=} \frac{1}{n^2} \sum_{i=1}^n \text{Var } X_i = \frac{n \cdot \sigma^2}{n^2} = \frac{\sigma^2}{n}.$$

Falls $X_i \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$, dann gilt wegen der Faltungsstabilität der Normalverteilung $\bar{X}_n \sim N(\cdot, \cdot)$, weil

$$\frac{1}{n} X_i \sim N\left(\frac{\mu}{n}, \frac{\sigma^2}{n^2}\right), \quad X_i \text{ u.i.v.}$$

Somit folgt aus 1) und 4) $\bar{X}_n \sim N\left(\mu, \frac{\sigma^2}{n}\right)$.

Damit ist der Satz bewiesen. □

Bemerkung 3.3.2

Aus Satz 3.3.1, 3) folgt

$$\begin{aligned} \mathbb{P}(|\bar{X}_n - \mu| > \varepsilon) &= 1 - \mathbb{P}(-\varepsilon \leq \bar{X}_n - \mu \leq \varepsilon) \\ &= 1 - \mathbb{P}\left(-\frac{\varepsilon\sqrt{n}}{\sigma} \leq \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \leq \frac{\varepsilon\sqrt{n}}{\sigma}\right) \\ &\underset{n \rightarrow \infty}{\approx} 1 - \left(\Phi\left(\frac{\varepsilon\sqrt{n}}{\sigma}\right) - \Phi\left(-\frac{\varepsilon\sqrt{n}}{\sigma}\right)\right) \\ &\stackrel{\Phi(-x)=1-\Phi(x)}{=} 1 - \left(\Phi\left(\frac{\varepsilon\sqrt{n}}{\sigma}\right) - 1 + \Phi\left(\frac{\varepsilon\sqrt{n}}{\sigma}\right)\right) \\ &= 1 - \left(2\Phi\left(\frac{\varepsilon\sqrt{n}}{\sigma}\right) - 1\right), \end{aligned}$$

wobei $\Phi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{t^2}{2}} dt$ die Verteilungsfunktion der $N(0, 1)$ -Verteilung ist. Insgesamt gilt also für großes n

$$\mathbb{P}(|\bar{X}_n - \mu| > \varepsilon) \approx 2 \left(1 - \Phi \left(\frac{\varepsilon \sqrt{n}}{\sigma} \right) \right).$$

3.3.3 Schätzer der Varianz

Seien X_i , $i = 1, \dots, n$ unabhängig identisch verteilt, $X_i \stackrel{d}{=} X$, $\mathbb{E}_\theta X^2 < \infty \quad \forall \theta \in \Theta$, $\theta = (\theta_1, \dots, \theta_m)^T$, $\theta_i = \sigma^2 = \text{Var}_\theta X$ für ein $i \in \{1, \dots, m\}$. Die *Stichprobenvarianz*

$$S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2$$

ist dann ein Schätzer für σ^2 . Falls der Erwartungswert $\mu = \mathbb{E}_\theta X$ der Stichprobenvariablen explizit benannt ist, so kann ein Schätzer für σ^2 auch als

$$\tilde{S}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2$$

definiert werden.

Wir werden nun die Eigenschaften von S_n^2 und $\tilde{S}_n^2$ untersuchen und sie miteinander vergleichen.

Satz 3.3.2

1. Die Stichprobenvarianz S_n^2 ist erwartungstreu für σ^2 :

$$\mathbb{E}_\theta S_n^2 = \sigma^2, \quad \theta \in \Theta.$$

2. Wenn $\mathbb{E}_\theta X^4 < \infty$, dann gilt

$$\text{Var}_\theta S_n^2 = \frac{1}{n} \left(\mu'_4 - \frac{n-3}{n-1} \sigma^4 \right),$$

wobei $\mu'_4 = \mathbb{E}_\theta (X - \mu)^4$.

Beweis 1. Aus Lemma 2.2.1 1), 2) folgt, dass

$$S_n^2 = \frac{1}{n-1} \left(\sum_{i=1}^n X_i^2 - n \bar{X}_n^2 \right),$$

und dass man o.B.d.A. $\mu = \mathbb{E}_\theta X_i = 0$ annehmen kann, woraus insbesondere $\mathbb{E}_\theta \bar{X}_n = 0$, $\theta \in \Theta$ folgt. Dann gilt

$$\begin{aligned} \mathbb{E}_\theta S_n^2 &= \frac{1}{n-1} \left(\sum_{i=1}^n \mathbb{E}_\theta X_i^2 - n \mathbb{E}_\theta \bar{X}_n^2 \right) = \frac{1}{n-1} \left(\sum_{i=1}^n \text{Var}_\theta X_i - n \text{Var}_\theta \bar{X}_n \right) \\ &\stackrel{\text{s. 3.3.1, 4)}}{=} \frac{1}{n-1} \left(n \sigma^2 - n \cdot \frac{\sigma^2}{n} \right) = \sigma^2, \quad \theta \in \Theta. \end{aligned}$$

2. Berechnen wir $\text{Var}_\theta S_n^2 = \mathbb{E}_\theta(S_n^2)^2 - (\mathbb{E}_\theta S_n^2)^2 = \mathbb{E}_\theta(S_n^2)^2 - \sigma^4$. Es gilt

$$\begin{aligned}\mathbb{E}_\theta S_n^4 &= \frac{1}{(n-1)^2} \mathbb{E}_\theta \left(\sum_{i=1}^n X_i^2 - n \bar{X}_n^2 \right)^2 \\ &= \frac{1}{(n-1)^2} \left(\underbrace{\mathbb{E}_\theta \left(\sum_{i=1}^n X_i^2 \right)^2}_{=I_1} - 2n \underbrace{\mathbb{E}_\theta \left(\bar{X}_n^2 \sum_{i=1}^n X_i^2 \right)}_{=I_2} + n^2 \underbrace{\mathbb{E}_\theta \bar{X}_n^4}_{=I_3} \right).\end{aligned}$$

Dabei gilt

$$\begin{aligned}I_1 &= \mathbb{E}_\theta \left(\sum_{i=1}^n X_i^2 \sum_{j=1}^n X_j^2 \right) = \mathbb{E}_\theta \left(\sum_{i=1}^n X_i^4 + \sum_{i \neq j} X_i^2 X_j^2 \right) = \sum_{i=1}^n \mathbb{E}_\theta X_i^4 + \sum_{i \neq j} \mathbb{E}_\theta (X_i^2 X_j^2) \\ &\stackrel{X_i \text{ u.i.v.}, \mu=0}{=} \sum_{i=1}^n \mu'_4 + \sum_{i \neq j} \text{Var}_\theta X_i \cdot \text{Var}_\theta X_j = n\mu'_4 + n(n-1)\sigma^4,\end{aligned}$$

$$\begin{aligned}I_2 &= \mathbb{E}_\theta \left(\left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 \sum_{j=1}^n X_j^2 \right) = \frac{1}{n^2} \mathbb{E}_\theta \left(\left(\sum_{i=1}^n X_i^2 + \sum_{i \neq j} X_i X_j \right) \sum_{j=1}^n X_j^2 \right) \\ &= \frac{1}{n^2} \mathbb{E}_\theta \left(\sum_{i=1}^n X_i^2 \sum_{j=1}^n X_j^2 \right) + \frac{1}{n^2} \mathbb{E}_\theta \left(\sum_{i \neq j} X_i X_j \sum_{k=1}^n X_k^2 \right) \\ &= \frac{1}{n^2} I_1 + \frac{1}{n^2} \sum_{i \neq j} \sum_k \underbrace{\mathbb{E}_\theta (X_i X_j X_k^2)}_{=0, \text{ da } X_i \text{ u.i.v. und } \mu=0} = \frac{I_1}{n^2} = \frac{\mu'_4 + (n-1)\sigma^4}{n},\end{aligned}$$

$$\begin{aligned}I_3 &= \mathbb{E}_\theta \left(\left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 \cdot \left(\frac{1}{n} \sum_{j=1}^n X_j \right)^2 \right) \\ &= \frac{1}{n^4} \mathbb{E}_\theta \left(\left(\sum_{k=1}^n X_k^2 + \sum_{i \neq j} X_i X_j \right) \cdot \left(\sum_{r=1}^n X_r^2 + \sum_{s \neq t} X_s X_t \right) \right) \\ &= \frac{1}{n^4} \mathbb{E}_\theta \left(\sum_{k,r=1}^n X_k^2 X_r^2 + 2 \sum_{k=1}^n X_k^2 \sum_{i \neq j} X_i X_j + \sum_{i \neq j} X_i X_j \sum_{s \neq t} X_s X_t \right) \\ &= \frac{1}{n^4} \left(\mathbb{E}_\theta \left(\sum_{k=1}^n X_k^4 \right) + \mathbb{E}_\theta \left(\sum_{k \neq r} X_k^2 X_r^2 \right) + 2 \underbrace{\mathbb{E}_\theta \left(\sum_{k=1}^n X_k^2 \sum_{i \neq j} X_i X_j \right)}_{=0, \text{ da } X_i \text{ u.i.v. und } \mu=0} \right. \\ &\quad \left. + 2 \underbrace{\mathbb{E}_\theta \left(\sum_{i \neq j} X_i^2 X_j^2 \right)}_{\text{weil } (i,j) \text{ und } (j,i) \text{ zählen}} + \underbrace{\mathbb{E}_\theta \left(\sum_{i \neq j \neq t} X_i^2 X_j X_t \right)}_{=0, \text{ da } X_j \text{ u.i.v. und } \mu=0} + \underbrace{\mathbb{E}_\theta \left(\sum_{i \neq j \neq s \neq t} X_i X_j X_s X_t \right)}_{=0, \text{ da } X_i \text{ u.i.v. und } \mu=0} \right) \\ &= \frac{1}{n^4} \left(n\mu'_4 + 3\mathbb{E}_\theta \sum_{i \neq j} X_i^2 X_j^2 \right) = \frac{n\mu'_4 + 3n(n-1)\sigma^4}{n^4} = \frac{\mu'_4 + 3(n-1)\sigma^4}{n^3}.\end{aligned}$$

Somit gilt insgesamt

$$\begin{aligned}\mathbb{E}_\theta S_n^4 &= \frac{1}{(n-1)^2} \left(n\mu'_4 + n(n-1)\sigma^4 - 2(\mu'_4 + (n-1)\sigma^4) + \frac{\mu'_4 + 3(n-1)\sigma^4}{n} \right) \\ &= \frac{(n^2 - 2n + 1)\mu'_4 + (n^2 - 2n + 3)(n-1)\sigma^4}{n(n-1)^2} \\ &= \frac{(n-1)^2}{(n-1)^2} \frac{\mu'_4}{n} + \frac{n^2 - 2n + 3}{n(n-1)} \sigma^4 = \frac{\mu'_4}{n} + \frac{n^2 - 2n + 3}{n(n-1)} \sigma^4\end{aligned}$$

und deshalb

$$\text{Var}_\theta S_n^2 = \frac{\mu'_4}{n} + \frac{n^2 - 2n + 3 - n^2 + n}{n(n-1)} \sigma^4 = \frac{\mu'_4}{n} - \frac{n-3}{n(n-1)} \sigma^4 = \frac{1}{n} \left(\mu'_4 - \frac{n-3}{n-1} \sigma^4 \right).$$

□

Satz 3.3.3 1. Der Schätzer $\tilde{S}_n^2$ für σ^2 ist erwartungstreu.

2. Es gilt $\text{Var}_\theta \tilde{S}_n^2 = 1/n(\mu'_4 - \sigma^4)$.

Beweis

1.
$$\mathbb{E}_\theta \tilde{S}_n^2 = \frac{1}{n} \sum_{i=1}^n \underbrace{\mathbb{E}_\theta (X_i - \mu)^2}_{=\text{Var}_\theta X_i} = \frac{1}{n} \sum_{i=1}^n \sigma^2 = \sigma^2.$$

2. Setzen wir wie in Satz 3.3.2 o.B.d.A. $\mu = 0$ voraus. Dann gilt

$$\begin{aligned}\text{Var}_\theta \tilde{S}_n^2 &= \mathbb{E}_\theta \left(\frac{1}{n} \sum_{i=1}^n X_i^2 \right)^2 - \left(\mathbb{E}_\theta \tilde{S}_n^2 \right)^2 = \frac{1}{n^2} \mathbb{E}_\theta \left(\sum_{i=1}^n X_i^2 \right)^2 - \sigma^4 \\ &= \frac{n\mu'_4 + n(n-1)\sigma^4}{n^2} - \sigma^4 = \frac{\mu'_4 + (n-1)\sigma^4}{n} - \sigma^4 = \frac{\mu'_4 - \sigma^4}{n}.\end{aligned}$$

I_1 Beweis S. 3.3.2

□

Folgerung 3.3.1

Der Schätzer $\tilde{S}_n^2$ für σ^2 ist besser als S_n^2 , weil beide erwartungstreu sind und

$$\text{Var}_\theta \tilde{S}_n^2 = \frac{\mu'_4 - \sigma^4}{n} < \frac{\mu'_4 - \frac{n-3}{n-1}\sigma^4}{n} = \text{Var}_\theta S_n^2.$$

Diese Eigenschaft von $\tilde{S}_n^2$ im Vergleich zu S_n^2 ist intuitiv klar, da man in $\tilde{S}_n^2$ mehr Informationen über die Verteilung der Stichprobenvariablen X_i (nämlich den bekannten Erwartungswert μ) eingesteckt hat.

Satz 3.3.4

Die Schätzer S_n^2 bzw. $\tilde{S}_n^2$ sind stark konsistent und asymptotisch normalverteilt:

$$\begin{aligned}S_n^2 &\xrightarrow[n \rightarrow \infty]{\text{f.s.}} \sigma^2, & \sqrt{n} \frac{S_n^2 - \sigma^2}{\sqrt{\mu'_4 - \sigma^4}} &\xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1), \\ \tilde{S}_n^2 &\xrightarrow[n \rightarrow \infty]{\text{f.s.}} \sigma^2, & \sqrt{n} \frac{\tilde{S}_n^2 - \sigma^2}{\sqrt{\mu'_4 - \sigma^4}} &\xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1),\end{aligned}$$

falls $\mu'_4 < \infty$.

Beweis Zeigen wir nur, dass S_n^2 die obigen Eigenschaften besitzt. Der Beweis für $\tilde{S}_n^2$ verläuft analog. Die starke Konsistenz von S_n^2 folgt aus dem starken Gesetz der großen Zahlen, nach dem

$$\frac{1}{n} \sum_{i=1}^n X_i^2 \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta X^2 \quad \text{und} \quad \bar{X}_n \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mu$$

gilt und somit auch

$$\bar{X}_n^2 \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mu^2.$$

Dann

$$\begin{aligned} S_n^2 &= \frac{1}{n-1} \left(\sum_{i=1}^n X_i^2 - n \bar{X}_n^2 \right) = \frac{n}{n-1} \left(\frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}_n^2 \right) \\ &\xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta X^2 - \mu^2 = \text{Var}_\theta X = \sigma^2, \end{aligned}$$

und die starke Konsistenz ist bewiesen. Um die asymptotische Normalverteilttheit zu beweisen, nehmen wir o.B.d.A. an, dass $\mu = \mathbb{E}_\theta X = 0$. Dann folgt mit Hilfe des Satzes von Slutsky (vgl. Sätze 6.8 - 6.9 aus dem WR-Skript)

$$\begin{aligned} \sqrt{n} \frac{S_n^2 - \sigma^2}{\sqrt{\mu'_4 - \sigma^4}} &= \sqrt{n} \frac{\frac{1}{n-1} \sum_{i=1}^n X_i^2 - \frac{n}{n-1} \bar{X}_n^2 - \sigma^2}{\sqrt{\mu'_4 - \sigma^4}} \\ &= \sqrt{n} \frac{1}{n-1} \frac{\sum_{i=1}^n X_i^2 - n\sigma^2}{\sqrt{\mu'_4 - \sigma^4}} - \underbrace{\frac{\sqrt{n} \frac{n}{n-1} \bar{X}_n^2}{\sqrt{\mu'_4 - \sigma^4}}}_{=R_n^1} - \underbrace{\left(1 - \frac{n}{n-1}\right) \frac{\sigma^2 \sqrt{n}}{\sqrt{\mu'_4 - \sigma^4}}}_{=R_n^2} \\ &\stackrel{d}{\underset{n \rightarrow \infty}{\rightsquigarrow}} \frac{\sum_{i=1}^n X_i^2 - n\sigma^2}{\sqrt{n(\mu'_4 - \sigma^4)}}, \end{aligned}$$

weil

$$R_n^2 = \left(1 - \frac{n}{n-1}\right) \sigma^2 \frac{\sqrt{n}}{\sqrt{\mu'_4 - \sigma^4}} = -\frac{\sigma^2}{\sqrt{\mu'_4 - \sigma^4}} \frac{\sqrt{n}}{n-1} \xrightarrow[n \rightarrow \infty]{\text{f.s.}} 0,$$

also auch stochastisch und in Verteilung. Es gilt

$$R_n^1 \sim \sqrt{n} \frac{\bar{X}_n^2}{\sqrt{\mu'_4 - \sigma^4}} \xrightarrow[n \rightarrow \infty]{d} 0,$$

weil

$$\mathbb{E}_\theta \left(\sqrt{n} \bar{X}_n^2 \right) \underset{\mu=0}{=} \sqrt{n} \text{Var}_\theta \bar{X}_n \underset{\text{S. 3.3.1, 4}}{=} \sqrt{n} \frac{\sigma^2}{n} = \frac{\sigma^2}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{} 0$$

und somit

$$\sqrt{n}(\bar{X}_n)^2 \xrightarrow[n \rightarrow \infty]{L^1} 0 \implies \sqrt{n}(\bar{X}_n)^2 \xrightarrow[n \rightarrow \infty]{\mathbb{P}} 0 \implies \sqrt{n}(\bar{X}_n)^2 \xrightarrow[n \rightarrow \infty]{d} 0.$$

Dann gilt

$$\lim_{n \rightarrow \infty} \sqrt{n} \frac{S_n^2 - \sigma^2}{\sqrt{\mu'_4 - \sigma^4}} \stackrel{d}{=} \lim_{n \rightarrow \infty} \frac{\sum_{i=1}^n X_i^2 - n\sigma^2}{\sqrt{n(\mu'_4 - \sigma^4)}} \stackrel{d}{=} Y \sim N(0, 1)$$

nach dem zentralen Grenzwertsatz für die Folge von unabhängigen identisch verteilten Zufallsvariablen $\{X_i^2\}_{i \in \mathbb{N}}$, weil $\mathbb{E}_\theta X_i^2 \stackrel{\mu=0}{=} \text{Var}_\theta X = \sigma^2$ und

$$\text{Var}_\theta X_i^2 = \mathbb{E}_\theta X^4 - (\mathbb{E}_\theta X^2)^2 = \mu'_4 - \sigma^4.$$

□

Folgerung 3.3.2

Es gilt

$$1. \quad \sqrt{n} \frac{\bar{X}_n - \mu}{S_n} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1)$$

und somit

$$2. \quad \mathbb{P} \left(\mu \in \left[\bar{X}_n - \frac{z_{1-\alpha/2} S_n}{\sqrt{n}}, \bar{X}_n + \frac{z_{1-\alpha/2} S_n}{\sqrt{n}} \right] \right) \xrightarrow[n \rightarrow \infty]{} 1 - \alpha \quad (3.3.1)$$

für ein $\alpha \in (0, 1)$, wobei z_α das α -Quantil der $N(0, 1)$ -Verteilung ist.

Bemerkung 3.3.3

Das Intervall in (3.3.1) nennt man *asymptotisches Konfidenz- oder Vertrauensintervall* für den Parameter μ . Falls α klein ist (z.B. $\alpha = 0,05$), so liegt μ mit einer asymptotisch großen Wahrscheinlichkeit $1 - \alpha$ im vorgegebenen Intervall. Diese Art der Schätzung von μ stellt eine Alternative zu den Punktschätzern dar und wird ausführlich in Kapitel 4 behandelt.

Beweis der Folgerung 3.3.2

1. Aus Satz 3.3.4 folgt

$$S_n^2 \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \sigma^2 \implies \frac{\sigma}{S_n} \xrightarrow[n \rightarrow \infty]{\text{f.s.}} 1 \implies \frac{\sigma}{S_n} \xrightarrow[n \rightarrow \infty]{d} 1$$

und somit nach der Verwendung des Satzes von Slutsky

$$\sqrt{n} \frac{\bar{X}_n - \mu}{S_n} = \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \cdot \frac{\sigma}{S_n} \xrightarrow[n \rightarrow \infty]{d} Y \cdot 1 = Y \sim N(0, 1),$$

wobei wir die asymptotische Normalverteiltheit von $\bar{X}_n$ benutzt haben.

2. Aus 1) folgt

$$\mathbb{P}_\theta \left(\sqrt{n} \frac{\bar{X}_n - \mu}{S_n} \in [z_{\alpha/2}, z_{1-\alpha/2}] \right) \xrightarrow[n \rightarrow \infty]{} \Phi(z_{1-\alpha/2}) - \Phi(z_{\alpha/2}) = 1 - \frac{\alpha}{2} - \frac{\alpha}{2} = 1 - \alpha.$$

Daraus folgt das Intervall (3.3.1) nach der Auflösung der Ungleichung

$$z_{\alpha/2} \leq \sqrt{n} \frac{\bar{X}_n - \mu}{S_n} \leq z_{1-\alpha/2}$$

bzgl. μ .

□

Betrachten wir weiterhin den wichtigen Spezialfall der normalverteilten Stichprobenvariablen X_i , $i = 1, \dots, n$, also $X \sim N(\mu, \sigma^2)$.

Satz 3.3.5

Falls $X_1, \dots, X_n$ normalverteilt sind mit Parametern μ und σ^2 , dann gilt

1. $\frac{(n-1)S_n^2}{\sigma^2} \sim \chi_{n-1}^2$,
2. $\frac{n\tilde{S}_n^2}{\sigma^2} \sim \chi_n^2$.

Beweis Beweisen wir den schwierigeren Fall 1, der Beweis im Fall 2 verläuft analog.

Da $X_i \sim N(\mu, \sigma^2)$, gilt, dass $\frac{X_i - \mu}{\sigma} \sim N(0, 1)$ unabhängige identisch verteilte Zufallsvariablen für $i = 1, \dots, n$ sind. Nach Lemma 2.2.1 gilt

$$\sum_{i=1}^n (X_i - \mu)^2 = \sum_{i=1}^n (X_i - \bar{X}_n)^2 + n(\bar{X}_n - \mu)^2$$

und somit

$$T_1 = \underbrace{\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma} \right)^2}_{\sim \chi_n^2} = \frac{n-1}{\sigma^2} S_n^2 + \underbrace{\left(\frac{\sqrt{n}(\bar{X}_n - \mu)}{\sigma} \right)^2}_{= T_2 \sim \chi_1^2 \text{ aus S. 3.3.1, 4)}.$$

In Lemma 3.3.2 wird bewiesen, dass S_n^2 und $\bar{X}_n$ unabhängig sind. Somit gilt

$$\varphi_{T_1}(s) = \varphi_{\frac{n-1}{\sigma^2} S_n^2}(s) \cdot \varphi_{T_2}(s), \quad \forall s \in \mathbb{R},$$

wobei $\varphi_Z(s)$ die charakteristische Funktion einer Zufallsvariablen Z ist. Da nach dem Satz 3.2.1

$$\varphi_{T_1}(s) = \frac{1}{(1 - 2is)^{n/2}}, \quad \varphi_{T_2}(s) = \frac{1}{(1 - 2is)^{1/2}},$$

folgt

$$\varphi_{\frac{n-1}{\sigma^2} S_n^2}(s) = \frac{\varphi_{T_1}(s)}{\varphi_{T_2}(s)} = \frac{1}{(1 - 2is)^{(n-1)/2}} = \varphi_{\chi_{n-1}^2}(s).$$

Aus dem Satz 3.2.1 und dem Eindeutigkeitssatz für charakteristische Funktionen (vgl. Folgerung 5.1 aus dem WR-Skript) folgt

$$\frac{n-1}{\sigma^2} S_n^2 \sim \chi_{n-1}^2.$$

□

Lemma 3.3.2

Falls $X \sim N(\mu, \sigma^2)$, $X_1, \dots, X_n$ unabhängige identisch verteilte Zufallsvariablen, $X_i \stackrel{d}{=} X$, dann sind $\bar{X}_n$ und S_n^2 unabhängig.

Dieses Lemma wird unter Anderem gebraucht, um folgendes Ergebnis zu beweisen:

Satz 3.3.6

Unter den Voraussetzungen von Lemma 3.3.2 gilt

$$\frac{\sqrt{n}(\bar{X}_n - \mu)}{S_n} \sim t_{n-1}.$$

Beweis von Lemma 3.3.2 Es folgt aus Lemma 2.2.1, dass

$$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X'_i + \mu \quad \text{und} \quad S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X'_i - \bar{X}_n)^2$$

für $X'_i = X_i - \mu$, $i = 1, \dots, n$. Somit kann wegen des Satzes 3.2.5 o.B.d.A. $\mu = 0$ und $\sigma^2 = 1$ angenommen werden. Um die Unabhängigkeit von $\bar{X}_n$ und S_n^2 zu zeigen, stellen wir S_n^2 in alternativer Form dar:

$$S_n^2 = \frac{1}{n-1} \left((X_1 - \bar{X}_n)^2 + \sum_{i=2}^n (X_i - \bar{X}_n)^2 \right) = \frac{1}{n-1} \left(\left(\sum_{i=2}^n (X_i - \bar{X}_n) \right)^2 + \sum_{i=2}^n (X_i - \bar{X}_n)^2 \right),$$

weil $\sum_{i=1}^n (X_i - \bar{X}_n) = 0$ nach Abschnitt 2.2.1. Somit gilt

$$S_n^2 = \tilde{\varphi}(X_2 - \bar{X}_n, \dots, X_n - \bar{X}_n),$$

wobei

$$\tilde{\varphi}(x_2, \dots, x_n) = \frac{1}{n-1} \left(\left(\sum_{i=2}^n x_i \right)^2 + \sum_{i=2}^n x_i^2 \right), \quad (x_2, \dots, x_n) \in \mathbb{R}^{n-1}.$$

Es genügt (nach Satz 3.2.5) zu zeigen, dass der Zufallsvektor $(X_2 - \bar{X}_n, \dots, X_n - \bar{X}_n)$ unabhängig von $\bar{X}_n$ ist. Sei $X = (X_1, \dots, X_n)^T$, X_i unabhängige identisch verteilte Zufallsvariablen mit $X_i \sim N(0, 1)$ nach unserer Annahme. Dann gilt

$$f_X(x_1, \dots, x_n) = \frac{1}{(2\pi)^{n/2}} \exp \left(-\frac{1}{2} \sum_{i=1}^n x_i^2 \right), \quad (x_1, \dots, x_n) \in \mathbb{R}^n$$

für die Dichte von X . Sei $\varphi = (\varphi_1, \dots, \varphi_n) : \mathbb{R}^n \rightarrow \mathbb{R}^n$ die lineare Abbildung mit

$$\begin{cases} \varphi_1(x) = \bar{x}_n, \\ \varphi_2(x) = x_2 - \bar{x}_n, \\ \vdots \\ \varphi_n(x) = x_n - \bar{x}_n, \end{cases} \quad x = (x_1, \dots, x_n) \in \mathbb{R}^n.$$

Um die Umkehrabbildung $\varphi^{-1} : (y_1, \dots, y_n) \mapsto (x_1, \dots, x_n)$ zu finden, setzen wir $y_i = \varphi_i(x)$, $i = 1, \dots, n$ und schreiben

$$\begin{cases} y_1 = \bar{x}_n \\ y_2 = x_2 - \bar{x}_n = x_2 - y_1 \\ \vdots \\ y_n = x_n - y_1 \end{cases}, \text{ woraus } \begin{cases} x_2 = y_1 + y_2 \\ \vdots \\ x_n = y_1 + y_n \\ x_2 + \dots + x_n = (n-1)y_1 + y_2 + \dots + y_n \\ x_1 + \dots + x_n = ny_1 = x_1 + (n-1)y_1 + y_2 + \dots + y_n \end{cases}$$

folgt und somit $x_1 = y_1 - \sum_{i=2}^n y_i$. Es gilt insgesamt

$$\begin{cases} \varphi_1^{-1}(y) = y_1 - \sum_{i=2}^n y_i, \\ \varphi_2^{-1}(y) = y_1 + y_2, \\ \vdots \\ \varphi_n^{-1}(y) = y_1 + y_n, \end{cases} \quad y = (y_1, \dots, y_n)^T \in \mathbb{R}^n.$$

Um den Dichtetransformationssatz 3.2.6 für $\varphi(X)$ zu verwenden, brauchen wir die Determinante der Jacobi-Matrix

$$J = \det \left(\frac{\partial \varphi_i^{-1}}{\partial y_j} \right)_{i,j=1,\dots,n} = \left| \begin{pmatrix} 1 & -1 & -1 & -1 & \dots & -1 \\ 1 & 1 & 0 & 0 & \dots & 0 \\ 1 & 0 & 1 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \ddots & \ddots & \vdots \\ 1 & 0 & \dots & 0 & 1 & 0 \\ 1 & 0 & \dots & \dots & 0 & 1 \end{pmatrix} \right|$$

$$= 1 \cdot 1 - (-1) \cdot 1 + (-1) \cdot (-1) - (-1) \cdot 1 + \dots = \underbrace{1 + \dots + 1}_n = n.$$

Somit gilt für die Dichte von $Y = \varphi(X) = (\bar{X}_n, X_2 - \bar{X}_n, \dots, X_n - \bar{X}_n)$

$$\begin{aligned} f_{\varphi(Y)}(y_1, \dots, y_n) &= f_X(\varphi^{-1}(y)) \cdot |J| = \frac{n}{(2\pi)^{n/2}} \exp \left\{ -\frac{1}{2} \left(y_1 - \sum_{i=2}^n y_i \right)^2 - \frac{1}{2} \sum_{i=2}^n (y_1 + y_i)^2 \right\} \\ &= \frac{n}{(2\pi)^{n/2}} \exp \left\{ -\frac{1}{2} \left(y_1^2 - 2y_1 \sum_{i=2}^n y_i + \left(\sum_{i=2}^n y_i \right)^2 + \sum_{i=2}^n y_i^2 + 2y_1 \sum_{i=2}^n y_i + (n-1)y_1^2 \right) \right\} \\ &= \frac{n}{(2\pi)^{n/2}} \exp \left\{ -\frac{1}{2} \left(ny_1^2 + \left(\sum_{i=2}^n y_i \right)^2 + \sum_{i=2}^n y_i^2 \right) \right\} \\ &= \underbrace{\left(\frac{n}{2\pi} \right)^{1/2} \exp \left\{ -\frac{1}{2} ny_1^2 \right\}}_{=f_{\varphi_1(X)}(y_1)} \cdot \underbrace{\left(\frac{n}{(2\pi)^{n-1}} \right)^{1/2} \exp \left\{ -\frac{1}{2} \left(\sum_{i=2}^n y_i^2 + \left(\sum_{i=2}^n y_i \right)^2 \right) \right\}}_{f_{(\varphi_2(X), \dots, \varphi_n(X))}(y_2, \dots, y_n)}, \end{aligned}$$

woraus die Unabhängigkeit von

$$\varphi_1(X) = \bar{X}_n \sim N \left(\mu, \frac{\sigma^2}{n} \right) \underset{\mu=0, \sigma^2=1}{=} N \left(0, \frac{1}{n} \right)$$

und

$$(\varphi_2(X), \dots, \varphi_n(X)) = (X_2 - \bar{X}_n, \dots, X_n - \bar{X}_n)$$

folgt. Somit sind auch $\bar{X}_n$ und $S_n^2 = \tilde{\varphi}(X_2 - \bar{X}_n, \dots, X_n - \bar{X}_n)$ unabhängig. $\square$

Beweis des Satzes 3.3.6 Aus den Sätzen 3.3.1, 4) und 3.3.5 folgt

$$\bar{X}_n \sim N \left(\mu, \frac{\sigma^2}{n} \right) \quad \text{und} \quad \frac{(n-1)S_n^2}{\sigma^2} \sim \chi_{n-1}^2,$$

also

$$Y_1 = \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \sim N(0, 1) \quad \text{und} \quad Y_2 = \frac{(n-1)S_n^2}{\sigma^2} \sim \chi_{n-1}^2.$$

Nach dem Lemma 3.3.2 und Satz 3.2.5 sind Y_1 und Y_2 unabhängig. Dann gilt

$$T = \frac{Y_1}{\sqrt{\frac{Y_2}{n-1}}} \sim t_{n-1}$$

nach der Definition einer t -Verteilung, wobei

$$T = \frac{\sqrt{n} \frac{\bar{X}_n - \mu}{\sigma}}{\sqrt{\frac{(n-1)S_n^2}{\sigma^2(n-1)}}} = \sqrt{n} \frac{\bar{X}_n - \mu}{S_n}.$$

Somit gilt

$$\sqrt{n} \frac{\bar{X}_n - \mu}{S_n} \sim t_{n-1}.$$

□

Bemerkung 3.3.4

Mit Hilfe des Satzes 3.3.6 kann folgendes Konfidenzintervall für den Erwartungswert μ einer normalverteilten Stichprobe $(X_1, \dots, X_n)$ bei unbekannter Varianz σ^2 ($X_i \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$) konstruiert werden:

$$\mathbb{P} \left(\mu \in \left[\bar{X}_n - \frac{t_{n-1, 1-\alpha/2}}{\sqrt{n}} S_n, \bar{X}_n + \frac{t_{n-1, 1-\alpha/2}}{\sqrt{n}} S_n \right] \right) = 1 - \alpha$$

für $\alpha \in (0, 1)$, denn

$$\begin{aligned} \mathbb{P} \left(\sqrt{n} \frac{\bar{X}_n - \mu}{S_n} \in \left[\underbrace{t_{n-1, \alpha/2}}_{=-t_{n-1, 1-\alpha/2} \text{ wg. Sym. } t\text{-Vert.}}, t_{n-1, 1-\alpha/2} \right] \right) &= F_{t_{n-1}}(t_{n-1, 1-\alpha/2}) - F_{t_{n-1}}(t_{n-1, \alpha/2}) \\ &= 1 - \frac{\alpha}{2} - \frac{\alpha}{2} = 1 - \alpha, \end{aligned} \quad (3.3.2)$$

wobei $t_{n-1, \alpha}$ das α -Quantil der t_{n-1} -Verteilung darstellt. Der Rest folgt aus (3.3.2) durch das Auflösen bzgl. μ .

3.3.4 Eigenschaften der Ordnungsstatistiken

In Abschnitt 2.2.2 haben wir bereits die Ordnungsstatistiken $x_{(1)}, \dots, x_{(n)}$ einer konkreten Stichprobe $(x_1, \dots, x_n)$ betrachtet. Wenn wir nun auf der Modellebene arbeiten, also eine Zufallsstichprobe $(X_1, \dots, X_n)$ von unabhängigen identisch verteilten Zufallsvariablen X_i mit Verteilungsfunktion $F(x)$ haben, welche Eigenschaften haben dann ihre Ordnungsstatistiken

$$X_{(1)}, \dots, X_{(n)}?$$

Satz 3.3.7

1. Die Verteilungsfunktion der Ordnungsstatistik $X_{(i)}$, $i = 1, \dots, n$ ist gegeben durch

$$F_{X_{(i)}}(x) = \sum_{k=i}^n \binom{n}{k} F^k(x) (1 - F(x))^{n-k}, \quad x \in \mathbb{R}. \quad (3.3.3)$$

2. Falls X_i eine diskrete Verteilung mit Wertebereich $E = \{\dots, a_{j-1}, a_j, a_{j+1}, \dots\}$ haben, $i = 1, \dots, n$, $a_i < a_j$ für $i < j$, dann gilt für die Zähldichte von $X_{(i)}$, $i = 1, \dots, n$:

$$\mathbb{P}(X_{(i)} = a_j) = \sum_{k=i}^n \binom{n}{k} \left(F^k(a_j) (1 - F(a_j))^{n-k} - F^k(a_{j-1}) (1 - F(a_{j-1}))^{n-k} \right),$$

wobei

$$F(a_j) = \sum_{a_k \in E, k \leq j} \mathbb{P}(X_i = a_k).$$

3. Falls X_i absolut stetig verteilt sind mit Dichte f , die stückweise stetig ist, dann ist auch $X_{(i)}$, $i = 1, \dots, n$ absolut stetig verteilt mit der Dichte

$$f_{X_{(i)}}(x) = \frac{n!}{(i-1)!(n-i)!} f(x) F^{i-1}(x) (1-F(x))^{n-i}, \quad x \in \mathbb{R}.$$

Beweis

1. Führen wir die Zufallsvariable

$$Y = \#\{i : X_i \leq x\} = \sum_{i=1}^n \mathbb{I}(X_i \leq x), \quad x \in \mathbb{R}$$

ein. Da $X_1, \dots, X_n$ unabhängig identisch verteilt mit Verteilungsfunktion F sind, gilt $Y \sim \text{Bin}(n, F(x))$. Weiterhin gilt

$$F_{X_{(i)}}(x) = \mathbb{P}(X_{(i)} \leq x) = \mathbb{P}(Y \geq i) = \sum_{k=i}^n \binom{n}{k} F^k(x) (1-F(x))^{n-k}, \quad x \in \mathbb{R}.$$

2. folgt aus 1) durch

$$\mathbb{P}(X_{(i)} = a_j) = \mathbb{P}(a_{j-1} < X_{(i)} \leq a_j) = F_{X_{(i)}}(a_j) - F_{X_{(i)}}(a_{j-1}) \quad \forall j, i.$$

3. Beweisen Sie 3) als Übungsaufgabe.

□

Bemerkung 3.3.5

1. Für $i = 1$ und $i = n$ sieht die Formel (3.3.3) besonders einfach aus:

$$\begin{aligned} F_{X_{(1)}}(x) &= 1 - (1 - F(x))^n, \quad x \in \mathbb{R} \\ F_{X_{(n)}}(x) &= F^n(x), \quad x \in \mathbb{R}. \end{aligned}$$

Diese Formeln lassen sich auch direkt herleiten:

$$\begin{aligned} F_{X_{(1)}}(x) &= \mathbb{P}\left(\min_{i=1, \dots, n} X_i \leq x\right) = 1 - \mathbb{P}\left(\min_{i=1, \dots, n} X_i > x\right) = 1 - \mathbb{P}(X_i \geq x, \quad \forall i = 1, \dots, n) \\ &= 1 - \prod_{i=1}^n \mathbb{P}(X_i > x) = 1 - (1 - F(x))^n, \\ F_{X_{(n)}}(x) &= \mathbb{P}\left(\max_{i=1, \dots, n} X_i \leq x\right) = \mathbb{P}(X_i \leq x, \quad \forall i = 1, \dots, n) \\ &= \prod_{i=1}^n \mathbb{P}(X_i \leq x) = F^n(x), \quad x \in \mathbb{R}. \end{aligned}$$

2. Falls X_i absolut stetig verteilt sind mit einer stückweise stetigen Dichte f , so lassen sich Formeln für die gemeinsame Dichte der Verteilung von $(X_{(i_1)}, \dots, X_{(i_k)})$, $i \leq k \leq n$ herleiten. Insbesondere gilt für $k = n$

$$f_{(X_{(1)}, \dots, X_{(n)})}(x_1, \dots, x_n) = \begin{cases} n! \cdot f(x_1) \cdot \dots \cdot f(x_n), & \text{falls } -\infty < x_1 < \dots < x_n < \infty, \\ 0, & \text{sonst.} \end{cases}$$

Übungsaufgabe 3.3.1

Zeigen Sie für $X_1, \dots, X_n$ unabhängig identisch verteilt, $X_i \sim U[0, \theta]$, $\theta > 0$, $i = 1, \dots, n$, dass

1. die Dichte von $X_{(i)}$ gleich

$$f_{X_{(i)}}(x) = \begin{cases} \frac{n!}{(i-1)!(n-i)!} \theta^{-n} x^{i-1} (\theta - x)^{n-i}, & x \in (0, \theta) \\ 0, & \text{sonst} \end{cases}$$

und

- 2.

$$\mathbb{E}X_{(i)}^k = \frac{\theta^k n! (i+k-1)!}{(n+k)!(i-1)!}, \quad k \in \mathbb{N}, \quad i = 1, \dots, n$$

sind. Insbesondere gilt $\mathbb{E}X_{(i)} = \frac{i}{n+1}\theta$ und $\text{Var} X_{(i)} = \frac{i(n-i+1)\theta^2}{(n+1)^2(n+2)}$.

3.3.5 Empirische Verteilungsfunktion

Im Folgenden betrachten wir die statistischen Eigenschaften der in Abschnitt 2.1.2 eingeführten empirischen Verteilungsfunktion $\hat{F}_n(x)$ einer Zufallsstichprobe $(X_1, \dots, X_n)$, wobei $X_i \stackrel{d}{=} X$ unabhängige identisch verteilte Zufallsvariablen mit Verteilungsfunktion $F(\cdot)$ sind.

Satz 3.3.8

Es gilt

1. $n\hat{F}_n(x) \sim \text{Bin}(n, F(x))$, $x \in \mathbb{R}$.
2. $\hat{F}_n(x)$ ist ein erwartungstreuer Schätzer für $F(x)$, $x \in \mathbb{R}$ mit

$$\text{Var} \hat{F}_n(x) = \frac{F(x)(1-F(x))}{n}.$$

3. $\hat{F}_n(x)$ ist stark konsistent.
4. $\hat{F}_n(x)$ ist asymptotisch normalverteilt:

$$\sqrt{n} \frac{\hat{F}_n(x) - F(x)}{\sqrt{F(x)(1-F(x))}} \xrightarrow{d} Y \sim N(0, 1), \quad \forall x : F(x) \in (0, 1).$$

Beweis 1. folgt aus der Darstellung

$$\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(X_i \leq x), \quad x \in \mathbb{R},$$

weil $\mathbb{I}(X_i \leq x) \sim \text{Bernoulli}(F(x))$, $\forall i = 1, \dots, n$. Somit ist

$$\sum_{i=1}^n \mathbb{I}(X_i \leq x) \sim \text{Bin}(n, F(x)).$$

2. Es folgt aus 1)

$$\begin{cases} \mathbb{E}(n\hat{F}_n(x)) = nF(x), & x \in \mathbb{R}, \\ \text{Var}(n\hat{F}_n(x)) = nF(x) \cdot (1 - F(x)), & x \in \mathbb{R}, \end{cases}$$

woraus $\mathbb{E}\hat{F}_n(x) = F(x)$ und $\text{Var}\hat{F}_n(x) = F(x)(1 - F(x))/n$ folgen.

3. Da $Y_i = \mathbb{I}(X_i \leq x)$, $i = 1, \dots, n$, $x \in \mathbb{R}$ unabhängige identisch verteilte Zufallsvariablen sind, gilt nach dem starken Gesetz der großen Zahlen

$$\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n Y_i \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}Y_i = F(x).$$

4. folgt aus der Anwendung des zentralen Grenzwertsatzes auf die oben genannte Folge $\{Y_i\}_{i \in \mathbb{N}}$. □

In Satz 3.3.8, 3) wird behauptet, dass

$$\hat{F}_n(x) \xrightarrow[n \rightarrow \infty]{\text{f.s.}} F(x), \quad \forall x \in \mathbb{R}.$$

Der nachfolgende Satz von Gliwenko-Cantelli behauptet, dass diese Konvergenz gleichmäßig in $x \in \mathbb{R}$ stattfindet. Um diesen Satz formulieren zu können, betrachten wir den *gleichmäßigen Abstand* zwischen $\hat{F}_n$ und F

$$D_n = \sup_{x \in \mathbb{R}} |\hat{F}_n(x) - F(x)|.$$

Dieser Abstand ist eine Zufallsvariable, die auch *Kolmogorow-Abstand* genannt wird. Er gibt den maximalen Fehler an, den man bei der Schätzung von $F(x)$ durch $\hat{F}_n(x)$ macht.

Übungsaufgabe 3.3.2

Zeigen Sie, dass

$$D_n = \max_{i \in \{1, \dots, n\}} \max \left\{ F(X_{(i)} - 0) - \frac{i-1}{n}, \frac{i}{n} - F(X_{(i)}) \right\}. \quad (3.3.4)$$

Beachten Sie dabei die Tatsache, dass $\hat{F}_n(x)$ eine Treppenfunktion mit Sprungstellen $X_{(i)}$, $i = 1, \dots, n$ ist.

Satz 3.3.9 (Gliwenko-Cantelli):

Es gilt $D_n \xrightarrow[n \rightarrow \infty]{\text{f.s.}} 0$.

Beweis Für alle $m \in \mathbb{N}$ wähle beliebige Zahlen $-\infty = z_0 < z_1 < \dots < z_{m-1} < z_m = \infty$. Dann gilt

$$D_n = \sup_{z \in \mathbb{R}} |\hat{F}_n(z) - F(z)| = \sup_{j=0, \dots, m-1} \sup_{z \in [z_j, z_{j+1})} |\hat{F}_n(z) - F(z)|.$$

Zeigen wir, dass $\forall m \in \mathbb{N}$ $z_0, \dots, z_m$ existieren, für die gilt

$$F(z_{j+1} - 0) - F(z_j) \leq \varepsilon = \frac{1}{m}. \quad (3.3.5)$$

Falls F stetig ist, genügt es, $z_j = F^{-1}(j/m)$, $j = 1 \dots m-1$ gleichzusetzen. Im allgemeinen Fall existieren $n < m/2$ Punkte x_j mit der Eigenschaft

$$F(x_j) - F(x_j - 0) > 2\varepsilon = 2/m$$

(weil $n \cdot 2\varepsilon \leq 1$ sein muss) und $k+1$ Punkte y_j zwischen diesen Punkten x_j mit Eigenschaft (3.3.5), wobei für k gilt:

$$n \cdot 2\varepsilon + (k+1)\varepsilon \leq 1 \implies 2n + k + 1 \leq m \implies k \leq m - 2n - 1.$$

Setzen wir $\{z_j\} = \{x_j\} \cup \{y_j\}$. Für alle $z \in [z_j, z_{j+1})$ gilt

$$\hat{F}_n(z) - F(z) \leq \hat{F}_n(z_{j+1} - 0) - F(z_j) \leq \hat{F}_n(z_{j+1} - 0) - F(z_{j+1} - 0) + \varepsilon,$$

weil aus (3.3.5) folgt, dass $-F(z_j) \leq \varepsilon - F(z_{j+1} - 0)$, $\forall j$.

Genauso gilt

$$\hat{F}_n(z) - F(z) \geq \hat{F}_n(z_j) - F(z_{j+1} - 0) \geq \hat{F}_n(z_j) - F(z_j) - \varepsilon,$$

weil aus (3.3.5) für alle j folgt, dass $-F(z_{j+1} - 0) \geq -F(z_j) - \varepsilon$ gilt. Für alle $m \in \mathbb{N}$, $j \in \{0, 1, \dots, m\}$ sei

$$A_{m,j} = \{\omega \in \Omega : \lim_{n \rightarrow \infty} \hat{F}_n(z_j) = F(z_j)\},$$

$$A'_{m,j} = \{\omega \in \Omega : \lim_{n \rightarrow \infty} \hat{F}_n(z_j - 0) = F(z_j - 0)\}.$$

Nach dem Satz 3.3.8, 3) gilt $\mathbb{P}(A_{m,j}) = 1$. Um $\mathbb{P}(A'_{m,j}) = 1$ zu zeigen, kann man die Verallgemeinerung von Aussage 3.3.8, 3) auf das Maß

$$\hat{F}_n(B) := \frac{1}{n} \sum_{i=1}^n \mathbb{I}(X_i \in B), \quad B \in \mathcal{B}_{\mathbb{R}}$$

benutzen: nach dem starken Gesetz der großen Zahlen gilt nämlich

$$\hat{F}_n(B) \xrightarrow[n \rightarrow \infty]{\text{f.s.}} F(B) = \mathbb{P}(X \in B), \quad B \in \mathcal{B}_{\mathbb{R}}.$$

Da $(-\infty, z_j) \in \mathcal{B}_{\mathbb{R}} \quad \forall j$, ist $\mathbb{P}(A'_{m,j}) = 1$ bewiesen $\forall m \forall j$. Für

$$A'_m = \bigcap_{j=0}^m (A_{m,j} \cap A'_{m,j})$$

gilt $\mathbb{P}(A'_m) = 1 \quad \forall m$, weil

$$\mathbb{P}(A'_m) = 1 - \mathbb{P}(\bar{A}'_m) = 1 - \mathbb{P}\left(\bigcup_{j=0}^m (\bar{A}_{m,j} \cup \bar{A}'_{m,j})\right) \geq 1 - \sum_{j=0}^m \left(\underbrace{\mathbb{P}(\bar{A}_{m,j})}_{=0} + \underbrace{\mathbb{P}(\bar{A}'_{m,j})}_{=0}\right) = 1.$$

Weiterhin: für $\varepsilon = 1/m \ \forall \omega \in A'_m \ \exists n(\omega, m) : \forall n > n(\omega, m) \ \forall j \in \{0, \dots, m-1\} \ \forall z \in [z_j, z_{j+1})$

$$\left. \begin{aligned} \hat{F}_n(z) - F(z) &\leq \underbrace{\hat{F}_n(z_{j+1} - 0) - F(z_{j+1} - 0)}_{< \varepsilon \text{ aus } A'_{m,j}} + \varepsilon < 2\varepsilon, \\ \hat{F}_n(z) - F(z) &\geq \underbrace{\hat{F}_n(z_j) - F(z_j)}_{> -\varepsilon \text{ aus } A_{m,j}} - \varepsilon > -2\varepsilon, \end{aligned} \right\} \implies |\hat{F}_n(z) - F(z)| < 2\varepsilon.$$

$$\implies D_n = \sup_{j=0, \dots, m-1} \sup_{z \in [z_j, z_{j+1})} |\hat{F}_n(z) - F(z)| < 2\varepsilon.$$

Nun wählen wir ein beliebiges $m \in \mathbb{N}$ und betrachten $A' = \bigcap_{m=1}^{\infty} A'_m$. Es folgt, dass $\mathbb{P}(A') = 1$ und $\forall \omega \in A' \ \exists n_0 : \forall n \geq n_0$

$$D_n < 2\varepsilon = \frac{2}{m} \quad \forall m \in \mathbb{N} \quad \implies \quad D_n \xrightarrow[n \rightarrow \infty]{\text{f.s.}} 0.$$

□

Satz 3.3.10 (Ungleichung von Dvoretzky-Kiefer-Wolfowitz):

Seien $X_1, \dots, X_n$ unabhängige identisch verteilte Zufallsvariablen mit Verteilungsfunktion F . Für alle $\varepsilon > 0$ gilt

$$\mathbb{P}(D_n > \varepsilon) \leq 2e^{-2n\varepsilon^2}.$$

(ohne Beweis)

Folgerung 3.3.3 (Konfidenzband für F):

Führen wir Statistiken

$$L(x) = \max\{\hat{F}_n(x) - \varepsilon_n, 0\} \text{ und } U(x) = \min\{\hat{F}_n(x) + \varepsilon_n, 1\}, \quad \varepsilon_n = \sqrt{\frac{1}{2n} \log\left(\frac{2}{\alpha}\right)}, \quad \alpha \in (0, 1)$$

ein. Dann gilt

$$\mathbb{P}(L(x) \leq F(x) \leq U(x) \quad \forall x \in \mathbb{R}) \geq 1 - \alpha \quad (3.3.6)$$

Beweis Beweisen Sie dieses Korollar als Übungsaufgabe! □

Bemerkung 3.3.6

Das simultane Konfidenzintervall $\{L(x) \leq F(x) \leq U(x), \quad x \in \mathbb{R}\}$ aus (3.3.6) heißt *Konfidenzband* für F zum Konfidenzniveau $1 - \alpha$ (vgl. Abb. 3.5).

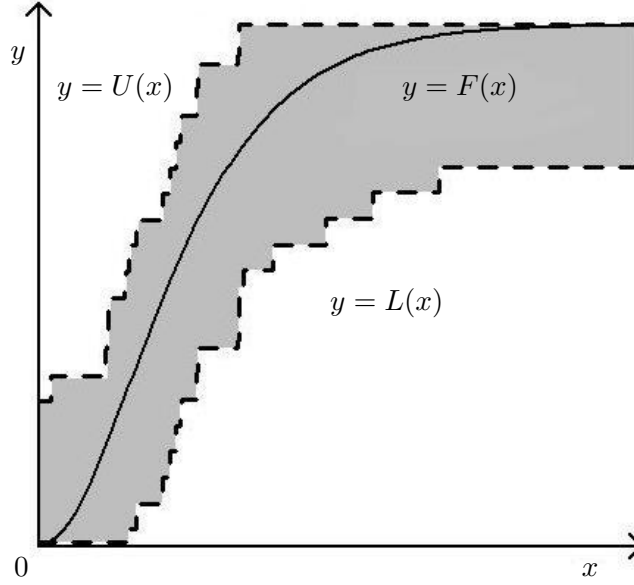
Falls die Verteilungsfunktion F stetig ist, kann man zeigen, dass die Zufallsvariable D_n nicht von F abhängt, also *verteilungsfrei* ist.

Satz 3.3.11

Für jede stetige Verteilungsfunktion F gilt

$$D_n \stackrel{d}{=} \sup_{y \in [0,1]} |\hat{G}_n(y) - y|, \text{ wobei } \hat{G}_n(y) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(Y_i \leq y), \quad y \in \mathbb{R}$$

die empirische Verteilungsfunktion der Zufallsstichprobe $(Y_1, \dots, Y_n)$ mit unabhängigen identisch verteilten Zufallsvariablen $Y_i \sim U[0, 1]$, $i = 1, \dots, n$ ist.

Abb. 3.5: Konfidenzband für F .

Beweis Zunächst definieren wir einen sogenannten *Konstanzbereich* $(a, b] \subset \mathbb{R}$ einer Verteilungsfunktion F als maximales Intervall mit der Eigenschaft $F(a) = F(b)$. Sei B die Vereinigung aller Konstanzbereiche von F . Auf B^C ist F eine monoton steigende eineindeutige Funktion. Damit folgt die Existenz ihrer Inversen $F^{-1} : (0, 1) \rightarrow B^C$. Gleichzeitig gilt

$$D_n = \sup_{x \in B^C} |\hat{F}_n(x) - F(x)|.$$

Führen wir $Y_i = F(X_i)$, $i = 1, \dots, n$ ein. Y_i sind unabhängig identisch verteilt und $Y_i \sim U[0, 1]$, denn

$$\mathbb{P}(Y_i \leq y) = \mathbb{P}(F(X_i) \leq y) = \mathbb{P}(X_i \leq F^{-1}(y)) = F(F^{-1}(y)) = y, \quad y \in (0, 1).$$

Somit gilt auch

$$\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(X_i \leq x) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(\underbrace{F(X_i)}_{Y_i} \leq F(x)) = \hat{G}_n(F(x)), \quad x \in B^C.$$

Hieraus folgt

$$\begin{aligned} D_n &= \sup_{x \in B^C} |\hat{F}_n(x) - F(x)| = \sup_{x \in B^C} |\hat{G}_n(F(x)) - F(x)| = \sup_{x \in \mathbb{R}} |\hat{G}_n(F(x)) - F(x)| \\ &= \sup_{y \in [0, 1]} |\hat{G}_n(y) - y|, \end{aligned}$$

wobei die letzte Gleichheit die Stetigkeit von F ausnützt. □

Folgerung 3.3.4

Falls F eine stetige Verteilungsfunktion ist, dann gilt

$$D_n \stackrel{d}{=} \max_{i=1, \dots, n} \max \left\{ Y_{(i)} - \frac{i-1}{n}, \frac{i}{n} - Y_{(i)} \right\},$$

wobei $Y_{(1)}, \dots, Y_{(n)}$ die Ordnungsstatistiken der auf $[0, 1]$ gleichverteilten Stichprobenvariablen $Y_1, \dots, Y_n$ sind.

Beweis Benutze dazu die Darstellung (3.3.4), den Satz 3.3.11 sowie die Tatsache, dass

$$F(x) = x, \quad x \in [0, 1]$$

für die Verteilungsfunktion der $U[0, 1]$ -Verteilung ist. $\square$

Folgende Ergebnisse werden ohne Beweis angegeben:

Bemerkung 3.3.7

1. Für die Zwecke des statistischen Testens (vgl. den Anpassungstest von Kolmogorow-Smirnow, Bemerkung 3.3.8, 3)) ist es notwendig, die Quantile der Verteilung von D_n zu nennen. Auf Grund der Komplexität der Verteilung von D_n ist es jedoch unmöglich, sie explizit anzugeben. Mit Hilfe des Satzes 3.3.11 ist es möglich, diese Quantile durch Monte-Carlo-Simulationen numerisch zu berechnen. Dazu simuliert man mehrere Stichproben $(Y_1, \dots, Y_n)$ von $U[0, 1]$ -verteilten Pseudozufallszahlen, bildet $\hat{G}_n(x)$ und berechnet D_n nach Folgerung 3.3.4.
2. Für stetige Verteilungsfunktionen F kann folgende Integraldarstellung von Verteilungsfunktion von D_n bewiesen werden:

$$\mathbb{P} \left(D_n \leq x + \frac{1}{2n} \right) = \begin{cases} 0, & x \leq 0, \\ \int_{\frac{1}{2n}-x}^{\frac{1}{2n}+x} \int_{\frac{3}{2n}-x}^{\frac{3}{2n}+x} \dots \int_{\frac{2n-1}{2n}-x}^{\frac{2n-1}{2n}+x} g(y_1, \dots, y_n) dy_n \dots dy_1, & 0 < x < \frac{2n-1}{2n}, \\ 1, & x \geq \frac{2n-1}{2n}. \end{cases}$$

wobei

$$g(y_1, \dots, y_n) = \begin{cases} n!, & 0 < y_1 < \dots < y_n < 1, \\ 0, & \text{sonst} \end{cases}$$

die Dichte der Ordnungsstatistiken $(Y_{(1)}, \dots, Y_{(n)})$ von $U[0, 1]$ -verteilten Stichprobenvariablen $(Y_1, \dots, Y_n)$ sind.

Satz 3.3.12 (Kolmogorow):

Falls die Verteilungsfunktion F der unabhängigen und identisch verteilten Stichprobenvariablen X_i , $i = 1, \dots, n$ stetig ist, dann gilt

$$\sqrt{n}D_n \xrightarrow[n \rightarrow \infty]{d} Y,$$

wobei Y eine Zufallsvariable mit der Verteilungsfunktion

$$K(x) = \begin{cases} \sum_{k=-\infty}^{\infty} (-1)^k e^{-2k^2 x^2} = 1 + 2 \sum_{k=1}^{\infty} (-1)^k e^{-2k^2 x^2}, & x > 0, \\ 0, & \text{sonst} \end{cases}$$

(Kolmogorow-Verteilung) ist.

Bemerkung 3.3.8

1. Die Verteilung von Kolmogorow ist die Verteilung des Maximums einer Brownschen Brücke, denn es gilt

$$Y \stackrel{d}{=} \sup_{t \in [0,1]} |w^\circ(t)|,$$

wobei $\{w^\circ(t), \quad t \in [0, 1]\}$ ein stochastischer Prozess ist, der die *Brownsche Brücke* genannt wird. Er wird als $w^\circ(t) = w(t) - w(1)t$, $t \in [0, 1]$ definiert, wobei $\{w(t), \quad t \in [0, 1]\}$ die Brownsche Bewegung ist (für die unter anderem $w(t) \sim N(0, t)$ gilt). Der Name „Brücke“ ist der Tatsache $w^\circ(0) = w^\circ(1) = 0$ zu verdanken.

2. Aus Satz 3.3.12 folgt

$$\mathbb{P}(\sqrt{n}D_n \leq x) \underset{n \rightarrow \infty}{\approx} K(x), \quad x \in \mathbb{R}.$$

Die daraus resultierende Näherungsformel

$$\mathbb{P}(D_n \leq x) \approx K(x\sqrt{n})$$

ist ab $n > 40$ praktisch brauchbar.

3. *Kolmogorow-Smirnow-Anpassungstest*: Mit Hilfe der Aussage des Satzes 3.3.12 ist es möglich, folgenden *asymptotischen Anpassungstest von Komogorow-Smirnow* zu entwickeln. Es wird die Haupthypothese $H_0 : F = F_0$ (die unbekannte Verteilungsfunktion der Stichprobenvariablen $X_1, \dots, X_n$ ist gleich F_0) gegen die Alternative $H_1 : F \neq F_0$ getestet. Dabei wird H_0 verworfen, falls

$$\sqrt{n}D_n \notin [k_{\alpha/2}, k_{1-\alpha/2}]$$

ist, wobei

$$D_n = \sup_{x \in \mathbb{R}} |\hat{F}_n(x) - F_0(x)|$$

und k_α das α -Quantil der Kolmogorow-Verteilung ist. Somit ist die Wahrscheinlichkeit, die richtige Hypothese H_0 zu verwerfen (Wahrscheinlichkeit des *Fehlers 1. Art*) asymptotisch gleich

$$\mathbb{P}\left(\sqrt{n}D_n \notin [k_{\alpha/2}, k_{1-\alpha/2}] \mid H_0\right) \underset{n \rightarrow \infty}{\longrightarrow} 1 - K(k_{1-\alpha/2}) + K(k_{\alpha/2}) = 1 - (1 - \alpha/2) + \alpha/2 = \alpha.$$

In der Praxis wird α klein gewählt, z.B. $\alpha \approx 0,05$. Somit ist im Fall, dass H_0 stimmt, die Wahrscheinlichkeit einer Fehlentscheidung in Folge des Testens klein.

Dieser Test ist nur ein Beispiel dessen, wie der Satz von Kolmogorow in der statistischen Testtheorie verwendet wird. Die allgemeine Philosophie des Testens wird in Kapitel 5 erläutert.

Mit Hilfe von $\hat{F}_n$ lassen sich sehr viele Schätzer durch die sogenannte *Plug-in-Methode* konstruieren. Dies werden wir jetzt näher erläutern: Sei $M = \{\text{Menge aller Verteilungsfunktionen}\}$.

Definition 3.3.9

Sei ein Parameter θ der Verteilungsfunktion F als Funktional $T : M \rightarrow \mathbb{R}$ von F gegeben: $\theta = T(F)$. Dann heißt $\hat{\theta} = T(\hat{F}_n)$ der *Plug-in-Schätzer* für θ .

Definition 3.3.10

Sei F eine beliebige Verteilungsfunktion. Das Funktional $T : M \rightarrow \mathbb{R}$ heißt *linear*, falls

$$T(aF_1 + bF_2) = aT(F_1) + bT(F_2) \quad \forall a, b \in \mathbb{R}, \quad F_1, F_2 \in M.$$

Betrachten wir eine spezielle Klasse der linearen Funktionalen

$$T(F) = \int_{\mathbb{R}} r(x) dF(x),$$

wobei $r(x)$ eine beliebige stetige Funktion ist. Beispiele für solche T sind

$$\mathbb{E}X^k = \int_{\mathbb{R}} x^k dF(x), \quad k \in \mathbb{N}.$$

Lemma 3.3.3

Der Plug-in Schätzer für $\theta = \int_{\mathbb{R}} r(x) dF(x)$ ist durch

$$\hat{\theta} = \int_{\mathbb{R}} r(x) d\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n r(x_i)$$

gegeben.

Übungsaufgabe 3.3.3

Beweisen Sie Lemma 3.3.3!

Beispiel 3.3.1 (Plug-in-Schätzer):

1. $\bar{X}_n$ ist ein Plug-in Schätzer für den Erwartungswert μ .
2. *Plug-in Schätzer* für $\sigma^2 = \text{Var } X$: Es gilt $\text{Var } X = \mathbb{E}X^2 - (\mathbb{E}X)^2$ und somit folgt

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n X_i^2 - \left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{n-1}{n} S_n^2.$$

3. *Schätzer für Schiefe und Wölbung* $\hat{\gamma}_1$ und $\hat{\gamma}_2$ (vgl. Abschnitt 2.2.4) sind Plug-in Schätzer:
Da der Koeffizient der Schiefe als

$$\gamma_1 = \mathbb{E} \left(\frac{X - \mu}{\sigma} \right)^3$$

definiert ist, wobei $\mu = \mathbb{E}X$, $\sigma^2 = \text{Var } X$, folgt

$$\hat{\gamma}_1 \stackrel{\mu \mapsto \bar{X}_n}{\underset{\sigma^2 \mapsto \hat{\sigma}^2}{\equiv}} \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^3}{(\hat{\sigma}_n^2)^{3/2}} = \frac{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^3}{\left(\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2 \right)^{3/2}}.$$

Die Konstruktion von $\hat{\gamma}_2$ erfolgt analog.

4. Der *empirische Korrelationskoeffizient* ϱ_{XY} ist ein Plug-in Schätzer:

$$\hat{\varrho}_{XY} = \frac{S_{XY}^2}{\sqrt{S_{XX}^2} \sqrt{S_{YY}^2}} = \frac{\sum_{i=1}^n (X_i - \bar{X}_n)(Y_i - \bar{Y}_n)}{\sqrt{\sum_{i=1}^n (X_i - \bar{X}_n)^2} \sqrt{\sum_{i=1}^n (Y_i - \bar{Y}_n)^2}};$$

in der Tat ist

$$\varrho_{XY} = \frac{\mathbb{E}(X - \mathbb{E}X)(Y - \mathbb{E}Y)}{\sqrt{\text{Var } X \cdot \text{Var } Y}} = \frac{\mathbb{E}(XY) - \mathbb{E}X \cdot \mathbb{E}Y}{\sqrt{(\mathbb{E}X^2 - (\mathbb{E}X)^2)(\mathbb{E}Y^2 - (\mathbb{E}Y)^2)}}$$

und somit gilt für die linearen Funktionale

$$T_1(F) = \int x dF(x), \quad T_2(F) = \int x^2 dF(x), \quad T_{12}(G) = \int xy dG(x, y),$$

$$\varrho_{XY} = \frac{T_{12}(F_{XY}) - T_1(F_X) \cdot T_1(F_Y)}{\sqrt{(T_2(F_X) - (T_1(F_X))^2)(T_2(F_Y) - (T_1(F_Y))^2)}}.$$

$\hat{\varrho}_{XY}$ bekommt man, in dem man F_X , F_Y und F_{XY} in T_1 , T_2 und T_{12} durch $\hat{F}_{n,X}$, $\hat{F}_{n,Y}$ und $\hat{F}_{n,XY}$ ersetzt:

$$\hat{\varrho}_{XY} = \frac{T_{12}(\hat{F}_{n,XY}) - T_1(\hat{F}_{n,X}) \cdot T_1(\hat{F}_{n,Y})}{\sqrt{\left(T_2(\hat{F}_{n,X}) - (T_1(\hat{F}_{n,X}))^2\right) \left(T_2(\hat{F}_{n,Y}) - (T_1(\hat{F}_{n,Y}))^2\right)}}.$$

3.4 Methoden zur Gewinnung von Punktschätzern

Sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen identisch verteilten Zufallsvariablen X_i mit Verteilungsfunktion $F \in \{F_\theta : \theta \in \Theta\}$, $\Theta \subset \mathbb{R}^m$ (Parametrisches Modell). Sei die Parametrisierung $\theta \mapsto F_\theta$ unterscheidbar, d.h. $F_\theta \neq F_{\theta'} \iff \theta \neq \theta'$.

Zielstellung: Konstruiere einen Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für $\theta = (\theta_1, \dots, \theta_m)$.

3.4.1 Momentenschätzer

Aus der Wahrscheinlichkeitsrechnung (Satz 4.8) folgt, dass unter gewissen Voraussetzungen (z.B. Gleichverteilung auf einem kompakten Intervall) an die Verteilung F diese Verteilung aus der Kenntnis von Momenten $\mathbb{E}X^k$, $k \in \mathbb{N}$ wiedergewonnen werden kann. Auf dieser Idee der Schätzung von F aus den Momenten basiert die von Karl Pearson am Ende des XIX. Jh. vorgeschlagene *Momentenmethode*.

Annahme: Es existiert ein $r \geq m$, so dass $\mathbb{E}_\theta |X_i|^r < \infty$. Seien die Momente $\mathbb{E}_\theta X_i^k = g_k(\theta)$, $k = 1, \dots, r$ als Funktionen des Parametervektors $\theta = (\theta_1, \dots, \theta_m) \in \Theta$ gegeben.

Momenten-Gleichungssystem: $\hat{\mu}_k = g_k(\theta)$, $k = 1, \dots, r$, wobei $\hat{\mu}_k = \frac{1}{n} \sum_{i=1}^n X_i^k$ die k -ten empirischen Momente sind.

Definition 3.4.1

Falls das obige Gleichungssystem eindeutig lösbar bzgl. θ ist, so heißt die Lösung $\hat{\theta}(X_1, \dots, X_n)$ *Momentenschätzer (M-Schätzer)* von θ .

Lemma 3.4.1

Falls die Funktion $g = (g_1, \dots, g_r) : \Theta \rightarrow C \subset \mathbb{R}^r$ eineindeutig und ihre Inverse $g^{-1} : C \rightarrow \Theta$ stetig ist, dann ist der M-Schätzer $\hat{\theta}(X_1, \dots, X_n)$ von θ stark konsistent.

Beweis Es gilt $\hat{\theta}(X_1, \dots, X_n) = g^{-1}(\hat{\mu}_1, \dots, \hat{\mu}_r) \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \theta$, weil $\hat{\mu}_k \xrightarrow[n \rightarrow \infty]{\text{f.s.}} g_k(\theta)$, $k = 1, \dots, r$ (starke Konsistenz der empirischen Momente) und g^{-1} stetig. $\square$

Bemerkung 3.4.1

1. Unter gewissen Regularitätsbedingungen an F_θ ist der M-Schätzer $\hat{\theta}(X_1, \dots, X_n)$ für θ asymptotisch normalverteilt:

$$\sqrt{n} \left(\hat{\theta}(X_1, \dots, X_n) - \theta \right) \xrightarrow[n \rightarrow \infty]{d} N(0, \Sigma),$$

wobei $N(0, \Sigma)$ die multivariate Normalverteilung mit Kovarianzmatrix

$$\Sigma = G^T \mathbb{E}(Y Y^T) G \quad \text{ist mit} \quad Y = (X, X^2, \dots, X^r)^T, \quad X \stackrel{d}{=} X_i,$$

und

$$G = \left(\frac{\partial g_i^{-1}}{\partial \theta_j} \right)_{\substack{i=1 \dots r, \\ j=1 \dots m}}.$$

2. Andere Eigenschaften gelten für M-Schätzer im Allgemeinen nicht. Zum Beispiel sind nicht alle M-Schätzer erwartungstreu (vgl. Beispiel 3.4.1, 1)).
3. Manchmal sind $r > m$ Gleichungen im Momentensystem notwendig, um einen M-Schätzer zu bekommen. Dies ist zum Beispiel dann der Fall, wenn manche Funktionen $g_i = \text{const}$ sind, d.h. sie enthalten keine Information über θ (vgl. Beispiel 3.4.1, 2)).

Beispiel 3.4.1

1. *Normalverteilung:* $X_i \stackrel{d}{=} X$, $i = 1, \dots, n$, $X \sim N(\mu, \sigma^2)$; Gesucht ist ein M-Schätzer für μ und σ^2 , also $\theta = (\mu, \sigma^2)$. Es gilt

$$\begin{aligned} g_1(\mu, \sigma^2) &= \mathbb{E}_\theta X = \mu, \\ g_2(\mu, \sigma^2) &= \mathbb{E}_\theta X^2 = \text{Var}_\theta X + (\mathbb{E}_\theta X)^2 = \sigma^2 + \mu^2. \end{aligned}$$

Somit ergibt sich das Gleichungssystem

$$\begin{cases} \frac{1}{n} \sum_{i=1}^n X_i = \mu, \\ \frac{1}{n} \sum_{i=1}^n X_i^2 = \mu^2 + \sigma^2. \end{cases}$$

Damit folgt

$$\begin{aligned} \hat{\mu} &= \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}_n, \\ \hat{\sigma}^2 &= \frac{1}{n} \sum_{i=1}^n X_i^2 - \hat{\mu}^2 = \frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i^2 - \bar{X}_n^2) \\ &= \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{n-1}{n} S_n^2. \end{aligned}$$

Das heißt, das die M-Schätzer $\hat{\mu} = \bar{X}_n$, $\hat{\sigma}^2 = \frac{n-1}{n} S_n^2$ sind. Dabei ist $\hat{\sigma}^2$ nicht erwartungstreu:

$$\mathbb{E}_\theta \hat{\sigma}^2 = \frac{n-1}{n} \cdot \mathbb{E}_\theta S_n^2 = \frac{n-1}{n} \sigma^2.$$

2. *Gleichverteilung:* $X_i \stackrel{d}{=} X$, $i = 1, \dots, n$, $X \sim U[-\theta, \theta]$, $\theta > 0$. Gesucht wird ein Momentenschätzer für θ . Es gilt

$$g_1(\theta) = \mathbb{E}_\theta X = 0,$$

$$g_2(\theta) = \mathbb{E}_\theta X^2 = \text{Var}_\theta X = \frac{(\theta - (-\theta))^2}{12} = \frac{(2\theta)^2}{12} = \frac{\theta^2}{3}.$$

Damit ergibt sich das Gleichungssystem

$$\begin{cases} \frac{1}{n} \sum_{i=1}^n X_i = 0 & \text{unbrauchbar,} \\ \frac{1}{n} \sum_{i=1}^n X_i^2 = \frac{\theta^2}{3}. \end{cases}$$

Es folgt, dass $\hat{\theta} = \sqrt{\frac{3}{n} \sum_{i=1}^n X_i^2}$ der Momentenschätzer für θ ist. Wir haben somit 2 Gleichungen für die Schätzung eines einzigen Parameters θ benötigt, d.h. $r = 2 > m = 1$.

3.4.2 Maximum-Likelihood-Schätzer

Diese wurden von Carl Friedrich Gauss (Anfang des XIX. Jh.) und Sir Ronald Fisher (1922) entdeckt. Seien entweder alle Verteilungen aus der parametrischen Familie $\{F_\theta : \theta \in \Theta\}$ diskret oder alle absolut stetig.

Definition 3.4.2

1. Falls die Stichprobenvariablen X_i , $i = 1, \dots, n$ absolut stetig verteilt mit Dichte $f_\theta(x)$ sind, dann heißt

$$L(x_1, \dots, x_n, \theta) = \prod_{i=1}^n f_\theta(x_i), \quad (x_1, \dots, x_n) \in \mathbb{R}^n, \quad \theta \in \Theta$$

die *Likelihood-Funktion* der Stichprobe $(x_1, \dots, x_n)$.

2. Falls die Stichprobenvariablen X_i , $i = 1, \dots, n$ diskret verteilt mit Zähldichte $p_\theta(x) = \mathbb{P}_\theta(X_i = x)$, $x \in C$ sind (C ist der Wertebereich von X), dann heißt

$$L(x_1, \dots, x_n, \theta) = \prod_{i=1}^n p_\theta(x_i), \quad (x_1, \dots, x_n) \in C^n, \quad \theta \in \Theta$$

die *Likelihood-Funktion* der Stichprobe $(x_1, \dots, x_n)$.

Nach dieser Definition gilt im

- *diskreten Fall* $L(x_1, \dots, x_n, \theta) = \mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n)$
- *absolut stetigen Fall*

$$L(x_1, \dots, x_n, \theta) \Delta x_1 \cdot \dots \cdot \Delta x_n = f_{(X_1, \dots, X_n), \theta}(x_1, \dots, x_n) \Delta x_1 \cdot \dots \cdot \Delta x_n$$

$$\approx \mathbb{P}_\theta(X_1 \in [x_1, x_1 + \Delta x_1], \dots, X_n \in [x_n, x_n + \Delta x_n]), \quad \Delta x_i \rightarrow 0, \quad i = 1, \dots, n.$$

Nun wird ein Schätzer für θ so gewählt, dass die Wahrscheinlichkeit

$$\mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n) \quad \text{bzw.} \quad \mathbb{P}_\theta(X_i \in [x_i, x_i + \Delta x_i], \quad i = 1, \dots, n)$$

maximal wird. $\implies$ Maximum-Likelihoodmethode:

Definition 3.4.3

Sei das Maximierungsproblem $L(x_1, \dots, x_n, \theta) \mapsto \max_{\theta \in \Theta}$ eindeutig lösbar. Dann heißt

$$\hat{\theta}(x_1, \dots, x_n) = \arg \max_{\theta \in \Theta} L(x_1, \dots, x_n, \theta)$$

der *Maximum-Likelihood-Schätzer* von θ (*ML-Schätzer*).

Bemerkung 3.4.2

1. In relativ wenigen Fällen ist ein ML-Schätzer $\hat{\theta}$ für θ explizit auffindbar. In diesen Fällen wird meistens der konstante Faktor von $L(x_1, \dots, x_n, \theta)$ weggeworfen und vom Rest der Logarithmus gebildet:

$$\log L(x_1, \dots, x_n, \theta) \quad (\text{die sog. Loglikelihood-Funktion}).$$

Dadurch wird

$$\prod_{i=1}^n f_\theta(x_i) \quad \text{bzw.} \quad \prod_{i=1}^n p_\theta(x_i)$$

zu einer Summe

$$\sum_{i=1}^n \log f_\theta(x_i) \quad \text{bzw.} \quad \sum_{i=1}^n \log p_\theta(x_i),$$

die leichter bzgl. θ zu differenzieren ist. Danach betrachtet man

$$\frac{\partial \log L(x_1, \dots, x_n, \theta)}{\partial \theta_j} = 0, \quad j = 1, \dots, m.$$

Dies ist die notwendige Bedingung eines Extremums von $\log L$ (und somit von L , weil $\log \nearrow$). Falls dieses System eindeutig lösbar ist, und die Lösung eine Maximum-Stelle ist, dann wird sie zum ML-Schätzer $\hat{\theta}(X_1, \dots, X_n)$ erklärt.

2. In den meisten praxisrelevanten Fällen sind ML-Schätzer jedoch nur numerisch auffindbar.

Beispiel 3.4.2

1. *Bernoulli-Verteilung*: $X_i \stackrel{d}{=} X$, $i = 1, \dots, n$, $X \sim \text{Bernoulli}(p)$, für ein $p \in [0, 1]$. Da

$$X = \begin{cases} 1, & \text{mit Wkt. } p \\ 0, & \text{sonst} \end{cases}$$

mit Zähldichte

$$p_\theta(x) = p^x (1-p)^{1-x}, \quad x \in \{0, 1\},$$

ist die *Likelihood-Funktion* der Stichprobe $(X_1, \dots, X_n)$ gegeben durch

$$L(x_1, \dots, x_n, \theta) = \prod_{i=1}^n p^{x_i} (1-p)^{1-x_i} = p^{\sum_{i=1}^n x_i} (1-p)^{n-\sum_{i=1}^n x_i} \stackrel{\text{def.}}{=} h(p).$$

- a) Falls $\sum_{i=1}^n x_i = 0$ ($\Leftrightarrow x_1 = x_2 = \dots = x_n = 0$), es folgt $h(p) = (1-p)^n \rightarrow \max_{p \in [0,1]} h(p) = 1$ bei $p = 0$. Dann ist der ML-Schätzer $\hat{p}(0, \dots, 0) = 0$.
- b) Falls $\sum_{i=1}^n x_i = n$ ($\Leftrightarrow x_1 = x_2 = \dots = x_n = 1$), es folgt $h(p) = p^n \rightarrow \max_{p \in [0,1]} h(p) = 1$ bei $p = 1$. Dann ist der ML-Schätzer $\hat{p}(1, 1, \dots, 1) = 1$.
- c) Falls $0 < \sum_{i=1}^n x_i < n$, dann gilt

$$\log L(x_1, \dots, x_n, p) = n\bar{x}_n \log p + n(1 - \bar{x}_n) \log(1 - p) = n \cdot g(p).$$

Da $g(p) \xrightarrow{p \rightarrow 0,1} -\infty$ und

$$\frac{\partial g(p)}{\partial p} = \frac{\bar{x}_n}{p} + \frac{1 - \bar{x}_n}{1 - p} \cdot (-1) = \frac{\bar{x}_n}{p} + \frac{\bar{x}_n - 1}{1 - p} = 0$$

$$\Leftrightarrow (1 - p)\bar{x}_n + (\bar{x}_n - 1)p = 0 \quad \Rightarrow \quad p = \bar{x}_n,$$

folgt aufgrund der Stetigkeit von g , dass g genau ein Extremum $\arg \max_p g(p) = \bar{x}_n$ besitzt.

Der ML-Schätzer ist also gegeben durch $\hat{p}(X_1, \dots, X_n) = \bar{X}_n$.

2. *Gleichverteilung:* $X \sim U[0, \theta]$, $\theta > 0$, $(X_1, \dots, X_n)$ unabhängig identisch verteilt, gesucht ist ein ML-Schätzer für θ . Es gilt

$$f_{X_i}(x) = 1/\theta \cdot \mathbb{I}(x \in [0, \theta]), \quad i = 1, \dots, n.$$

Somit ist die Likelihood-Funktion durch

$$\begin{aligned} L(x_1, \dots, x_n, \theta) &= \begin{cases} (1/\theta)^n, & 0 \leq x_1, \dots, x_n \leq \theta \\ 0, & \text{sonst} \end{cases} \\ &= \begin{cases} (1/\theta)^n, & \text{falls } \min\{x_1, \dots, x_n\} \geq 0, \quad \max\{x_1, \dots, x_n\} \leq \theta \\ 0, & \text{sonst} \end{cases} \\ &= g(\theta), \quad \theta > 0 \end{aligned}$$

gegeben. Damit folgt $\hat{\theta} = \arg \max_{\theta > 0} g(\theta) = \max\{x_1, \dots, x_n\} = x_{(n)}$, wodurch der ML-

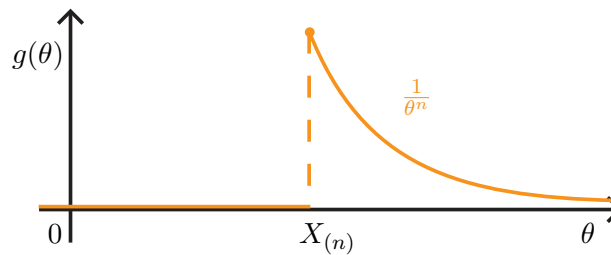


Abb. 3.6: Illustration der Funktion g .

Schätzer durch $\hat{\theta}(X_1, \dots, X_n) = X_{(n)}$ gegeben ist.

Nun wollen wir zeigen, dass ML-Schätzer unter gewissen Voraussetzungen schwach konsistent und asymptotisch normalverteilt sind.

Definition 3.4.4

Sei

$$L(x, \theta) = \begin{cases} f_\theta(x), & \text{im absolut stetigen Fall,} \\ p_\theta(x), & \text{im diskreten Fall} \end{cases}$$

die Likelihood-Funktion von x . Für $\theta, \theta' \in \Theta$ und $X \stackrel{d}{=} X_i$, $\mathbb{P}_\theta(L(X, \theta') = 0) = 0$ definieren wir die *Information* (Abstand) $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'})$ von *Kullback-Leibler* im absolut stetigen Fall als

$$H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = \mathbb{E}_\theta \log L(X, \theta) - \mathbb{E}_\theta \log L(X, \theta') = \int_{\mathbb{R}} \log \frac{L(x, \theta)}{L(x, \theta')} \cdot L(x, \theta) dx.$$

Für den Fall $\mathbb{P}_\theta(L(X, \theta') = 0) > 0$ setzen wir $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = \infty$. Im diskreten Fall betrachte statt des Integrals die Summe über die nicht trivialen $p_\theta(x)$.

Wir werden gleich zeigen, dass $H(\cdot, \cdot)$ die Eigenschaften $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = 0 \iff \theta = \theta'$ und $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) \geq 0 \quad \forall \theta, \theta' \in \Theta$ besitzt. Es ist allerdings offensichtlich, dass $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'})$ nicht symmetrisch bzgl. θ und θ' ist. Somit ist $H(\cdot, \cdot)$ keine Metrik.

Lemma 3.4.2

Es gilt

1. $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'})$ ist wohldefiniert und ≥ 0 .
2. Falls $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = 0$, dann gilt $\theta = \theta'$.

Beweis Wir betrachten zum Beispiel den Fall absolut stetiger $\mathbb{P}_\theta$, $\theta \in \Theta$ (diskreter Fall folgt analog).

1. Definieren wir

$$f(x) = \begin{cases} \frac{L(x, \theta)}{L(x, \theta')}, & \text{falls } L(x, \theta') > 0, \\ 1, & \text{sonst.} \end{cases}$$

Betrachten wir den Fall $\mathbb{P}_\theta(L(X, \theta') = 0) = 0$, so folgt $\mathbb{P}_\theta(L(X, \theta') > 0) = 1$. Ansonsten ist $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = \infty > 0$, also positiv und wohldefiniert. Dann folgt mit Wahrscheinlichkeit 1, dass $L(x, \theta) = f(x) \cdot L(x, \theta')$. Sei $g(x) = 1 - x + x \log x$, $x > 0$. Man kann zeigen, dass g konvex mit $g(x) \geq 0$ ist. Tatsächlich, es gilt

$$g'(x) = -1 + \log x + 1 = \log x, \quad g''(x) = 1/x > 0.$$

Somit besitzt g genau eine Nullstelle bei $x = 1$, die gleichzeitig ihr Minimum ist. Betrachten wir $g(f(X))$, $X \sim L(x, \theta')$. Dann gilt

$$\begin{aligned} 0 &\leq \mathbb{E}_{\theta'} g(f(X)) = 1 - \mathbb{E}_{\theta'} f(X) + \mathbb{E}_{\theta'} (f(X) \log f(X)) \\ &= 1 - \int \frac{L(x, \theta)}{L(x, \theta')} \cdot L(x, \theta') dx + \int \frac{L(x, \theta)}{L(x, \theta')} \cdot \log \frac{L(x, \theta)}{L(x, \theta')} \cdot L(x, \theta') dx = H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}). \end{aligned}$$

Somit gilt $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) \geq 0$, was zu zeigen war.

2. Falls $H(\mathbb{P}_\theta, \mathbb{P}_{\theta'}) = 0 \implies \mathbb{E}_{\theta'} g(f(X)) = 0$, $g(f(X)) \geq 0$. Somit folgt θ' -fast sicher $g(f(X)) = 0 \implies f(X) \stackrel{\theta'\text{-f.s.}}{=} 1$, damit entweder $L(x, \theta') = 0$ oder $L(x, \theta) = L(x, \theta')$ für alle x und daher $\mathbb{P}_\theta = \mathbb{P}_{\theta'}$.

□

Satz 3.4.1 (Schwache Konsistenz von ML-Schätzern):

Sei $m = 1$ und Θ ein offenes Intervall aus $\mathbb{R}$. Sei $L(x_1, \dots, x_n, \theta)$ *unimodal*, d.h. für $\hat{\theta}$ ML-Schätzer für θ gilt

$$\begin{cases} \forall \theta < \hat{\theta}(x_1, \dots, x_n) & \implies L(x_1, \dots, x_n, \theta) \text{ ist steigend} \\ \forall \theta > \hat{\theta}(x_1, \dots, x_n) & \implies L(x_1, \dots, x_n, \theta) \text{ ist fallend} \end{cases}$$

(d.h. es existiert genau ein $\max_{\theta \in \Theta} L(x_1, \dots, x_n, \theta)$). Dann gilt $\hat{\theta}(X_1, \dots, X_n) \xrightarrow[n \rightarrow \infty]{P} \theta$.

Beweis Es ist zu zeigen, dass

$$\mathbb{P}_\theta \left(\left| \hat{\theta}(X_1, \dots, X_n) - \theta \right| > \varepsilon \right) \xrightarrow[n \rightarrow \infty]{} 0, \quad \varepsilon > 0. \quad (3.4.1)$$

Wählen wir beliebiges $\varepsilon > 0$: $\theta \pm \varepsilon \in \Theta$. Dann gilt $H(\mathbb{P}_\theta, \mathbb{P}_{\theta \pm \varepsilon}) > \delta > 0$, wegen der Unterscheidbarkeit der Parametrisierung von $\mathbb{P}_\theta$ und Lemma 3.4.2. Betrachten wir $\{|\hat{\theta} - \theta| \leq \varepsilon\}$. Um (3.4.1) zu zeigen, ist es hinreichend, eine untere Schranke für $\mathbb{P}_\theta(|\hat{\theta} - \theta| \leq \varepsilon)$ zu konstruieren, die für $n \rightarrow \infty$ gegen 1 konvergiert. Es gilt

$$\begin{aligned} \{|\hat{\theta} - \theta| < \varepsilon\} &\stackrel{\text{Unimod}}{\supseteq} \{L(X_1, \dots, X_n, \theta - \varepsilon) < L(X_1, \dots, X_n, \theta) < L(X_1, \dots, X_n, \theta + \varepsilon)\} \\ &= \left\{ \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta \pm \varepsilon)} > 1 \right\} \stackrel{\delta > 0 \implies e^{n\delta} > 1}{\supseteq} \left\{ \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta \pm \varepsilon)} > e^{n\delta} \right\} \\ &= \left\{ \frac{1}{n} \log \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta \pm \varepsilon)} > \delta \right\} = A_+ \cap A_-, \end{aligned}$$

wobei

$$A_\pm = \left\{ \frac{1}{n} \log \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta \pm \varepsilon)} > \delta \right\}.$$

Somit gilt also

$$\mathbb{P}_\theta \left(|\hat{\theta} - \theta| < \varepsilon \right) \geq \mathbb{P}_\theta(A_+ \cap A_-) = \mathbb{P}_\theta(A_+) + \mathbb{P}_\theta(A_-) - \mathbb{P}_\theta(A_+ \cup A_-).$$

Wenn wir zeigen können, dass

$$\lim_{n \rightarrow \infty} \mathbb{P}_\theta(A_\pm) = 1, \quad (3.4.2)$$

dann folgt daraus

$$1 \geq \lim_{n \rightarrow \infty} \mathbb{P}_\theta(A_+ \cup A_-) \geq \lim_{n \rightarrow \infty} \mathbb{P}_\theta(A_\pm) = 1 \implies \lim_{n \rightarrow \infty} \mathbb{P}_\theta(A_+ \cup A_-) = 1$$

und

$$1 \geq \lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(|\hat{\theta} - \theta| < \varepsilon \right) \geq 1 + 1 - 1 = 1,$$

womit folgt, dass

$$\lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(|\hat{\theta} - \theta| > \varepsilon \right) \leq 1 - \underbrace{\lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(|\hat{\theta} - \theta| < \varepsilon \right)}_{=1} = 0,$$

d.h. $\hat{\theta} \xrightarrow[n \rightarrow \infty]{P} \theta$.

Jetzt zeigen wir, dass $\mathbb{P}_\theta(A_+) \xrightarrow[n \rightarrow \infty]{} 1$ (für $\mathbb{P}_\theta(A_-)$ ist es analog).

1. Sei $H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) < \infty$. Sei

$$f(x) = \begin{cases} \frac{L(x, \theta)}{L(x, \theta + \varepsilon)}, & \text{falls } L(x, \theta + \varepsilon) > 0, \\ 1, & \text{sonst.} \end{cases}$$

Dann folgt aus Definition 3.4.4, dass $\mathbb{P}_\theta(L(X_1, \theta + \varepsilon) > 0) = 1$. Weiter gilt

$$\begin{aligned} \frac{1}{n} \log \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta + \varepsilon)} &= \frac{1}{n} \sum_{i=1}^n \log \frac{L(X_i, \theta)}{L(X_i, \theta + \varepsilon)} = \frac{1}{n} \sum_{i=1}^n \log f(X_i) \\ &\xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta \log f(X_1) = \int L(x, \theta) \cdot \log \frac{L(x, \theta)}{L(x, \theta + \varepsilon)} dx = H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) > \delta > 0 \end{aligned}$$

nach dem starken Gesetz der großen Zahlen, weil $\log f(X_1) \in L^1(\Omega, \mathcal{F}, \mathbb{P})$ wegen

$$\mathbb{E}_\theta \log f(X_1) = H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) < \infty \implies \mathbb{P}(A_+) \xrightarrow[n \rightarrow \infty]{} 1.$$

2. Sei $H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) = \infty$ und $\mathbb{P}_\theta(L(X_1, \theta + \varepsilon) = 0) = 0$, dann folgt

$$f(x) \stackrel{\text{f.s.}}{=} \frac{L(x, \theta)}{L(x, \theta + \varepsilon)}$$

bzgl. der Verteilung P_{X_1} . Es gilt $\log \min\{f(X_1), c\} \in L^1(\Omega, \mathcal{F}, \mathbb{P})$ für alle $c > 0$. Somit folgt wie in Punkt 1:

$$\frac{1}{n} \sum_{i=1}^n \log \min\{f(X_i), c\} \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta \log \min\{f(X_1), c\} \in (0, \infty) \xrightarrow[c \rightarrow \infty]{} H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) = \infty$$

und damit

$$\begin{aligned} A_+ &\supset \left\{ \frac{1}{n} \sum_{i=1}^n \log \min\{f(X_i), c\} > \delta \right\} \\ \implies \mathbb{P}(A_+) &\geq \mathbb{P} \left(\frac{1}{n} \sum_{i=1}^n \log \min\{f(X_i), c\} > \delta \right) \xrightarrow[n \rightarrow \infty]{} 1. \end{aligned}$$

3. Sei $H(\mathbb{P}_\theta, \mathbb{P}_{\theta+\varepsilon}) = \infty$ und $\mathbb{P}_\theta(L(X_1, \theta + \varepsilon) = 0) = a > 0$, dann folgt

$$\begin{aligned} \mathbb{P}_\theta \left(\frac{1}{n} \log \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta + \varepsilon)} = \infty \right) &= 1 - \mathbb{P} \left(\frac{1}{n} \log \frac{L(X_1, \dots, X_n, \theta)}{L(X_1, \dots, X_n, \theta + \varepsilon)} < \infty \right) \\ &= 1 - \mathbb{P} \left(\bigcap_{i=1}^n \{L(X_i, \theta + \varepsilon) > 0\} \right) \stackrel{X_i \text{ u.i.v.}}{=} 1 - (1 - a)^n \xrightarrow[n \rightarrow \infty]{} 1 \end{aligned}$$

Insgesamt also $\mathbb{P}(A_+) \xrightarrow[n \rightarrow \infty]{} 1$.

□

Definition 3.4.5

Sei $X = (X_1, \dots, X_n)$ eine Zufallsstichprobe von unabhängigen identisch verteilten Zufallsvariablen $X_i \sim F_\theta$, $\theta \in \Theta$. Sei $L(x, \theta)$ die Likelihood-Funktion von X_i . Dann heißt der Ausdruck

$$I(\theta) = \mathbb{E}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \theta) \right)^2, \quad \theta \in \Theta \quad (3.4.3)$$

die *Fisher-Information* der Stichprobe $(X_1, \dots, X_n)$.

Es wird in Zukunft vorausgesetzt, dass $0 < I(\theta) < \infty$. Wir stellen nun einige Bedingungen auf, die für die asymptotische Normalverteiltetheit von ML-Schätzern notwendig sind.

1. $\Theta \subset \mathbb{R}$ ist ein offenes Intervall ($m = 1$).
2. Es gelte $\mathbb{P}_\theta \neq \mathbb{P}_{\theta'}$ genau dann, wenn $\theta \neq \theta'$.
3. Die Familie $\{\mathbb{P}_\theta, \theta \in \Theta\}$, $\theta \in \Theta$ bestehe nur aus diskreten oder nur aus absolut stetigen Verteilungen, also nicht aus Mischungen von diskreten und absolut stetigen Verteilungen.
4. $B = \text{supp } L(x, \theta) = \{x \in \mathbb{R} : L(x, \theta) > 0\}$ hängt nicht von $\theta \in \Theta$ ab. Dabei heißt supp (von englisch „support“) der „Träger“ einer Funktion f und ist definiert als

$$\text{supp } f = \{x \in \mathbb{R} : f(x) \neq 0\}$$

und die Likelihood-Funktion $L(x, \theta)$ ist durch

$$L(x, \theta) = \begin{cases} p(x, \theta), & \text{im diskreten Fall,} \\ f(x, \theta), & \text{im absolut stetigen Fall} \end{cases} \quad (3.4.4)$$

gegeben, wobei $p(x, \theta)$ bzw. $f(x, \theta)$ die Wahrscheinlichkeitsfunktion bzw. Dichte von $\mathbb{P}_\theta$ ist.

5. Die Abbildung $L(x, \theta)$ ist dreimal stetig differenzierbar und es gilt

$$0 = \frac{d^k}{d\theta^k} \int_B L(x, \theta) dx = \int_B \frac{\partial^k}{\partial \theta^k} L(x, \theta) dx, \quad k = 1, 2, \theta \in \Theta.$$

Da das Integral über die Dichte $L(x, \theta)$ gleich 1 ist, ist die Ableitung gleich 0. Dabei sind im diskreten Fall die Integrale durch Summen zu ersetzen.

6. Für alle $\theta_0 \in \Theta$ existiert eine Konstante $\delta_{\theta_0} > 0$ und eine messbare Funktion $g_{\theta_0} : B \rightarrow [0, \infty)$, so dass

$$\left| \frac{\partial^3 \log L(x, \theta)}{\partial \theta^3} \right| \leq g_{\theta_0}(x), \quad \forall x \in B, \quad |\theta - \theta_0| < \delta_{\theta_0},$$

wobei $\mathbb{E}_{\theta_0} g_{\theta_0}(X_1) < \infty$.

Bemerkung 3.4.3

Es gilt folgende Relation:

$$n \cdot I(\theta) = \text{Var}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \dots, X_n, \theta) \right),$$

wobei

$$L(X_1, \dots, X_n, \theta) = \prod_{i=1}^n L(X_i, \theta) \quad (3.4.5)$$

die Likelihood-Funktion der Stichprobe $(X_1, \dots, X_n)$ ist mit $L(X_i, \theta)$ nach (3.4.4).

Beweis Es gilt

$$\frac{\partial}{\partial \theta} \log L(X_1, \dots, X_n, \theta) = \frac{\partial}{\partial \theta} \sum_{i=1}^n \log L(X_i, \theta) = \sum_{i=1}^n \frac{\partial}{\partial \theta} \log L(X_i, \theta) = \sum_{i=1}^n \frac{L'(X_i, \theta)}{L(X_i, \theta)}.$$

Ferner

$$\mathbb{E}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \dots, X_n, \theta) \right) = \sum_{i=1}^n \mathbb{E}_\theta \frac{L'(X_i, \theta)}{L(X_i, \theta)} = \sum_{i=1}^n \int_B \frac{L'(X, \theta)}{L(X, \theta)} \cdot L(X, \theta) dx \stackrel{5)}{=} 0.$$

Insgesamt gilt also

$$\begin{aligned} \text{Var}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \dots, X_n, \theta) \right) &= \text{Var}_\theta \left(\sum_{i=1}^n \frac{\partial}{\partial \theta} \log L(X_i, \theta) \right) \\ &\stackrel{X_i \text{ unabh.}}{=} \sum_{i=1}^n \text{Var}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_i, \theta) \right) \stackrel{X_i \text{ ident. vert.}}{=} n \cdot \text{Var}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \theta) \right) \\ &= n \cdot \mathbb{E}_\theta \left(\frac{\partial}{\partial \theta} \log L(X_1, \theta) \right)^2 = n \cdot I(\theta). \end{aligned}$$

□

Beispiel 3.4.3

Seien $X_i \sim N(\mu, \theta^2)$, $i = 1, \dots, n$. Für $\theta = \mu$ zeigen wir bei benannten σ^2 , dass $\mathbb{I}(\mu) = \frac{1}{\sigma^2}$.

In der Tat, $L(X_1, \mu) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{(X_1 - \mu)^2}{2\sigma^2} \right\}$,
 $\log L(X_1, \mu) = -\log(\sqrt{2\pi}\sigma) - (X_1 - \mu)^2/(2\sigma^2)$,
 $\frac{\partial \log L(X_1, \mu)}{\partial \mu} = -\frac{2(X_1 - \mu)}{2\sigma^2} \cdot (-1) = \frac{X_1 - \mu}{\sigma^2}$, somit gilt

$$I(\mu) = \mathbb{E}_\mu \left(\frac{\partial \log L(X_1, \mu)}{\partial \mu} \right)^2 = \frac{1}{\sigma^4} \mathbb{E}_\mu (X_1 - \mu)^2 = \frac{1}{\sigma^4} \cdot \sigma^2 = \frac{1}{\sigma^2}.$$

Damit ist $\text{Var}_\mu \left(\frac{\partial}{\partial \mu} \log L(X_1, \dots, X_n, \mu) \right) = \frac{n}{\sigma^2} = \frac{1}{\text{Var}_\mu(\hat{\mu})}$ nach Bemerkung 3.4.3 und Satz 3.3.1, 4), wobei $\hat{\mu} = \bar{X}_n$.

Es bedeutet, dass wenig Information über μ (kleine Werte von $I(\mu)$) eine große Streuung bei der Schätzung von μ bedeutet und umgekehrt.

Satz 3.4.2

Sei $(X_1, \dots, X_n)$ eine Stichprobe von Zufallsvariablen, für die die Bedingungen 1) bis 6) erfüllt sind und $0 < I(\theta) < \infty$, $\theta \in \Theta$. Falls $\hat{\theta}(X_1, \dots, X_n)$ ein schwach konsistenter ML-Schätzer für θ ist, dann ist $\hat{\theta}(X_1, \dots, X_n)$ asymptotisch normalverteilt:

$$\sqrt{n \cdot I(\theta)} \left(\hat{\theta}(X_1, \dots, X_n) - \theta \right) \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1).$$

Beweis Führen wir die Bezeichnung $l_n(\theta) = \log L(X_1, \dots, X_n, \theta)$, $\theta \in \Theta$ ein. Sei

$$l_n^{(k)}(\theta) = \frac{d^k}{d\theta^k} l_n(\theta), \quad k = 1, 2, 3.$$

Ist $\hat{\theta}$ ein ML-Schätzer, so folgt $l_n^{(1)}(\hat{\theta}) = 0$. Schreiben wir die Taylor-Entwicklung von $l_n^{(1)}(\hat{\theta})$ in der Umgebung von θ auf:

$$0 = l_n^{(1)}(\hat{\theta}) = l_n^{(1)}(\theta) + (\hat{\theta} - \theta) \cdot l_n^{(2)}(\theta) + (\hat{\theta} - \theta)^2 \cdot \frac{l_n^{(3)}(\theta^*)}{2},$$

wobei θ^* zwischen θ und $\hat{\theta}$ liegt. Dabei ist

$$-(\hat{\theta} - \theta) \left(l_n^{(2)}(\theta) + (\hat{\theta} - \theta) \frac{l_n^{(3)}(\theta^*)}{2} \right) = l_n^{(1)}(\theta) \implies \sqrt{n}(\hat{\theta} - \theta) = \frac{\frac{l_n^{(1)}(\theta)}{\sqrt{n}}}{-\frac{l_n^{(2)}(\theta)}{n} - (\hat{\theta} - \theta) \frac{l_n^{(3)}(\theta^*)}{2n}}$$

Falls wir zeigen können, dass

1. $\frac{l_n^{(1)}(\theta)}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{d} N(0, I(\theta)),$
2. $-\frac{l_n^{(2)}(\theta)}{n} \xrightarrow[n \rightarrow \infty]{\text{f.s.}} I(\theta),$
3. $(\hat{\theta} - \theta) \xrightarrow[n \rightarrow \infty]{P} 0 \quad \text{und} \quad \frac{l_n^{(3)}(\theta^*)}{2n}$

beschränkt ist, das heißt

$$\exists c > 0 : \quad \lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(\left| \frac{l_n^{(3)}(\theta^*)}{2n} \right| < c \right) = 1,$$

dann konvergiert der Ausdruck

$$(\hat{\theta} - \theta) \cdot \frac{l_n^{(3)}(\theta^*)}{2n} \xrightarrow[n \rightarrow \infty]{P} 0, \text{ weil } \left| \frac{l_n^{(3)}(\theta^*)}{n} \right| \leq g_\theta(X_1) \text{ integrierbar}$$

und somit gilt

$$\sqrt{n}(\hat{\theta} - \theta) = \frac{\frac{l_n^{(1)}(\theta)}{\sqrt{n}}}{-\frac{l_n^{(2)}(\theta)}{n} - (\hat{\theta} - \theta) \frac{l_n^{(3)}(\theta^*)}{2n}} \xrightarrow[n \rightarrow \infty]{d} Z_1 \sim N\left(0, \frac{1}{I(\theta)}\right)$$

nach dem Satz von Slutsky. Damit folgt $\sqrt{n}\sqrt{I(\theta)}(\hat{\theta} - \theta) \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1)$

1. Es gilt

$$\frac{l_n^{(1)}(\theta)}{\sqrt{n}} = \frac{\sum_{i=1}^n \frac{\partial}{\partial \theta} \log L(X_i, \theta)}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{d} Y_1 \sim N\left(0, \underbrace{\text{Var}_\theta\left(\frac{\partial}{\partial \theta} \log L(X_i, \theta)\right)}_{=I(\theta)}\right)$$

nach dem zentralen Grenzwertsatz, weil $\frac{\partial}{\partial \theta} \log L(X_i, \theta)$ unabhängig identisch verteilte Zufallsvariablen mit Erwartungswert 0 (siehe Bemerkung 3.4.3) sind.

2.

$$\begin{aligned}
-\frac{1}{n} l_n^{(2)}(\theta) &= -\frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta^2} \log L(X_i, \theta) = \frac{1}{n} \sum_{i=1}^n \frac{\left(L^{(1)}(X_i, \theta) \right)^2 - L(X_i, \theta) \cdot L^{(2)}(X_i, \theta)}{(L(X_i, \theta))^2} \\
&= \frac{1}{n} \sum_{i=1}^n \left(\frac{L^{(1)}(X_i, \theta)}{L(X_i, \theta)} \right)^2 - \frac{1}{n} \sum_{i=1}^n \frac{L^{(2)}(X_i, \theta)}{L(X_i, \theta)} \\
&\xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta \left(\frac{L^{(1)}(X_1, \theta)}{L(X_1, \theta)} \right)^2 - \mathbb{E}_\theta \left(\frac{L^{(2)}(X_1, \theta)}{L(X_1, \theta)} \right) = I(\theta) - 0 = I(\theta)
\end{aligned}$$

nach dem Gesetz der großen Zahlen, wobei

$$L^{(k)}(X_i, \theta) = \frac{\partial^k}{\partial \theta^k} L(X_i, \theta)$$

und

$$\mathbb{E}_\theta \left(\frac{L^{(2)}(X_1, \theta)}{L(X_1, \theta)} \right) = \int_B \frac{\partial^2}{\partial \theta^2} L(x, \theta) dx \stackrel{5)}{=} \frac{d^2}{d\theta^2} \int_B L(x, \theta) dx = 0.$$

3. $\hat{\theta} \xrightarrow[n \rightarrow \infty]{P} \theta$, weil $\hat{\theta}$ schwach konsistent ist. Zeigen wir, dass

$$\frac{l_n^{(3)}(\theta^*)}{n} (\hat{\theta} - \theta) \xrightarrow[n \rightarrow \infty]{P} 0.$$

Aus $\hat{\theta} \xrightarrow[n \rightarrow \infty]{P} \theta$ folgt für alle $\varepsilon > 0$

$$\mathbb{P}(|\hat{\theta} - \theta| \leq \varepsilon) \xrightarrow[n \rightarrow \infty]{} 1.$$

Damit folgt, dass mit asymptotisch großer Wahrscheinlichkeit $|\hat{\theta} - \theta| \leq \delta$, $\delta > 0$ gilt, welches aus der Bedingung 6) folgt. Damit gilt, dass für alle θ : $|\hat{\theta} - \theta| < \delta$

$$\left| \frac{l_n^{(3)}(\theta)}{n} \right| \leq \frac{1}{n} \sum_{i=1}^n \underbrace{\left| \frac{\partial^3}{\partial \theta^3} \log L(X_i, \theta) \right|}_{\leq g_\theta(X_i)} \leq \frac{1}{n} \sum_{i=1}^n g_\theta(X_i) \xrightarrow[n \rightarrow \infty]{\text{f.s.}} \mathbb{E}_\theta g_\theta(X_1) < \infty.$$

So folgt, dass eine Konstante $c > 0$ existiert, sodass

$$\mathbb{P}_\theta \left(\left| \frac{l_n^{(3)}(\theta^*)}{n} \right| < c \right) \xrightarrow[n \rightarrow \infty]{} 1 \quad \text{und somit} \quad \frac{l_n^{(3)}(\theta^*)}{n} (\hat{\theta} - \theta) \xrightarrow[n \rightarrow \infty]{P} 0.$$

Der Beweis ist beendet.

□

3.4.3 Bayes-Schätzer

Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe, wobei X_i unabhängige identisch verteilte Zufallsvariablen mit Verteilungsfunktion F_θ , $\theta \in \Theta$ sind. Sei F_θ entweder eine diskrete oder eine absolut stetige Verteilung. Sei aber auch θ eine Zufallsvariable $\tilde{\theta}$ mit Verteilung $Q(\cdot)$ auf dem Messraum $(\Theta, \mathcal{B}_\Theta)$, die entweder diskret mit Zähldichte $q(\cdot)$ oder absolut stetig mit Dichte $q(\cdot)$ ist. Nach wie vor werden beide Fälle gemeinsam betrachtet, dabei entsprechen sich die Summation und Integration im diskreten bzw. absolut stetigen Fall.

Definition 3.4.6

Die Verteilung $Q(\cdot)$ heißt *a-priori-Verteilung* des Parameters θ (von $\tilde{\theta}$) (a-priori bedeutet hier „vor dem Experiment $(X_1, \dots, X_n)$ “).

Definition 3.4.7

Die *a-posteriori-Verteilung* des Parameters θ (von $\tilde{\theta}$) ist gegeben durch die (Zähl-)Dichte

$$q_{X_1, \dots, X_n}(\theta, X_1, \dots, X_n) = \begin{cases} \mathbb{P}(\tilde{\theta} = \theta \mid X_1 = x_1, \dots, X_n = x_n), & \text{falls die Verteilung } Q \text{ diskret ist,} \\ f_{\tilde{\theta} \mid X_1, \dots, X_n}(\theta, x_1, \dots, x_n), & \text{falls die Verteilung } Q \text{ absolut stetig ist.} \end{cases}$$

Dabei ist

$$\begin{aligned} \mathbb{P}(\tilde{\theta} = \theta \mid X_1 = x_1, \dots, X_n = x_n) &= \frac{\mathbb{P}(\tilde{\theta} = \theta, X_1 = x_1, \dots, X_n = x_n)}{\mathbb{P}(X_1 = x_1, \dots, X_n = x_n)} \\ &= \frac{\mathbb{P}_\theta(X_i = x_i, i = 1, \dots, n) \cdot q(\theta)}{\sum_{\theta_1 \in \Theta} \mathbb{P}_{\theta_1}(X_i = x_i, i = 1, \dots, n) \cdot q(\theta_1)} \end{aligned}$$

die *Bayesche Formel*, bzw.

$$f_{\tilde{\theta} \mid X_1, \dots, X_n}(\theta, x_1, \dots, x_n) = \frac{f_{(\tilde{\theta}, X_1, \dots, X_n)}(\theta, x_1, \dots, x_n)}{f_{X_1, \dots, X_n}(x_1, \dots, x_n)} = \frac{L(x_1, \dots, x_n, \theta) \cdot q(\theta)}{\int_{\Theta} L(x_1, \dots, x_n, \theta_1) \cdot q(\theta_1) d\theta_1},$$

mit $L(x_1, \dots, x_n, \theta)$ nach (3.4.5).

Definition 3.4.8

Eine *Verlustfunktion* $V : \Theta^2 \rightarrow \mathbb{R}_+$ ist eine Θ^2 -messbare Funktion.

Verlustfunktionen spielen in unseren Betrachtungen folgende Rolle: $\mathbb{E}_* V(\tilde{\theta}, a)$ stellt den *erwarteten Verlust* (mittleres Risiko) dar, der bei der Schätzung des Parameters θ durch a entsteht. Dabei stellt $\mathbb{E}_*$ den Erwartungswert bezüglich der *a-posteriori-Verteilung* von $\tilde{\theta}$ dar. Es sind offensichtlich die konkreten Stichprobenwerte $x_1, \dots, x_n$ in die a-posteriori-Verteilung eingegangen, deshalb ist $\mathbb{E}_* V(\tilde{\theta}, a)$ eine Funktion von a und $x_1, \dots, x_n$:

$$\mathbb{E}_* V(\tilde{\theta}, a) = \varphi(x_1, \dots, x_n, a).$$

Definition 3.4.9

Ein Schätzer $\hat{\theta}$ heißt *Bayes-Schätzer* des Parameters θ , falls

$$\hat{\theta}(x_1, \dots, x_n) = \arg \min_a \mathbb{E}_* V(\tilde{\theta}, a) \quad (3.4.6)$$

existiert und eindeutig ist.

Bemerkung 3.4.4

1. Manchmal gilt $\hat{\theta} \notin \Theta$, was mit der Existenz des Minimums von $\varphi(x_1, \dots, x_n, a)$ auf Θ zu tun hat.
2. Der Name „Bayesscher Ansatz“ stammt von dem englischen Mathematiker Thomas Bayes (1702–1761), der die Bayessche Formel

$$\mathbb{P}(B_i|A) = \frac{\mathbb{P}(A|B_i) \cdot \mathbb{P}(B_i)}{\sum_j \mathbb{P}(A|B_j) \cdot \mathbb{P}(B_j)} \quad (3.4.7)$$

nur ideenhaft eingeführt hat. Der eigentliche Entdecker der Formel (3.4.7) ist Pierre-Simon Laplace (1749–1827) (Ende des XVIII. Jahrhunderts). Diese Formel wurde bei der Herleitung der *a-posteriori-Verteilung* von $\tilde{\theta}$ implizit benutzt.

3. Die Vorgehensweise in Definition 3.4.9 ist in konkreten praxisrelevanten Fällen meistens nur numerisch möglich. Es gibt sehr wenige Beispiele für analytische Lösungen des in (3.4.6) gestellten Minimierungsproblems.

Beispiel 3.4.4 (Quadratische Verlustfunktion):

Ist $V(\theta_1, \theta_2) = (\theta_1 - \theta_2)^2$, so ist

$$\arg \min_a (\varphi(x_1, \dots, x_n, a)) = \arg \min_a (\mathbb{E}_*(\tilde{\theta} - a)^2) = \arg \min_a (\mathbb{E}_*\tilde{\theta}^2 - 2a\mathbb{E}_*\tilde{\theta} + a^2) = \mathbb{E}_*\tilde{\theta}$$

und daher der *Bayes-Schätzer* $\hat{\theta}(x_1, \dots, x_n)$ für θ durch $\mathbb{E}_*\tilde{\theta}$ gegeben.

Beispiel 3.4.5 (Bernoulli-Verteilung):

Sei $(X_1, \dots, X_n)$ eine unabhängig identisch verteilte Stichprobe von $X_i \sim \text{Bernoulli}(p)$, $p \in (0, 1)$. Weiter sei die a-priori-Verteilung

$$\tilde{p} \sim \text{Beta}(\alpha, \beta), \quad \alpha, \beta > 0, \text{ mit Zähldichte } q(p) = \frac{p^{\alpha-1}(1-p)^{\beta-1}}{B(\alpha, \beta)} \cdot \mathbb{I}(p \in [0, 1]),$$

die a-posteriori-Verteilung von $\tilde{p}$ ist dann gleich

$$q^*(p) = f_{\tilde{p}|X_1=x_1, \dots, X_n=x_n}(p) = \frac{\mathbb{P}_p(X_1 = x_1, \dots, X_n = x_n) \cdot q(p)}{\int_0^1 \mathbb{P}_{p_1}(X_1 = x_1, \dots, X_n = x_n) \cdot q(p_1) dp_1}.$$

Es ist immer möglich die a-posteriori-Verteilung nicht bezüglich des Vektors $(X_1, \dots, X_n)$, sondern bezüglich einer Funktion $g(X_1, \dots, X_n)$, zu berechnen (*Komplexitätsreduktion*).

Hier ist $Y = g(X_1, \dots, X_n) = \sum_{i=1}^n X_i$ die Gesamtanzahl aller Erfolge in n Experimenten, wobei

$$X_i = \begin{cases} 1, & \text{mit Wahrscheinlichkeit } p, \\ 0, & \text{sonst.} \end{cases}$$

Daher gilt für die a-posteriori-Verteilung bzgl. Y :

$$\begin{aligned} q^*(p) = f_{\tilde{p}|Y=k}(p) &= \frac{\mathbb{P}_p(Y = k) \cdot q(p)}{\int_0^1 \mathbb{P}_{p_1}(Y = k) q(p_1) dp_1} \\ &\stackrel{Y \sim \text{Bin}(n, p)}{\underset{\text{falls } \tilde{p}=p}{=}} \frac{\binom{n}{k} p^k (1-p)^{n-k} \cdot (B(\alpha, \beta))^{-1} \cdot p^{\alpha-1} (1-p)^{\beta-1}}{\frac{\binom{n}{k}}{B(\alpha, \beta)} \cdot \int_0^1 p_1^{k+\alpha-1} (1-p_1)^{n-k+\beta-1} dp_1} \\ &= \frac{p^{k+\alpha-1} (1-p)^{n-k+\beta-1}}{B(k+\alpha, n-k+\beta)}, \quad p \in [0, 1]. \end{aligned}$$

Daher ist die a-posteriori-Verteilung von $\tilde{p}$ unter der Bedingung $Y = k$ durch

$$\text{Beta}(k + \alpha, n - k + \beta)$$

gegeben.

Für den *Bayes-Schätzer* gilt:

$$\begin{aligned}\hat{p}(x_1, \dots, x_n) &= \mathbb{E}_* \tilde{p} = \int_0^1 p \cdot q^*(p) dp = \frac{\int_0^1 p^{k+\alpha} (1-p)^{n-k+\beta-1} dp}{B(k+\alpha, n-k+\beta)} \\ &= \frac{B(k+\alpha+1, n-k+\beta)}{B(k+\alpha, n-k+\beta)} = \dots = \frac{k+\alpha}{\alpha+\beta+n} = \frac{\sum_{i=1}^n x_i + \alpha}{\alpha+\beta+n} = \frac{\alpha + n\bar{x}_n}{\alpha+\beta+n}.\end{aligned}$$

Interpretation:

$$\hat{p}(X_1, \dots, X_n) = \underbrace{\frac{n}{\alpha+\beta+n}}_{=:c_1} \bar{X}_n + \underbrace{\frac{\alpha+\beta}{\alpha+\beta+n} \cdot \frac{\alpha}{\alpha+\beta}}_{=:c_2} = c_1 \cdot \bar{X}_n + c_2 \cdot \mathbb{E}_{apr} \tilde{\theta},$$

wobei $c_1 + c_2 = 1$ ist. Dies heißt, dass die Bayessche Methode einen Mittelweg zwischen dem Schätzer $\mathbb{E}_{apr} \tilde{\theta}$ (in Abwesenheit der Information über die Stichprobe $(X_1, \dots, X_n)$) und dem M-Schätzer $\bar{X}_n$ (in Abwesenheit der a-priori-Information über die Verteilung von $\tilde{p}$) für p einschlägt.

3.4.4 Resampling-Methoden zur Gewinnung von Punktschätzern

Sei $(X_1, \dots, X_n)$ eine Stichprobe im parametrischen Modell. Gesucht ist ein Schätzer $\hat{\theta}$ für den Parameter θ . Um diesen Schätzer zu konstruieren, werden bei Resampling-Methoden neue Stichproben $(X_1^*, \dots, X_n^*)$ durch das unabhängige Ziehen mit Zurücklegen aus der alten Stichprobe $(X_1, \dots, X_n)$ generiert und auf ihrer Basis Mittelwerte, Stichprobenvarianzen und andere Schätzer gebildet. Dabei ist die Dimension m des Parameterraums Θ beliebig.

Wir werden im Folgenden die *Resampling-Methoden*

1. *Jackknife* (dt. „Taschenmesser“, weist auf Mittel, die jedem immer zur Hand sein sollten)
2. *Bootstrap* (engl. „self-sufficient“, dt. „mit eigenen Ressourcen“)

betrachten.

1. *Jackknife-Methoden zur Schätzung der Varianz bzw. der Verzerrung von Schätzern:*

Als einführendes Beispiel betrachten wir $\theta = \mathbb{E}X = \mu$ bzw. $\theta = \text{Var } X = \sigma^2$ und ihre (erwartungstreue) Schätzer $\hat{\mu} = \bar{X}_n$ bzw. $\hat{\sigma}^2 = S_n^2$.

Wie wir bereits wissen, gilt

$$\text{Var } \hat{\mu} = \frac{\sigma^2}{n}, \quad \text{Var } \hat{\sigma}^2 = \frac{1}{n} \left(\mu'_4 - \frac{n-3}{n-1} \sigma^4 \right).$$

Nun ist ein Schätzer für die Varianz von $\hat{\mu}$ bzw. $\hat{\sigma}^2$ gesucht. Dazu verwenden wir die Plug-in Methode

$$\widehat{\text{Var}} \hat{\mu} = \frac{S_n^2}{n}, \quad \widehat{\text{Var}} \hat{\sigma}^2 = \frac{1}{n} \left(\hat{\mu}'_4 - \frac{n-3}{n-1} S_n^4 \right),$$

wobei $\hat{\mu}'_4$ das vierte zentrierte empirische Moment ist.

Im Allgemeinen sind jedoch *keine* Formeln von $\text{Var } \hat{\theta}$ bekannt. Hier kommt nun die *Jackknife*-Methode zum Einsatz:

- Sei $X_{[i]}$ die Stichprobe $(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n)$, $i = 1, \dots, n$. Falls

$$\hat{\theta}(X_1, \dots, X_n) = \varphi_n(X_1, \dots, X_n),$$

so bilden wir

$$\hat{\theta}_{[i]} = \varphi_{n-1}(X_{[i]}), \quad \bar{\theta}_{[\cdot]} = \frac{1}{n} \sum_{i=1}^n \hat{\theta}_{[i]}, \quad \widehat{\text{Var}}_{jn}(\hat{\theta}) \stackrel{\text{def.}}{=} \frac{n-1}{n} \sum_{i=1}^n \left(\hat{\theta}_{[i]} - \bar{\theta}_{[\cdot]} \right)^2.$$

Definition 3.4.10

Der Schätzer $\bar{\theta}_{[\cdot]}$ bzw. $\widehat{\text{Var}}_{jn}(\hat{\theta})$ heißt *Jackknife-Schätzer* für den Erwartungswert bzw. die Varianz des Schätzers $\hat{\theta}$ von θ .

Beispiel 3.4.6

Sei $\theta = \mu$, $\hat{\theta} = \hat{\mu} = \bar{X}_n$, so gilt

$$\varphi_n(x_1, \dots, x_n) = \frac{1}{n} \sum_{i=1}^n x_i,$$

womit folgt, dass

$$\begin{aligned} \hat{\theta}_{[i]} &= \frac{1}{n-1} \sum_{j \neq i} X_j = \frac{1}{n-1} \left(-X_i + \sum_{j=1}^n X_j \right) = \frac{n}{n-1} \bar{X}_n - \frac{1}{n-1} X_i, \quad \forall i = 1, \dots, n, \\ \bar{\theta}_{[\cdot]} &= \frac{1}{n} \sum_{i=1}^n \hat{\theta}_{[i]} = \frac{n}{n-1} \bar{X}_n - \frac{1}{n(n-1)} \sum_{i=1}^n X_i = \frac{n \cdot \bar{X}_n}{n-1} - \frac{\bar{X}_n}{n-1} = \frac{n-1}{n-1} \bar{X}_n = \bar{X}_n. \end{aligned}$$

Daher ist ein *Jackknife-Schätzer* für μ gleich $\bar{X}_n$.

Konstruieren wir nun einen *Jackknife-Schätzer der Varianz*:

$$\begin{aligned} \widehat{\text{Var}}_{jn}(\hat{\theta}) &= \frac{n-1}{n} \sum_{i=1}^n \left(\frac{n}{n-1} \bar{X}_n - \frac{1}{n-1} X_i - \bar{X}_n \right)^2 = \frac{n-1}{n} \sum_{i=1}^n \left(\frac{1}{n-1} (\bar{X}_n - X_i) \right)^2 \\ &= \frac{n-1}{n(n-1)^2} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{1}{n} S_n^2, \end{aligned}$$

wobei dies genau der Plug-in Schätzer der Varianz von $\hat{\mu}$ ist.

- *Jackknife-Schätzer für die Verzerrung eines Schätzers*

Sei $\hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ . Der Bias von $\hat{\theta}$ ist $\mathbb{E}_\theta \hat{\theta} - \theta = \text{Bias}(\hat{\theta})$.

Definition 3.4.11

Ein *Jackknife-Schätzer der Verzerrung (Bias)* von $\hat{\theta}$ ist durch

$$\widehat{\text{Bias}}_{jn}(\hat{\theta}) = (n-1)(\bar{\theta}_{[\cdot]} - \hat{\theta})$$

gegeben.

An folgenden Beispielen wird klar, dass der oben beschriebene Vorgang zur Verringerung der Verzerrung beiträgt:

Der *Schätzer*

$$\tilde{\theta} = \hat{\theta} - \widehat{\text{Bias}}_{jn}(\hat{\theta}) = n\hat{\theta} - (n-1)\bar{\theta}_{[\cdot]} \quad (3.4.8)$$

hat in der Regel einen *kleineren Bias* als $\hat{\theta}$. Dabei ist wiederum

$$\hat{\theta}_{[i]} = \varphi_{n-1}(X_{[i]}) \quad \text{und} \quad \bar{\theta}_{[\cdot]} = \frac{1}{n} \sum_{i=1}^n \hat{\theta}_{[i]} \quad \text{mit} \quad \hat{\theta}(X_1, \dots, X_n) = \varphi_n(X_1, \dots, X_n) .$$

Beispiel 3.4.7

- a) Ist $\theta = \mathbb{E}X_i = \mu$, so ist $\hat{\theta} = \bar{X}_n$ ein unverzerrter Schätzer für μ . Was ist der Bias-korrigierte Schätzer $\tilde{\mu}$? (Dieser sollte schließlich nicht schlechter werden!)
- Es gilt $\bar{\theta}_{[\cdot]} = \bar{X}_n$, daher ist der Bias-Schätzer von Jackknife $\widehat{\text{Bias}}_{jn}(\hat{\theta}) = (n-1)(\bar{X}_n - \bar{X}_n) = 0$ und somit $\tilde{\theta} = \hat{\theta} - 0 = \bar{X}_n$. Wir haben also gesehen, dass die Jackknife-Methode die unverzerrten Schätzer (zumindest in diesem Beispiel) richtig behandelt, indem sie keinen zusätzlichen Bias einbaut.
- b) $\theta = \sigma^2 = \text{Var}X_i$, $\hat{\theta} = \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2$ ein verzerrter M-Schätzer der Varianz. Was ist $\tilde{\theta}$ in diesem Fall?

Übungsaufgabe 3.4.1

Zeigen Sie, dass $\tilde{\theta} = S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{n}{n-1} \hat{\sigma}^2$ ein erwartungstreuer Schätzer der Varianz ist. Somit wird der Bias von $\hat{\sigma}^2$ durch die Anwendung der Jackknife-Methode vollständig beseitigt.

Beweisidee: Zeigen Sie hierzu zunächst, dass

$$\widehat{\text{Bias}}_{jn}(\hat{\theta}) = -\frac{1}{n(n-1)} \sum_{i=1}^n (X_i - \bar{X}_n)^2 .$$

Bemerkung 3.4.5

Die Beispiele 3.4.7 a), b), in denen sich der Jackknife-Schätzer analytisch bestimmen ließ, sind eher eine Ausnahme als die Regel. In den meisten Fällen erfolgt die Bias-Reduktion mit Hilfe der Monte-Carlo-Methoden auf Basis der Formel (3.4.8).

2. Bootstrap-Schätzer:

Die Bootstrap-Methode besteht in dem Erzeugen einer neuen Stichprobe $(X_1^*, \dots, X_n^*)$, die aus einer approximativen Verteilung $\hat{F}$ der Stichprobenvariablen X_i , $i = 1, \dots, n$ gewonnen wird. Seien $\mathbb{E}_*$ und Var_* die wahrscheinlichkeitstheoretischen Größen, die auf dem Verteilungsgesetz $\mathbb{P}_*$ der neuen Stichprobe $(X_1^*, \dots, X_n^*)$ beruhen. Dabei gibt es folgende Möglichkeiten, $\hat{F}$ zu konstruieren:

- i) $\hat{F}(x) = \hat{F}_n(x)$ die empirische Verteilungsfunktion von X_i , falls X_i unabhängig identisch verteilt sind.
- ii) $\hat{F}$ ist ein parametrischer Schätzer von F , der parametrischen Verteilungsfunktion von X_i . Das heißt, falls $X_i \sim F_\theta$, $i = 1, \dots, n$ für ein $\theta \in \Theta$ und $\hat{\theta} = \hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ ist, so setzen wir $\hat{F} = F_{\hat{\theta}}$ (Plug-in Methode).

Definition 3.4.12

Ein *Bootstrap-Schätzer* für den Erwartungswert (bzw. *Bias* oder *Varianz*) von Schätzer $\hat{\theta}(X_1, \dots, X_n)$ ist gegeben durch

- a) $\hat{\mathbb{E}}_{boot}(\hat{\theta}) = \mathbb{E}_* \hat{\theta}(X_1^*, \dots, X_n^*)$.
- b) $\widehat{\text{Bias}}_{boot}(\hat{\theta}) = \hat{\mathbb{E}}_{boot} \hat{\theta} - \hat{\theta}$.
- c) $\widehat{\text{Var}}_{boot}(\hat{\theta}) = \text{Var}_*(\hat{\theta}(X_1^*, \dots, X_n^*))$.

Beispiel 3.4.8

Sei $\theta = \mu = \mathbb{E}X_i$ und $\hat{F} = \hat{F}_n$ die empirische Verteilungsfunktion. Wie generiert man eine Stichprobe $X_1^*, \dots, X_n^*$, wobei $X_i^* \sim \hat{F}_n$?

$\hat{F}_n$ gewichtet jede Beobachtung x_i der ursprünglichen Stichprobe mit dem Gewicht $1/n$, deshalb genügt es, einen der Einträge $(x_1, \dots, x_n)$ auszuwählen (mit Wahrscheinlichkeit $1/n$, Urnenmodell „Ziehen mit Zurücklegen“), um X_j^* , $j = 1, \dots, n$ zu generieren.

Bootstrap-Schätzer für den Erwartungswert von $\hat{\mu} = \bar{X}_n$:

$$\hat{\mathbb{E}}_{boot} \hat{\mu} = \mathbb{E}_* \left(\frac{1}{n} \sum_{i=1}^n X_i^* \right) \stackrel{X_i^* \text{ u.i.v.}}{=} \frac{1}{n} \cdot n \mathbb{E}_*(X_1^*) = \int x d\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}_n.$$

Somit folgt $\widehat{\text{Bias}}_{boot} \hat{\mu} = 0$.

$$\widehat{\text{Var}}_{boot}(\hat{\mu}) = \text{Var}_* \left(\frac{1}{n} \sum_{i=1}^n X_i^* \right) \stackrel{X_i^* \text{ u.i.v.}}{=} \frac{1}{n^2} \cdot n \cdot \text{Var}_*(X_1^*) = \frac{1}{n} \cdot \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{\hat{\sigma}^2}{n},$$

ein Plug-in Schätzer für $\text{Var} \bar{X}_n = \sigma^2/n$.

Monte-Carlo-Methoden zur numerischen Berechnung von Bootstrap-Schätzern:

Was kann man tun, wenn keine expliziten Formeln für z.B. $\widehat{\text{Var}}_{boot}(\hat{\theta})$ vorliegen (der Regelfall in der Statistik)?

Generiere M unabhängige Stichproben $(X_{i1}^*, \dots, X_{in}^*)$, $i = 1, \dots, M$ nach der Regel i) oder ii) mit Hilfe der Monte-Carlo-Simulation. Dann berechne

$$\hat{\theta}_i = \hat{\theta}(X_{i1}^*, \dots, X_{in}^*), \quad i = 1, \dots, M \quad \text{und setze} \quad \hat{\mathbb{E}}_{boot} \hat{\theta} \approx \frac{1}{M} \sum_{i=1}^M \hat{\theta}_i.$$

Ähnlich gewinnt man approximative Bootstrap-Schätzer für $\widehat{\text{Bias}} \hat{\theta}$ und $\widehat{\text{Var}} \hat{\theta}$:

$$\widehat{\text{Bias}}_{boot} \hat{\theta} \approx \hat{\mathbb{E}}_{boot} \hat{\theta} - \hat{\theta}, \quad \widehat{\text{Var}}_{boot} \hat{\theta} \approx \frac{1}{M-1} \sum_{i=1}^M (\hat{\theta}_i - \hat{\mathbb{E}}_{boot} \hat{\theta})^2.$$

Mehr sogar, man kann die Verteilungsfunktion von X_{ij}^* durch die empirische Verteilungsfunktion bestimmen:

$$\hat{F}_{boot}(x) = \frac{1}{M} \sum_{i=1}^M \frac{1}{n} \sum_{j=1}^n \mathbb{I}(X_{ij}^* \leq x), \quad x \in \mathbb{R}.$$

Ferner lassen sich mit Hilfe von oben genannten Methoden *Bootstrap-Konfidenzintervalle* für $\hat{\theta}$ ableiten:

Dafür lassen sich Quantile $\hat{F}_{\hat{\theta}}^{-1}(\alpha_1)$ und $\hat{F}_{\hat{\theta}}^{-1}(\alpha_2)$ der Verteilung von $\hat{\theta}(X_1^*, \dots, X_n^*)$ aus der Stichprobe $(\hat{\theta}_1, \dots, \hat{\theta}_M)$ empirisch bestimmen. Damit gilt

$$\mathbb{P}\left(\hat{F}_{\hat{\theta}}^{-1}(\alpha_1) \leq \hat{\theta}(X_1^*, \dots, X_n^*) \leq \hat{F}_{\hat{\theta}}^{-1}(\alpha_2)\right) \approx 1 - \alpha_1 - \alpha_2 = 1 - \alpha,$$

wobei $\alpha = \alpha_1 + \alpha_2$ klein ist. Beachte dabei, dass man hofft, dass X_i^* sehr ähnlich verteilt ist wie X_i und somit

$$\mathbb{P}\left(\hat{F}_{\hat{\theta}}^{-1}(\alpha_1) \leq \hat{\theta}(X_1, \dots, X_n) \leq \hat{F}_{\hat{\theta}}^{-1}(\alpha_2)\right) \approx 1 - \alpha_1 - \alpha_2 = 1 - \alpha$$

gilt.

3.5 Weitere Güteeigenschaften von Punktschätzern

3.5.1 Ungleichung von Cramér-Rao

Sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen identisch verteilten Zufallsvariablen X_i mit Verteilungsfunktion F_θ , $\theta \in \Theta$. Sei $\hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ . Falls $\hat{\theta}$ erwartungstreu ist, dann misst man die Güte eines anderen erwartungstreuen Schätzers $\tilde{\theta}$ von θ am Wert seiner Varianz. Das bedeutet, falls $\text{Var}_\theta \tilde{\theta} < \text{Var}_\theta \hat{\theta}$, dann ist der Schätzer $\tilde{\theta}$ besser. Wir werden uns nun mit der Frage befassen, ob immer wieder neue, bessere Schätzer $\tilde{\theta}$ mit immer kleinerer Varianz konstruiert werden können. Die Antwort hierauf ist unter gewissen Voraussetzungen negativ. Die untere Schranke der Varianz $\text{Var}_\theta \hat{\theta}$ hierzu liefert der Satz von Cramér-Rao.

Sei $L(x, \theta)$ die Likelihood-Funktion von X_i , d.h.

$$L(x, \theta) = \begin{cases} \mathbb{P}_\theta(x), & \text{im diskreten Fall,} \\ f_\theta(x), & \text{im stetigen Fall} \end{cases}$$

und $L(x_1, \dots, x_n, \theta) = \prod_{i=1}^n L(x_i, \theta)$ die Likelihood-Funktion von der gesamten Stichprobe $(X_1, \dots, X_n)$. Es gelten die Bedingungen 1) bis 5), die für die asymptotische Normalverteiltheit von ML-Schätzern auf Seite 80 gestellt wurden, wobei die Bedingung 5) für $k = 1$ gilt.

Satz 3.5.1 (Ungleichung von Cramér-Rao):

Sei $\hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ mit den folgenden Eigenschaften:

1. $\mathbb{E}_\theta \hat{\theta}^2(X_1, \dots, X_n) < \infty \quad \forall \theta \in \Theta$.
2. Für alle $\theta \in \Theta$ existiert

$$\frac{d}{d\theta} \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n) = \begin{cases} \int_{\mathbb{R}} \dots \int_{\mathbb{R}} \hat{\theta}(x_1, \dots, x_n) \frac{\partial}{\partial \theta} L(x_1, \dots, x_n, \theta) dx_1 \dots dx_n, & \text{im stetigen Fall,} \\ \sum_{x_1, \dots, x_n} \hat{\theta}(x_1, \dots, x_n) \frac{\partial}{\partial \theta} L(x_1, \dots, x_n, \theta), & \text{im disk. Fall.} \end{cases}$$

Dann gilt

$$\text{Var}_\theta \hat{\theta}(X_1, \dots, X_n) \geq \frac{\left(\frac{d}{d\theta} \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n)\right)^2}{n \cdot I(\theta)}, \quad \theta \in \Theta,$$

wobei $I(\theta)$ die Fisher-Information aus (3.4.3) ist.

Beweis Führen wir die Funktion

$$\varphi_\theta(x_1, \dots, x_n) = \frac{\partial}{\partial \theta} \log L(x_1, \dots, x_n, \theta)$$

ein. In Bemerkung 3.4.3 haben wir bewiesen, dass

$$\mathbb{E}_\theta \varphi_\theta(X_1, \dots, X_n) = 0, \quad \text{Var}_\theta \varphi_\theta(X_1, \dots, X_n) = n \cdot I(\theta).$$

Wenden wir die Ungleichung von Cauchy-Schwarz auf $\text{Cov}_\theta(\varphi_\theta(X_1, \dots, X_n), \hat{\theta}(X_1, \dots, X_n))$ an:

$$\begin{aligned} \text{Cov}_\theta(\varphi_\theta(X_1, \dots, X_n), \hat{\theta}(X_1, \dots, X_n)) &= \mathbb{E}_\theta(\varphi_\theta(X_1, \dots, X_n) \cdot \hat{\theta}(X_1, \dots, X_n)) - 0 \\ &\leq \sqrt{\text{Var}_\theta \varphi_\theta(X_1, \dots, X_n)} \sqrt{\text{Var}_\theta \hat{\theta}(X_1, \dots, X_n)} \end{aligned}$$

Somit folgt

$$\text{Var}_\theta \hat{\theta}(X_1, \dots, X_n) \geq \frac{\overbrace{\left(\mathbb{E}_\theta(\varphi_\theta(X_1, \dots, X_n) \cdot \hat{\theta}(X_1, \dots, X_n)) \right)^2}^{=:A}}{\text{Var}_\theta \varphi_\theta(X_1, \dots, X_n)} = \frac{A^2}{n \cdot I(\theta)}.$$

Es bleibt zu zeigen, dass

$$A = \frac{d}{d\theta} \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n).$$

Wir zeigen die Aussage für den absolut stetigen Fall (im diskreten Fall sind die Integrale durch Summen zu ersetzen):

$$\begin{aligned} A &= \int \frac{\partial}{\partial \theta} \log L(x_1, \dots, x_n, \theta) \cdot \hat{\theta}(x_1, \dots, x_n) \cdot L(x_1, \dots, x_n, \theta) dx_1 \dots dx_n \\ &= \int \frac{\partial}{\partial \theta} L(x_1, \dots, x_n, \theta) \cdot \hat{\theta}(x_1, \dots, x_n) dx_1 \dots dx_n \stackrel{\text{Vor. 2)}}{=} \frac{d}{d\theta} \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n). \end{aligned}$$

□

Folgerung 3.5.1

Falls $\hat{\theta}$ ein erwartungstreuer Schätzer für θ ist und die Voraussetzungen des Satzes 3.5.1 erfüllt sind, so gilt

$$\text{Var}_\theta \hat{\theta}(X_1, \dots, X_n) \geq \frac{1}{n \cdot I(\theta)}.$$

Beweis Wende die Ungleichung von Cramér-Rao an $\hat{\theta}$ mit

$$\frac{d}{d\theta} (\mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n)) = \frac{d}{d\theta} \theta = 1$$

an.

□

An folgenden Beispielen werden wir sehen, dass der Schätzer $\bar{X}_n$ des Erwartungswertes μ in der Klasse aller Schätzer für μ , die die Voraussetzungen des Satzes 3.5.1 erfüllen, die kleinste Varianz besitzt. Somit ist $\bar{X}_n$ der beste erwartungstreue Schätzer in dieser Klasse für mindestens zwei parametrische Familien von Verteilungen:

- Normalverteilung und
- Poisson-Verteilung.

Beispiel 3.5.1

1. $X_i \sim N(\mu, \sigma^2)$, $\hat{\mu} = \bar{X}_n$ als Schätzer für μ . Dabei ist $\hat{\mu}$ erwartungstreu mit $\text{Var} \hat{\mu} = \sigma^2/n$. Zeigen wir, dass die Cramér-Rao-Schranke für die Varianz eines erwartungstreuen Schätzers $\hat{\theta}$ für μ ebenso gleich σ^2/n ist. Prüfen wir zunächst die Voraussetzungen des Satzes 3.5.1:

Zeigen wir, dass

$$0 = \frac{d}{d\mu} \int_{\mathbb{R}} L(x, \mu) dx = \int_{\mathbb{R}} \frac{\partial}{\partial \mu} L(x, \mu) dx \quad \text{mit} \quad L(x, \mu) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} :$$

$$\begin{aligned} \frac{\partial}{\partial \mu} L(x, \mu) &= \frac{2(x-\mu)}{2\sigma^2} \cdot \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} = \frac{x-\mu}{\sigma^2} \cdot L(x, \mu), \\ \int_{\mathbb{R}} \frac{\partial}{\partial \mu} L(x, \mu) dx &= \mathbb{E} \left(\frac{X-\mu}{\sigma^2} \right) = 0. \end{aligned}$$

Zeigen wir weiterhin die Gültigkeit der Bedingung 2) des Satzes 3.5.1:

$$\frac{d}{d\mu} \mathbb{E} \bar{X}_n = \frac{d}{d\mu} (\mu) = 1 \stackrel{?}{=} \frac{1}{n} \int_{\mathbb{R}^n} (x_1 + \dots + x_n) \frac{\partial}{\partial \mu} \left(\prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x_i-\mu}{\sigma}\right)^2} \right) dx_1 \dots dx_n.$$

Induktion bzgl. n :

- Induktionsanfang $n = 1$:

$$\int_{\mathbb{R}} x \frac{\partial}{\partial \mu} L(x, \mu) dx = \int_{\mathbb{R}} \frac{x(x-\mu)}{\sigma^2} L(x, \mu) dx = \frac{1}{\sigma^2} \left(\mathbb{E}_{\mu} X^2 - \mu^2 \right) = \frac{\text{Var}_{\mu} X}{\sigma^2} = 1.$$

- Induktionshypothese: Für n gilt

$$\int_{\mathbb{R}^n} (x_1 + \dots + x_n) \cdot \frac{\partial}{\partial \mu} L(x_1, \dots, x_n, \mu) dx_1 \dots dx_n = n.$$

- Induktionsschritt $n \rightarrow n+1$:

$$A = \int_{\mathbb{R}^{n+1}} (x_1 + \dots + x_{n+1}) \frac{\partial}{\partial \mu} \underbrace{L(x_1, \dots, x_{n+1}, \mu)}_{=L(x_1, \dots, x_n, \mu) \cdot L(x_{n+1}, \mu)} dx_1 \dots dx_{n+1} \stackrel{?}{=} n+1.$$

Dabei gilt für A :

$$\begin{aligned}
A &= \int_{\mathbb{R}^{n+1}} (x_1 + \dots + x_n) \cdot \left(\frac{\partial}{\partial \mu} L(x_1, \dots, x_n, \mu) \cdot L(x_{n+1}, \mu) + L(x_1, \dots, x_n, \mu) \cdot \right. \\
&\quad \cdot \left. \frac{\partial}{\partial \mu} L(x_{n+1}, \mu) \right) dx_1 \dots dx_n dx_{n+1} + \int_{\mathbb{R}^{n+1}} x_{n+1} \left(\frac{\partial}{\partial \mu} L(x_1, \dots, x_n, \mu) \cdot \right. \\
&\quad \cdot L(x_{n+1}, \mu) + L(x_1, \dots, x_n, \mu) \cdot \left. \frac{\partial}{\partial \mu} L(x_{n+1}, \mu) \right) dx_1 \dots dx_n dx_{n+1} \\
&= n \cdot \underbrace{\int_{\mathbb{R}} L(x_{n+1}, \mu) dx_{n+1}}_{=1} + \int_{\mathbb{R}^n} (x_1 + \dots + x_n) \cdot L(x_1, \dots, x_n, \mu) dx_1 \dots dx_n \cdot \\
&\quad \cdot \underbrace{\int_{\mathbb{R}} \frac{\partial}{\partial \mu} L(x_{n+1}, \mu) dx_{n+1}}_{=0} + \int_{\mathbb{R}} x_{n+1} L(x_{n+1}, \mu) dx_{n+1} \cdot \\
&\quad \cdot \underbrace{\int_{\mathbb{R}^n} \frac{\partial}{\partial \mu} L(x_1, \dots, x_n, \mu) dx_1 \dots dx_n}_{=0} + \underbrace{\int_{\mathbb{R}} x_{n+1} \frac{\partial}{\partial \mu} L(x_{n+1}, \mu) dx_{n+1}}_{=\frac{d}{d\mu} \mathbb{E}_\mu X = \frac{d}{d\mu} \mu = 1} \cdot \\
&\quad \cdot \underbrace{\int_{\mathbb{R}^n} L(x_1, \dots, x_n, \mu) dx_1 \dots dx_n}_{=1} = n + 1.
\end{aligned}$$

Nachdem alle Voraussetzungen erfüllt sind, berechnen wir die Schranke

$$\frac{1}{n \cdot I(\mu)} \quad \text{mit} \quad I(\mu) = \mathbb{E}_\mu \left(\frac{\partial}{\partial \mu} \log L(X, \mu) \right)^2.$$

Es folgt aus dem Beispiel 3.4.3, dass

$$I(\mu) = \frac{1}{\sigma^2} \implies n \cdot I(\mu) = \frac{n}{\sigma^2}.$$

Insgesamt gilt also

$$\text{Var}_\mu \hat{\theta} \geq \frac{1}{\frac{n}{\sigma^2}} = \frac{\sigma^2}{n} = \text{Var}_\mu \bar{X}_n$$

für einen beliebigen erwartungstreuen Schätzer $\hat{\theta}$ für μ , der die Voraussetzungen des Satzes 3.5.1 erfüllt.

2. Das zweite Beispiel sei folgende Übungsaufgabe:

Übungsaufgabe 3.5.1

Seien $X_i \sim \text{Poisson}(\lambda)$, $i = 1, \dots, n$. Zeigen Sie, dass die Schranke von Cramér-Rao

$$\frac{1}{n \cdot I(\lambda)} = \frac{\lambda}{n} = \text{Var}_\lambda \bar{X}_n$$

ist. Dies bedeutet, dass auch hier $\bar{X}_n$ der beste erwartungstreue Schätzer ist, der die Voraussetzungen des Satzes 3.5.1 erfüllt.

An Hand des nächsten Beispiels wollen wir zeigen, dass die Konstruktion von Schätzern mit einer Varianz, die kleiner als die Cramér-Rao-Schranke ist, möglich ist, falls die Voraussetzungen von Satz 3.5.1 nicht erfüllt sind.

Beispiel 3.5.2

Seien $X_i \sim U[0, \theta]$, $\theta > 0$. Dann ist die Bedingung „ $\text{supp} f_\theta(x) = [0, \theta]$ unabhängig von θ “ verletzt und auch eine weitere Bedingung:

$$0 \neq \int_{\mathbb{R}} \frac{\partial}{\partial \theta} L(x, \theta) dx = \int_0^\theta \left(\frac{1}{\theta} \right)' dx = -\frac{1}{\theta^2} \cdot \theta = -\frac{1}{\theta}.$$

Sei $\hat{\theta}$ ein erwartungstreuer Schätzer für θ , so würde nach der Ungleichung von Cramér-Rao folgen, dass $\text{Var}_\theta \hat{\theta} \geq (n \cdot I(\theta))^{-1}$, wobei

$$I(\theta) = \mathbb{E} \left(\frac{\partial}{\partial \theta} \log L(X, \theta) \right)^2 = \int_0^\theta \frac{1}{\theta} \left(\frac{\partial}{\partial \theta} \log \left(\frac{1}{\theta} \right) \right)^2 dx = \frac{1}{\theta} \int_0^\theta dx \cdot \left(-\frac{1}{\theta} \right)^2 = \frac{1}{\theta^2}.$$

Damit hätten wir

$$\text{Var}_\theta \hat{\theta} \geq \frac{\theta^2}{n}.$$

Betrachten wir

$$\hat{\theta}(X_1, \dots, X_n) = \frac{n+1}{n} \max\{X_1, \dots, X_n\} = \frac{n+1}{n} X_{(n)}.$$

Zeigen wir, dass

$$\mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n) = \theta \quad \text{und} \quad \text{Var}_\theta \hat{\theta}(X_1, \dots, X_n) < \frac{\theta^2}{n}.$$

Berechnen wir dazu $\mathbb{E}_\theta X_{(n)}^k$, $k \in \mathbb{N}$. Es gilt

$$\begin{aligned} F_{X_{(n)}}(x) &= F_{X_i}^n(x) = \begin{cases} \frac{x^n}{\theta^n}, & x \in [0, \theta], \\ 1, & x \geq \theta, \\ 0, & x < 0, \end{cases} \\ f_{X_{(n)}}(x) &= F'_{X_{(n)}}(x) = \frac{nx^{n-1}}{\theta^n} \cdot \mathbb{I}(x \in [0, \theta]), \\ \mathbb{E}_\theta X_{(n)}^k &= \int_0^\theta x^k \frac{nx^{n-1}}{\theta^n} dx = \frac{n}{\theta^n} \int_0^\theta x^{n+k-1} dx = \frac{n \cdot \theta^{n+k}}{\theta^n \cdot (n+k)} = \frac{n\theta^k}{n+k}. \end{aligned}$$

Damit folgt

$$\mathbb{E}_\theta \hat{\theta} = \frac{n+1}{n} \cdot \mathbb{E}_\theta X_{(n)} = \frac{n+1}{n} \cdot \frac{n\theta}{n+1} = \theta,$$

das heißt, $\hat{\theta}$ ist erwartungstreu. Weiterhin gilt

$$\begin{aligned} \text{Var}_\theta \hat{\theta} &= \left(\frac{n+1}{n} \right)^2 \cdot \text{Var}_\theta X_{(n)} = \left(\frac{n+1}{n} \right)^2 \cdot \left(\frac{n\theta^2}{n+2} - \frac{n^2\theta^2}{(n+1)^2} \right) \\ &= \frac{(n+1)^2}{n^2} \cdot \frac{n(n+1)^2 - n^2(n+2)}{(n+2)(n+1)^2} \cdot \theta^2 \\ &= \frac{\theta^2}{n(n+2)} (n^2 + 2n + 1 - n^2 - 2n) = \frac{\theta^2}{n(n+2)} \end{aligned}$$

und somit

$$\text{Var}_\theta \hat{\theta} = \frac{\theta^2}{n(n+2)} < \frac{\theta^2}{n}.$$

3.5.2 Bedingte Erwartung

Seien X und Y zwei Zufallsvariablen, wobei Y eine absolut stetige Verteilung besitzt. Dann folgt $\mathbb{P}(Y = y) = 0 \quad \forall y \in \mathbb{R}$. Deshalb kann die bedingte Wahrscheinlichkeit $\mathbb{P}(X \in B | Y = y)$ auf dem gewöhnlichen Wege

$$\mathbb{P}(X \in B | Y = y) = \frac{\mathbb{P}(X \in B, Y = y)}{\mathbb{P}(Y = y)}$$

nicht definiert werden. Aus der Praxis ist aber eine Reihe von Fragestellungen bekannt (z.B. Bayessche Analyse), in denen Wahrscheinlichkeiten $\mathbb{P}(X \in B | Y = y)$ ausgewertet werden müssen. Deswegen werden wir eine neue Definition der bedingten Wahrscheinlichkeit geben, die solche Situationen berücksichtigt. Diese Definition erfolgt durch die Definition der bedingten Erwartung.

Schema:

1. Es wird die bedingte Erwartung von der Zufallsvariablen X bzgl. der σ -Algebra $\mathcal{B}$ als Zufallsvariable $\mathbb{E}(X|\mathcal{B})$ eingeführt, wobei $\mathcal{B}$ eine Teil- σ -Algebra von $\mathcal{F}$ und $(\Omega, \mathcal{F}, \mathbb{P})$ der Wahrscheinlichkeitsraum ist.
2. Die bedingte Erwartung von X unter der Bedingung Y wird als $\mathbb{E}(X|Y) = \mathbb{E}(X|\sigma_Y)$ eingeführt, wobei σ_Y die von Y erzeugte σ -Algebra ist.
3. $\mathbb{P}(X \in B | Y = y)$ wird als Zufallsvariable $\mathbb{E}(\mathbb{I}(X \in B) | Y)$ auf der Menge $\{\omega \in \Omega : Y(\omega) = y\}$ eingeführt.

Gehen wir nun dieses Schema im Detail durch:

1. Sei $(\Omega, \mathcal{F}, \mathbb{P})$ ein Wahrscheinlichkeitsraum und $\mathcal{B}$ eine Teil- σ -Algebra von $\mathcal{F}$, d.h. $\mathcal{B} \subseteq \mathcal{F}$.

Definition 3.5.1

Der *bedingte Erwartungswert* einer Zufallsvariablen X definiert auf dem Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P})$ bezüglich einer σ -Algebra $\mathcal{B} \subseteq \mathcal{F}$ ist in dem Fall $\mathbb{E}|X| < \infty$ als eine $\mathcal{B}$ -messbare Zufallsvariable Y definiert, die die Eigenschaft

$$\int_B Y(\omega) \mathbb{P}(d\omega) = \int_B X \mathbb{P}(d\omega), \quad \forall B \in \mathcal{B}$$

besitzt. Dabei wird die Bezeichnung $Y = \mathbb{E}(X|\mathcal{B})$ verwendet.

Warum existiert diese Zufallsvariable Y ?

- Zerlegen wir X in den positiven X_+ und negativen X_- Anteil $X = X_+ - X_-$ und beweisen die Existenz von $\mathbb{E}(X_\pm|\mathcal{B})$. Danach setzen wir $\mathbb{E}(X|\mathcal{B}) = \mathbb{E}(X_+|\mathcal{B}) - \mathbb{E}(X_-|\mathcal{B})$.
- Somit genügt es zu zeigen, dass der Erwartungswert $\mathbb{E}(X|\mathcal{B})$ einer nicht negativen Zufallsvariablen $X \geq 0$ fast sicher existiert.

- Sei $Q(B) = \int_B X(\omega) \mathbb{P}(d\omega)$. Man kann zeigen, dass $Q(\cdot)$ ein Maß auf $(\Omega, \mathcal{F})$ ist. Dabei folgt aus $\mathbb{P}(B) = 0$ die Gleichheit $Q(B) = 0$ für $B \in \mathcal{B}_{\mathbb{R}}$ (bzw. $B \in \mathcal{B}$). Somit ist Q absolut stetig bzgl. $\mathbb{P}$. Weiter existiert nach dem Satz von Radon-Nikodym eine Dichte $Y(\omega)$, die messbar bzgl. $\mathcal{B}$ ist und für die

$$Q(B) = \int_B Y(\omega) \mathbb{P}(d\omega) \implies Y(\omega) = \mathbb{E}(X|\mathcal{B})$$

gilt.

Bemerkung 3.5.1

Aus der obigen Beweisskizze wird ersichtlich, dass $Y(\omega) = \mathbb{E}(X|\mathcal{B})$ nur $\mathbb{P}$ -fast sicher definiert ist. Somit kann man mehrere Versionen von $Y(\omega)$ angeben, die sich auf einer Menge der Wahrscheinlichkeit 0 unterscheiden.

Satz 3.5.2 (Eigenschaften des bedingten Erwartungswertes):

Seien X und Y Zufallsvariablen auf dem Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P})$ mit der Eigenschaft $\mathbb{E}|X| < \infty$, $\mathbb{E}|Y| < \infty$ und $\mathbb{E}|XY| < \infty$ (dies kann noch ein wenig abgeschwächt werden, ist hier allerdings ausreichend). Seien $\mathcal{B}$, $\mathcal{B}_1$ und $\mathcal{B}_2$ Teil- σ -Algebren von $\mathcal{F}$. Es gelten folgende Eigenschaften (im fast sicheren Sinne):

- $\mathbb{E}(X|\{\emptyset, \Omega\}) = \mathbb{E}X$, $\mathbb{E}(X|\mathcal{F}) = X$ fast sicher.
- Falls $X \leq Y$ fast sicher, dann gilt ebenso $\mathbb{E}(X|\mathcal{B}) \leq \mathbb{E}(Y|\mathcal{B})$ fast sicher.
- Es gilt $\mathbb{E}(XY|\mathcal{B}) = X \cdot \mathbb{E}(Y|\mathcal{B})$, falls X $\mathcal{B}$ -messbar ist.
- $\mathbb{E}(c|\mathcal{B}) = c$ für $c = \text{const.}$
- Es gilt $\mathbb{E}(\mathbb{E}(X|\mathcal{B}_2)|\mathcal{B}_1) = \mathbb{E}(X|\mathcal{B}_1)$ und $\mathbb{E}(\mathbb{E}(X|\mathcal{B}_1)|\mathcal{B}_2) = \mathbb{E}(X|\mathcal{B}_1)$, falls $\mathcal{B}_1 \subseteq \mathcal{B}_2$.
- Falls X unabhängig von $\mathcal{B}$ ist (d.h., die σ -Algebren $\sigma_X = X^{-1}(\mathcal{B}_{\mathbb{R}})$ und $\mathcal{B}$ sind unabhängig), dann gilt $\mathbb{E}(X|\mathcal{B}) = \mathbb{E}X$.
- $\mathbb{E}(\mathbb{E}(X|\mathcal{B})) = \mathbb{E}X$.

Ohne Beweis (siehe Beweis in [26]).

Beispiel 3.5.3

Sei $\mathcal{B} = \sigma(\{A_1, \dots, A_n\})$, wobei $\{A_1, \dots, A_n\}$ eine messbare Zerlegung des Wahrscheinlichkeitsraumes $(\Omega, \mathcal{F}, \mathbb{P})$ ist, d.h. $\bigcup_{i=1}^n A_i = \Omega$, $A_i \cap A_j = \emptyset$, $i \neq j$, $\mathbb{P}(A_i) > 0$, $i = 1, \dots, n$. Was ist $\mathbb{E}(X|\mathcal{B})$? Da $\mathbb{E}(X|\mathcal{B})$ $\mathcal{B}$ -messbar ist, können wir die allgemeine Form der Funktionen ausnutzen, die messbar bzgl. einer endlich erzeugten σ -Algebra $\mathcal{B} = \sigma(\{A_1, \dots, A_n\})$ sind: $\mathbb{E}(X(\omega)|\mathcal{B}) = \sum_{i=1}^n k_i \mathbb{I}(\omega \in A_i)$ (ohne Beweis).

Berechnen wir k_i : Aus der Definition 3.5.1 folgt für $B = A_j$

$$\begin{aligned} \int_B \mathbb{E}(X|\mathcal{B}) \mathbb{P}(d\omega) &= \int_{A_j} \sum_{i=1}^n k_i \cdot \mathbb{I}(\omega \in A_i) \mathbb{P}(d\omega) = k_j \cdot \mathbb{P}(A_j) \\ &= \int_B X \mathbb{P}(d\omega) = \int_{A_j} X \mathbb{P}(d\omega) = \mathbb{E}(X \cdot \mathbb{I}_{A_j}) \\ &\implies k_j = \frac{\mathbb{E}(X \cdot \mathbb{I}_{A_j})}{\mathbb{P}(A_j)}, \quad j = 1, \dots, n. \\ &\implies \mathbb{E}(X(\omega)|\mathcal{B}) = \frac{\mathbb{E}(X \cdot \mathbb{I}_{A_j})}{\mathbb{P}(A_j)}, \quad \text{falls } \omega \in A_j, \quad j = 1, \dots, n. \end{aligned}$$

2. Bedingte Erwartung bzgl. einer Zufallsvariablen Y :**Definition 3.5.2**

Seien X und Y zwei Zufallsvariablen definiert auf dem Wahrscheinlichkeitsraum $(\Omega, \mathcal{F}, \mathbb{P})$. Der *bedingte Erwartungswert von X unter der Bedingung Y* wird als $\mathbb{E}(X|Y) = \mathbb{E}(X|\sigma_Y)$ eingeführt, wobei σ_Y die von Y erzeugte σ -Algebra ist: $\sigma_Y = Y^{-1}(\mathcal{B}_{\mathbb{R}})$.

Lemma 3.5.1

Es existiert genau eine Borel-messbare Funktion $g : \mathbb{R} \rightarrow \mathbb{R}$, für die gilt, dass $\mathbb{E}(X|Y) = g(Y)$ fast sicher (Ohne Beweis).

Daher wird die Schreibweise $\mathbb{E}(X|Y = y)$ als $g(y)$ verstanden: $\mathbb{E}(X|Y = y) = g(y)$ oder $\mathbb{E}(X|Y = y)$ ist der Wert von $\mathbb{E}(X|Y)$ auf der Menge $\{\omega \in \Omega : Y(\omega) = y\}$.

3. Bedingte Wahrscheinlichkeit bzgl. einer σ -Algebra bzw. einer Zufallsvariable.**Definition 3.5.3**

Die *bedingte Wahrscheinlichkeit von $A \in \mathcal{F}$ unter der Bedingung $\mathcal{B}$* ist gegeben durch $\mathbb{P}(A|\mathcal{B}) = \mathbb{E}(\mathbb{I}_A|\mathcal{B})$ fast sicher. Analog dazu definieren wir $\mathbb{P}(A|Y) = \mathbb{E}(\mathbb{I}_A|Y)$ für eine Zufallsvariable Y .

Bemerkung 3.5.2

Die so definierte Familie von Zufallsvariablen $\mathbb{P}(\cdot|\mathcal{B})$ erfüllen (fast sicher) nicht die Eigenschaften eines Maßes: Es gilt

$$0 \leq \mathbb{P}(A|\mathcal{B}) \leq 1, \quad \forall A \in \mathcal{F} \text{ fast sicher,}$$

aber die Eigenschaft der σ -Additivität

$$\mathbb{P}\left(\bigcup_{i=1}^{\infty} A_i \middle| \mathcal{B}\right) \stackrel{\text{f.s.}}{=} \sum_{i=1}^{\infty} \mathbb{P}(A_i|\mathcal{B})$$

für disjunkte $\{A_i\}$ hängt von der Version $\mathbb{P}(\cdot|\mathcal{B})$ ab. Das bedeutet, es existiert kein $M \in \mathcal{F} : \mathbb{P}(M) = 0$, so dass die obige Eigenschaft für alle $\omega \in M^C$ gilt.

3.5.3 Suffizienz

Sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen identisch verteilten Zufallsvariablen X_i mit Verteilungsfunktion F_θ , $\theta \in \Theta \subseteq \mathbb{R}^m$. Wenn man von der vollen Information $\{X_1 = x_1, \dots, X_n = x_n\}$ zum Schätzer $\hat{\theta}(X_1, \dots, X_n)$ des Parameters θ übergeht, dann entsteht durch die Abbildung

$$\hat{\theta} : \mathbb{R}^n \rightarrow \mathbb{R}^m, \quad m \ll n$$

ein Informationsverlust, weil man normalerweise $(X_1, \dots, X_n)$ nicht aus $\hat{\theta}(X_1, \dots, X_n)$ zurückrechnen kann. Die sogenannten *suffizienten* Schätzer minimieren diesen Informationsverlust im stochastischen Sinne:

Definition 3.5.4

1. Seien Zufallsvariablen $X_1, \dots, X_n$ und $\hat{\theta}(X_1, \dots, X_n)$ diskret verteilt. Ein Schätzer $\hat{\theta}$ des Parameters θ heißt *suffizient*, falls

$$\mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n \mid \hat{\theta}(X_1, \dots, X_n) = t)$$

nicht von θ abhängt, solange $x_1, \dots, x_n$ und t in den Trägern der Zähldichten von $(X_1, \dots, X_n)$ bzw. $\hat{\theta}(X_1, \dots, X_n)$ liegen.

2. Falls $X_1, \dots, X_n$ und $\hat{\theta}(X_1, \dots, X_n)$ absolut stetig verteilt sind, dann heißt der Schätzer $\hat{\theta}$ *suffizient* für θ , falls die Wahrscheinlichkeit

$$\mathbb{P}\left((X_1, \dots, X_n) \in B \mid \hat{\theta}(X_1, \dots, X_n) = t\right)$$

für beliebige $B \in \mathcal{B}_{\mathbb{R}^n}$ und $t \in \text{supp} f_{\hat{\theta}}$ nicht von $\theta \in \Theta$ abhängt, wobei $f_{\hat{\theta}}$ die Dichte von $\hat{\theta}$ ist.

Bemerkung 3.5.3

1. Betrachten wir im diskreten Fall die bedingte Likelihood-Funktion

$$L_{\hat{\theta}}(x_1, \dots, x_n, \theta) = \mathbb{P}_{\theta}\left(X_1 = x_1, \dots, X_n = x_n \mid \hat{\theta}(X_1, \dots, X_n) = t\right).$$

Aus Definition 3.5.4 folgt, dass wir keinen neuen ML-Schätzer für θ aus dieser bedingten Likelihood $L_{\theta}(x_1, \dots, x_n, \theta)$ gewinnen werden können, da sie nicht von θ abhängt. Das heißt, der Schätzer $\hat{\theta}$ enthält bereits die volle Information über θ , die man aus der Stichprobe $(x_1, \dots, x_n)$ gewinnen kann.

2. Falls $g : \mathbb{R}^m \rightarrow \mathbb{R}^m$ eine bijektive Borel-messbare Abbildung und $\hat{\theta}(X_1, \dots, X_n)$ ein suffizienter Schätzer von $\theta \in \Theta \subset \mathbb{R}^m$ ist, dann ist der Schätzer $g(\hat{\theta}(X_1, \dots, X_n))$ auch ein suffizienter Schätzer für θ . Dies wird aus der Tatsache ersichtlich, dass

$$\left\{\omega \in \Omega : g\left(\hat{\theta}(X_1, \dots, X_n)\right) = t\right\} = \left\{\omega \in \Omega : \hat{\theta}(X_1, \dots, X_n) = g^{-1}(t)\right\}, \quad \forall t.$$

Lemma 3.5.2 (Suffizienz):

Seien Zufallsvariablen $X_1, \dots, X_n$ und $\hat{\theta}(X_1, \dots, X_n)$ entweder alle diskret oder absolut stetig verteilt mit den Likelihood-Funktionen

$$L(x_1, \dots, x_n, \theta) = \begin{cases} \mathbb{P}_{\theta}(X_1 = x_1, \dots, X_n = x_n), & \text{im diskreten Fall,} \\ f_{X_1, \dots, X_n}(x_1, \dots, x_n, \theta), & \text{im absolut stetigen Fall,} \end{cases}$$

$$L_{\hat{\theta}}(t, \theta) = \begin{cases} \mathbb{P}_{\theta}(\hat{\theta}(X_1, \dots, X_n) = t), & \text{im diskreten Fall,} \\ f_{\hat{\theta}}(t, \theta), & \text{im absolut stetigen Fall.} \end{cases}$$

Sei der Träger von L gegeben durch

$$\text{supp} L = \{(x_1, \dots, x_n) \in \mathbb{R}^n : L(x_1, \dots, x_n, \theta) > 0\}$$

Der Schätzer $\hat{\theta}$ ist suffizient für θ genau dann, wenn

$$\frac{L(x_1, \dots, x_n, \theta)}{L_{\hat{\theta}}(\hat{\theta}(x_1, \dots, x_n), \theta)} \quad (3.5.1)$$

nicht von θ abhängig ist für alle $(x_1, \dots, x_n) \in \text{supp} L : \hat{\theta}(x_1, \dots, x_n) \in \text{supp} L_{\hat{\theta}}$.

Beweis Wir beweisen lediglich den diskreten Fall:

„ $\Rightarrow$ “ Ist $\hat{\theta}$ suffizient, so überprüfen wir, ob damit folgt, dass (3.5.1) von θ abhängt für alle $(x_1, \dots, x_n) \in \mathbb{R}$, $t \in \mathbb{R}$ und $\theta \in \Theta$, so dass $(x_1, \dots, x_n) \in \text{supp} L$. Es gilt:

$$\begin{aligned} \mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n \mid \hat{\theta}(X_1, \dots, X_n) = t) \\ = \frac{\mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n, \hat{\theta}(X_1, \dots, X_n) = t)}{\mathbb{P}_\theta(\hat{\theta}(X_1, \dots, X_n) = t)} \\ = \begin{cases} 0, & \text{falls } \hat{\theta}(x_1, \dots, x_n) \neq t \\ \frac{\mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n)}{\mathbb{P}_\theta(\hat{\theta}(X_1, \dots, X_n) = \hat{\theta}(x_1, \dots, x_n))}, & \text{falls } \hat{\theta}(x_1, \dots, x_n) = t. \end{cases} \end{aligned}$$

Somit hängt (3.5.1) nicht von θ ab.

„ $\Leftarrow$ “ Folgt aus dem 1. Fall durch Betrachtung von hinten.

□

Beispiel 3.5.4

1. *Bernoulli-Verteilung:* Seien $X_i \sim \text{Bernoulli}(p)$, $p \in [0, 1]$, $i = 1, \dots, n$, $\hat{p} = \bar{X}_n$ ein erwartungstreuer Schätzer für p . Wir zeigen nun, dass $\hat{p}$ suffizient ist. Es gilt

$$\hat{p} = \bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i = \frac{1}{n} Y,$$

wobei $Y \sim \text{Bin}(n, p)$. Es genügt nach Bemerkung 3.5.3 2) zu zeigen, dass Y ein suffizienter Schätzer für p ist. Es gilt für $x_i \in \{0, 1\}$, $i = 1, \dots, n$

$$\mathbb{P}(X_1 = x_1, \dots, X_n = x_n) = \prod_{i=1}^n p^{x_i} (1-p)^{1-x_i} = p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i}.$$

Definieren wir nun L_Y als

$$L_Y(y, p) = \binom{n}{y} p^y (1-p)^{n-y}, \quad y = 0, \dots, n.$$

Setzen wir nun statt y die Summe $\sum_{i=1}^n x_i$ ein und betrachten

$$\frac{L(x_1, \dots, x_n, p)}{L_Y(\sum_{i=1}^n x_i, p)} = \frac{p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i}}{\binom{n}{\sum_{i=1}^n x_i} p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i}} = \frac{1}{\binom{n}{\sum_{i=1}^n x_i}}.$$

Dies hängt offensichtlich nicht von p ab, womit folgt aus Lemma 3.5.2, dass Y und somit $\hat{p}$ suffizient sind.

2. *Normalverteilung mit bekannter Varianz:* Seien $X_i \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$, σ^2 bekannt. So ist $\hat{\mu} = \bar{X}_n$ ein erwartungstreuer Schätzer für μ . Zeigen wir nun, dass $\hat{\mu}$ suffizient ist: Betrachten wir

$$\begin{aligned} L(x_1, \dots, x_n, \mu) &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2} \left(\frac{x_i - \mu}{\sigma}\right)^2\right) \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \cdot \exp\left(-\frac{\sum_{i=1}^n (x_i - \mu)^2}{2\sigma^2}\right) \end{aligned}$$

und nach Lemma 2.2.1

$$= \frac{1}{(2\pi\sigma^2)^{n/2}} \cdot \exp\left(-\frac{\sum_{i=1}^n (x_i - \bar{x}_n)^2 + n(\bar{x}_n - \mu)^2}{2\sigma^2}\right).$$

Ferner gilt bekanntermaßen $\hat{\mu} \sim N(\mu, \sigma^2/n)$, und somit

$$\begin{aligned} L_{\hat{\mu}}(x, \mu) &= \frac{\sqrt{n}}{\sqrt{2\pi}\sigma} \cdot \exp\left(-\frac{n}{2} \left(\frac{x - \mu}{\sigma}\right)^2\right), \\ \frac{L(x_1, \dots, x_n, \mu)}{L_{\hat{\mu}}(\bar{x}_n, \mu)} &= \frac{\frac{1}{(2\pi\sigma^2)^{n/2}} \cdot \exp\left(-\frac{\sum_{i=1}^n (x_i - \bar{x}_n)^2 + n(\bar{x}_n - \mu)^2}{2\sigma^2}\right)}{\frac{\sqrt{n}}{\sqrt{2\pi}\sigma} \cdot \exp\left(-\frac{n(\bar{x}_n - \mu)^2}{2\sigma^2}\right)} \\ &= \frac{1}{\sqrt{n}(2\pi\sigma^2)^{n/2-1}} \cdot \exp\left(-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \bar{x}_n)^2\right), \end{aligned}$$

was von μ unabhängig ist. Somit folgt nach Lemma 3.5.2, dass $\hat{\mu} = \bar{X}_n$ ein suffizienter Schätzer für μ ist.

Mit Hilfe des nächsten Satzes von Neyman-Fisher wird es möglich sein zu zeigen, dass bei unbekannter Varianz der Schätzer $(\bar{X}_n, S_n^2)$ für (μ, σ^2) suffizient ist.

Satz 3.5.3 (Faktorisierungssatz von Neyman-Fisher):

Unter den Voraussetzungen von Lemma 3.5.2 ist $\hat{\theta}(X_1, \dots, X_n)$ ein suffizienter Schätzer für θ genau dann, wenn zwei messbare Funktionen $g: \mathbb{R}^n \times \Theta \rightarrow \mathbb{R}$ und $h: \mathbb{R}^n \rightarrow \mathbb{R}$ existieren, so dass folgende Faktorisierung der Likelihood-Funktion $L(x_1, \dots, x_n, \theta)$ der Stichprobe $(X_1, \dots, X_n)$ gilt:

$$L(x_1, \dots, x_n, \theta) = g(\hat{\theta}(x_1, \dots, x_n), \theta) \cdot h(x_1, \dots, x_n), \quad (x_1, \dots, x_n) \in \text{supp} L, \quad \theta \in \Theta.$$

Beweis Wir beweisen nur den diskreten Fall.

1. Falls $\hat{\theta}$ suffizient ist, dann hängt nach Lemma 3.5.2

$$\frac{L(x_1, \dots, x_n, \theta)}{\underbrace{L_{\hat{\theta}}(\hat{\theta}(x_1, \dots, x_n), \theta)}_{=g(\hat{\theta}(x_1, \dots, x_n), \theta)}} = h(x_1, \dots, x_n)$$

nicht von θ ab. Somit bekommen wir die Faktorisierung von Neyman-Fisher.

2. Sei nun $L(x_1, \dots, x_n, \theta) = g(\hat{\theta}(x_1, \dots, x_n), \theta) \cdot h(x_1, \dots, x_n)$ für alle $(x_1, \dots, x_n) \in \text{supp} L$, $\theta \in \Theta$. Führen wir eine Menge

$$C = \{(y_1, \dots, y_n) \in \mathbb{R}^n : \hat{\theta}(y_1, \dots, y_n) = \hat{\theta}(x_1, \dots, x_n)\} = \hat{\theta}^{-1}(\hat{\theta}(x_1, \dots, x_n))$$

ein. So gilt

$$\begin{aligned}
 \frac{\mathbb{P}_\theta(X_1 = x_1, \dots, X_n = x_n)}{\underbrace{L_\theta(\hat{\theta}(x_1, \dots, x_n), \theta)}_{=\mathbb{P}_\theta(\hat{\theta}(X_1, \dots, X_n) = \hat{\theta}(x_1, \dots, x_n))}} &= \frac{g(\hat{\theta}(x_1, \dots, x_n), \theta) \cdot h(x_1, \dots, x_n)}{\sum_{(y_1, \dots, y_n) \in C} \mathbb{P}_\theta(X_1 = y_1, \dots, X_n = y_n)} \\
 &= \frac{g(\hat{\theta}(x_1, \dots, x_n), \theta) \cdot h(x_1, \dots, x_n)}{\sum_{(y_1, \dots, y_n) \in C} \underbrace{g(\hat{\theta}(y_1, \dots, y_n), \theta)}_{=\hat{\theta}(x_1, \dots, x_n)} \cdot h(y_1, \dots, y_n)} \\
 &= \frac{h(x_1, \dots, x_n)}{\sum_{(y_1, \dots, y_n) \in C} h(y_1, \dots, y_n)},
 \end{aligned}$$

welches nicht von θ abhängt. Daher ist $\hat{\theta}$ nach Lemma 3.5.2 suffizient. □

Beispiel 3.5.5

1. *Poisson-Verteilung:* Seien $X_i \sim \text{Poisson}(\lambda)$, $\lambda > 0$, $\hat{\lambda} = \bar{X}_n$ ein erwartungstreuer Schätzer für λ . Zeigen wir, dass $\hat{\lambda}$ suffizient ist. Es gilt für $x_i \in \{0, 1, 2, \dots\}$, $i = 1, \dots, n$

$$\begin{aligned}
 L(x_1, \dots, x_n, \lambda) &= \prod_{i=1}^n e^{-\lambda} \frac{\lambda^{x_i}}{x_i!} = \frac{e^{-n\lambda} \cdot \lambda^{\sum_{i=1}^n x_i}}{x_1! \cdot \dots \cdot x_n!} = \frac{e^{-n\lambda} \lambda^{n\bar{x}_n}}{x_1! \cdot \dots \cdot x_n!}, \\
 &= g(\bar{x}_n, \lambda) \cdot h(x_1, \dots, x_n),
 \end{aligned}$$

wobei $g(\bar{x}_n, \lambda) = e^{-n\lambda} \cdot \lambda^{n\bar{x}_n}$, $h(x_1, \dots, x_n) = \frac{1}{x_1! \cdot \dots \cdot x_n!}$ ist. Somit ist $\hat{\lambda} = \bar{X}_n$ nach Satz 3.5.3 suffizient.

2. *Exponentialverteilung:* Seien $X_i \sim \text{Exp}(\lambda)$, $\lambda > 0$, $\hat{\lambda} = \bar{X}_n^{-1}$ ein Momentenschätzer für λ , der zwar nicht erwartungstreu ist, jedoch stark konsistent, denn $\bar{X}_n \xrightarrow[n \rightarrow \infty]{f.s.} \mathbb{E}X_i = \frac{1}{\lambda}$ nach dem starken Gesetz der großen Zahlen. Zeigen wir, dass $\hat{\lambda}$ suffizient ist. Für $x_1 \geq 0, \dots, x_n \geq 0$ gilt

$$\begin{aligned}
 L(x_1, \dots, x_n, \lambda) &= \prod_{i=1}^n \lambda e^{-\lambda x_i} = \lambda^n e^{-\lambda \sum_{i=1}^n x_i} = \lambda^n e^{-\lambda n \bar{x}_n} \\
 &= \lambda^n e^{-\frac{\lambda n}{\hat{\lambda}}} = g(\hat{\lambda}, \lambda) \cdot \underbrace{h(x_1, \dots, x_n)}_{=1},
 \end{aligned}$$

wobei $g(\hat{\lambda}, \lambda) = \lambda^n e^{-\frac{\lambda n}{\hat{\lambda}}}$ und $h(x_1, \dots, x_n) \equiv 1$ ist. Somit ist $\hat{\lambda}$ nach dem Satz 3.5.3 suffizient.

Übungsaufgabe 3.5.2

Zeigen Sie mit Hilfe des Satzes 3.5.3, dass der Schätzer $(\bar{X}_n, S_n^2)$ suffizient für (μ, σ^2) im Falle der normal und unabhängig identisch verteilten Stichprobe $(X_1, \dots, X_n)$, $X_i \sim N(\mu, \sigma^2)$ ist.

Bemerkung 3.5.4

Der Vorteil des Satzes von Neyman-Fisher ist, dass man für die Überprüfung der Suffizienzeigenschaft von $\hat{\theta}$ die Likelihood-Funktion von $\hat{\theta}$ nicht explizit zu kennen braucht. Dies ist insbesondere in den Fällen vorteilhaft, in denen der Schätzer $\hat{\theta}$ kompliziert ist und seine Likelihood-Funktion nicht analytisch angegeben werden kann (bzw. unbekannt ist).

3.5.4 Vollständigkeit**Definition 3.5.5**

Ein Schätzer $\hat{\theta}(X_1, \dots, X_n)$ des Parameters $\theta \in \Theta \subset \mathbb{R}^m$ heißt *vollständig*, falls für beliebige messbare Funktionen $g: \mathbb{R}^m \rightarrow \mathbb{R}$ mit der Eigenschaft $\mathbb{E}_\theta g(\hat{\theta}(X_1, \dots, X_n)) = 0$, $\theta \in \Theta$ folgt

$$g(\hat{\theta}(X_1, \dots, X_n)) \equiv 0. \quad P_\theta - \text{f.s. für alle } \theta \in \Theta.$$

Bemerkung 3.5.5

1. Seien $g_1, g_2: \mathbb{R}^m \rightarrow \mathbb{R}$ Funktionen, für die $\forall \theta \in \Theta$ gilt

$$\mathbb{E}_\theta \left| g_i(\hat{\theta}(X_1, \dots, X_n)) \right| < \infty, \quad \mathbb{E}_\theta g_1(\hat{\theta}(X_1, \dots, X_n)) = \mathbb{E}_\theta g_2(\hat{\theta}(X_1, \dots, X_n)),$$

wobei $\hat{\theta}$ vollständig ist. So folgt aus der Definition 3.5.5

$$g_1(\hat{\theta}(X_1, \dots, X_n)) = g_2(\hat{\theta}(X_1, \dots, X_n))$$

fast sicher (nehme $g = g_1 - g_2$).

Fazit: Die Eigenschaft der Vollständigkeit erlaubt aus dem Vergleich der Schätzer $g_1(\hat{\theta})$ und $g_2(\hat{\theta})$ im Mittel eine Aussage über ihre fast sichere Gleichheit zu machen.

2. Falls $\hat{\theta}$ ein vollständiger Schätzer für θ ist, dann ist auch $g(\hat{\theta})$ ein vollständiger Schätzer für θ für eine beliebige messbare Funktion $g: \mathbb{R}^m \rightarrow \mathbb{R}^m$.

Beispiel 3.5.6

1. *Bernoulli-Verteilung:* Seien $X_i \sim \text{Bernoulli}(p)$, $p \in [0, 1]$. Zeigen wir, dass $\hat{p} = \bar{X}_n$ vollständig ist:

Sei g eine beliebige Funktion $\mathbb{R} \rightarrow \mathbb{R}$. Es genügt zu zeigen, dass $Y = \sum_{i=1}^n X_i$ vollständig ist. Es gilt $Y \sim \text{Bin}(n, p)$, womit folgt, dass

$$\mathbb{E}_p g(Y) = \sum_{k=0}^n g(k) \binom{n}{k} p^k (1-p)^{n-k}.$$

Weiter gilt $\mathbb{E}_p g(Y) = 0$ genau dann, wenn

$$\sum_{k=0}^n g(k) \binom{n}{k} \left(\underbrace{\frac{p}{1-p}}_{=t} \right)^k = p_n(t) = 0$$

für $p \in (0, 1)$, also $t \in (0, \infty)$. $p_n(t)$ ist ein Polynom des Grades n , womit folgt

$$g(k) \binom{n}{k} = 0 \quad \text{für alle } k \quad \implies \quad g(k) = 0, \quad k = 0, \dots, n \quad \implies \quad g(Y) = 0 \quad \mathbb{P}_p\text{-fast sicher.}$$

Somit ist Y vollständig und daher auch $\hat{p} = \bar{X}_n$.

2. *Gleichverteilung:* Sei $X_i \sim U[0, \theta]$, $i = 1, \dots, n$. Wie wir bereits gezeigt haben, ist der Schätzer $\hat{\theta}(X_1, \dots, X_n) = \frac{n+1}{n} X_{(n)}$ erwartungstreu. Zeigen wir nun, dass er ein vollständiger Schätzer ist. Es genügt zu zeigen, dass $X_{(n)} = \max_{i=1, \dots, n} X_i$ vollständig ist. Es ist zu zeigen, dass für alle messbaren $g : \mathbb{R} \rightarrow \mathbb{R}$ aus $\mathbb{E}_\theta g(X_{(n)}) = 0$ folgt $g(X_{(n)}) = 0$ fast sicher. Die Dichte von $X_{(n)}$ ist nach Beispiel 3.5.2 gegeben durch $f_{X_{(n)}}(x) = \frac{nx^{n-1}}{\theta^n} \cdot \mathbb{I}_{[0, \theta]}(x)$.

$$\begin{aligned} 0 &= \frac{d}{d\theta} \mathbb{E}_\theta g(X_{(n)}) = \frac{d}{d\theta} \int_0^\theta g(x) f_{X_{(n)}}(x) dx = \frac{d}{d\theta} \frac{1}{\theta^n} \int_0^\theta nx^{n-1} g(x) dx \\ &= -n \frac{1}{\theta^{n+1}} \int_0^\theta g(x) nx^{n-1} dx + \frac{1}{\theta^n} n \theta^{n-1} g(\theta) = -\frac{n}{\theta} \underbrace{\mathbb{E}_\theta g(X_{(n)})}_{=0} + \frac{n}{\theta} g(\theta) \\ &= \frac{n}{\theta} g(\theta) = 0 \text{ für fast alle } \theta > 0 \implies g(x) = 0, \quad x > 0. \end{aligned}$$

Daher gilt $g(X_{(n)}) = 0$ fast sicher.

3.5.5 Bester erwartungstreuer Schätzer

Aus Definition 3.3.7 folgt: Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe, $X_i \sim F_\theta$, $\theta \in \Theta \subset \mathbb{R}$ ($m = 1$), X_i unabhängig identisch verteilte Zufallsvariablen. Dann heißt $\hat{\theta}(X_1, \dots, X_n)$ bester erwartungstreuer Schätzer, falls

$$\mathbb{E}_\theta \hat{\theta}^2(X_1, \dots, X_n) < \infty \quad \mathbb{E}_\theta \hat{\theta}(X_1, \dots, X_n) = \theta, \quad \theta \in \Theta.$$

und $\hat{\theta}$ die minimale Varianz unter allen erwartungstreuen Schätzern besitzt.

Lemma 3.5.3 (Eindeutigkeit der besten erwartungstreuen Schätzer):

Falls $\hat{\theta}$ ein bester erwartungstreuer Schätzer für θ ist, dann ist er eindeutig bestimmt.

Beweis Sei $\hat{\theta} = \hat{\theta}(X_1, \dots, X_n)$ ein bester erwartungstreuer Schätzer für θ und $\tilde{\theta}$ ein weiterer bester erwartungstreuer Schätzer für θ . Zeigen wir, dass $\hat{\theta} = \tilde{\theta}$.

Ex adverso: Nehmen wir an, dass $\hat{\theta} \neq \tilde{\theta}$ ist und betrachten $\theta^* = 1/2(\hat{\theta} + \tilde{\theta})$. Offensichtlich ist θ^* erwartungstreu. Untersuchen wir

$$\text{Var}_\theta \theta^* = \frac{1}{4} \text{Var}_\theta (\hat{\theta} + \tilde{\theta}) = \frac{1}{4} \text{Var}_\theta \hat{\theta} + \frac{1}{4} \text{Var}_\theta \tilde{\theta} + \frac{1}{2} \text{Cov}_\theta (\hat{\theta}, \tilde{\theta}).$$

Da $\hat{\theta}, \tilde{\theta}$ beste erwartungstreue Schätzer sind und mit der Ungleichung von Cauchy-Schwarz $|\text{Cov}_\theta (\hat{\theta}, \tilde{\theta})| \leq \sqrt{\text{Var}_\theta \hat{\theta} \cdot \text{Var}_\theta \tilde{\theta}} = \text{Var}_\theta \hat{\theta}$ gilt, folgt

$$\text{Var}_\theta \theta^* \leq \frac{1}{2} \text{Var}_\theta \hat{\theta} + \frac{1}{2} \text{Var}_\theta \hat{\theta} = \text{Var}_\theta \hat{\theta}.$$

Da $\hat{\theta}$ der beste erwartungstreue Schätzer ist folgt $\text{Var}_\theta \theta^* = \text{Var}_\theta \hat{\theta}$ und somit $\varrho(\hat{\theta}, \tilde{\theta}) = 1 \implies \hat{\theta}$ und $\tilde{\theta}$ sind linear abhängig, d.h. es existieren Konstanten a und b , für die gilt $\hat{\theta} = a\tilde{\theta} + b$. Es folgt $a = 1$ aus $\text{Var}_\theta \hat{\theta} = a^2 \text{Var}_\theta \tilde{\theta} = \text{Var}_\theta \hat{\theta}$ und $b = 0$, weil $\hat{\theta}$ und $\tilde{\theta}$ erwartungstreu sind: $\theta = \mathbb{E}_\theta \hat{\theta} = \mathbb{E}_\theta \tilde{\theta} + b = \theta + b$. Das bedeutet, dass $\hat{\theta} = \tilde{\theta}$. $\square$

Lemma 3.5.4

Ein erwartungstreuer Schätzer $\hat{\theta}$, dessen zweites Moment endlich ist, ist genau dann der beste erwartungstreue Schätzer für θ , wenn $\text{Cov}_\theta (\hat{\theta}, \varphi) = 0$, $\theta \in \Theta$ für eine beliebige Stichprobenfunktion $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}$ mit der Eigenschaft $\mathbb{E}_\theta \varphi(X_1, \dots, X_n) = 0$, $\forall \theta \in \Theta$.

Beweis Wir beweisen den Satz für beide Richtungen getrennt:

„ $\Rightarrow$ “ Sei $\hat{\theta}$ der beste erwartungstreue Schätzer für θ , $\varphi(X_1, \dots, X_n)$ eine Stichprobenfunktion mit $\mathbb{E}_\theta \varphi(X_1, \dots, X_n) = 0$, $\forall \theta \in \Theta$. So ist zu zeigen, dass $\text{Cov}_\theta(\hat{\theta}, \varphi) = \mathbb{E}_\theta(\hat{\theta}\varphi) = 0$, $\theta \in \Theta$ gilt.

Definieren wir $\tilde{\theta} = \hat{\theta} + a\varphi$, $a \in \mathbb{R}$. Berechnen wir

$$\text{Var}_\theta \tilde{\theta} = \text{Var}_\theta \hat{\theta} + a^2 \text{Var}_\theta \varphi + 2a \text{Cov}_\theta(\hat{\theta}, \varphi)$$

für $a \in \mathbb{R}$. Sei $g(a) = a^2 \text{Var}_\theta \varphi + 2a \text{Cov}_\theta(\varphi, \hat{\theta})$. Falls $\text{Cov}_\theta(\varphi, \hat{\theta}) \neq 0$, dann existiert ein $a \in \mathbb{R}$ mit $g(a) < 0$. Da $\tilde{\theta}$ ein erwartungstreuer Schätzer für θ ist ($\mathbb{E}_\theta \tilde{\theta} = \mathbb{E}_\theta \hat{\theta} + a \mathbb{E}_\theta \varphi = \theta + 0 = \theta$) folgt $\text{Var}_\theta \tilde{\theta} \geq \text{Var}_\theta \hat{\theta}$ für alle $a \in \mathbb{R}$. Dies ist jedoch ein Widerspruch mit $g(a) < 0$ für ein $a \in \mathbb{R}$. Damit folgt $\text{Cov}_\theta(\varphi, \hat{\theta}) = 0$, $\theta \in \Theta$.

„ $\Leftarrow$ “ Sei $\hat{\theta}$ erwartungstreu, $\mathbb{E}_\theta \hat{\theta}^2 < \infty$, $\theta \in \Theta$, $\text{Cov}_\theta(\varphi, \hat{\theta}) = 0$, $\theta \in \Theta$, falls $\mathbb{E}_\theta \varphi = 0$, $\theta \in \Theta$. Sei $\tilde{\theta}$ ein anderer erwartungstreuer Schätzer für θ . Zeigen wir, dass $\text{Var}_\theta \tilde{\theta} \geq \text{Var}_\theta \hat{\theta}$. Es gilt

$$\tilde{\theta} = \hat{\theta} + \underbrace{(\tilde{\theta} - \hat{\theta})}_{=: \varphi}, \quad \mathbb{E}_\theta \varphi = \mathbb{E}_\theta \tilde{\theta} - \mathbb{E}_\theta \hat{\theta} = \theta - \theta = 0, \quad \forall \theta \in \Theta.$$

Somit

$$\text{Var}_\theta \tilde{\theta} = \text{Var}_\theta \hat{\theta} + \underbrace{\text{Var}_\theta \varphi}_{\geq 0} + 2 \underbrace{\text{Cov}_\theta(\hat{\theta}, \varphi)}_{=0} \geq \text{Var}_\theta \hat{\theta},$$

woraus folgt, dass $\hat{\theta}$ der beste Erwartungstreuer Schätzer für θ ist. □

Satz 3.5.4 (Lehmann-Scheffé):

Sei $\hat{\theta}$ ein erwartungstreuer vollständiger und suffizienter Schätzer für θ , $\mathbb{E}_\theta \hat{\theta}^2 < \infty$, $\forall \theta \in \Theta$. Dann ist $\hat{\theta}$ der beste erwartungstreue Schätzer für θ .

Beweis Nach Lemma 3.5.4 ist zu zeigen, dass $\text{Cov}_\theta(\hat{\theta}, \varphi) = \mathbb{E}_\theta(\hat{\theta}\varphi) = 0$, $\theta \in \Theta$, falls $\mathbb{E}_\theta \varphi = 0$, $\theta \in \Theta$. Es ist

$$\mathbb{E}_\theta(\hat{\theta}\varphi) = \mathbb{E}_\theta(\mathbb{E}(\hat{\theta}\varphi|\hat{\theta})) \stackrel{\hat{\theta} \text{ } \sigma(\hat{\theta})\text{-messbar}}{=} \mathbb{E}_\theta(\hat{\theta} \cdot \mathbb{E}_\theta(\varphi|\hat{\theta})) = \mathbb{E}_\theta(\hat{\theta} \cdot g(\hat{\theta})) \stackrel{?}{=} 0,$$

falls $g(\hat{\theta}) = 0$ fast sicher. Da $\hat{\theta}$ suffizient ist, ist $g(t) = \mathbb{E}_\theta(\varphi|\hat{\theta} = t)$ unabhängig von θ . Betrachten wir $\mathbb{E}_\theta g(\hat{\theta})$. Wir wollen zeigen, dass $\mathbb{E}_\theta g(\hat{\theta}) = 0$, $\theta \in \Theta$. Daraus und aus der Vollständigkeit von $\hat{\theta}$ wird folgen, dass $g(\hat{\theta}) = 0$ fast sicher für alle $\theta \in \Theta$.

$$\mathbb{E}_\theta g(\hat{\theta}) = \mathbb{E}_\theta(\mathbb{E}_\theta(\varphi|\hat{\theta})) = \mathbb{E}_\theta \varphi = 0$$

nach Voraussetzung. Somit folgt $\mathbb{E}_\theta(\varphi\hat{\theta}) = 0$ und $\hat{\theta}$ ist unkorreliert mit φ : $\mathbb{E}_\theta \varphi = 0$, $\theta \in \Theta$, womit folgt, dass nach Lemma 3.5.4 $\hat{\theta}$ der beste erwartungstreue Schätzer ist. □

Satz 3.5.5

Sei $\hat{\theta}$ ein erwartungstreuer Schätzer für θ , $\mathbb{E}_\theta \hat{\theta}^2 < \infty$, $\theta \in \Theta$. Sei $\tilde{\theta}$ ein vollständiger und suffizienter Schätzer für θ . Dann ist der Schätzer $\theta^* = \mathbb{E}(\hat{\theta}|\tilde{\theta})$ der beste erwartungstreue Schätzer für θ .

Beweis 1. Zeigen wir, dass $\mathbb{E}_\theta \theta^{*2} < \infty \quad \forall \theta \in \Theta$. Es gilt

$$\mathbb{E}_\theta (\theta^{*2}) = \mathbb{E}_\theta \left(\mathbb{E}(\hat{\theta} | \tilde{\theta}) \right)^2 \leq \mathbb{E}_\theta \left(\mathbb{E}(\hat{\theta}^2 | \tilde{\theta}) \right) = \mathbb{E}_\theta \hat{\theta}^2 < \infty,$$

da mit der Ungleichung von Jensen für bedingte Erwartung gilt

$$f(\mathbb{E}(X | \mathcal{B})) \stackrel{\text{f.s.}}{\leq} \mathbb{E}(f(X) | \mathcal{B})$$

für jede Zufallsvariable X , σ -Algebra $\mathcal{B}$ und konvexe Funktion f .

2. Zeigen wir, dass θ^* erwartungstreu ist: $\mathbb{E}_\theta \theta^* = \mathbb{E}_\theta(\mathbb{E}(\hat{\theta} | \tilde{\theta})) = \mathbb{E}_\theta \hat{\theta} = \theta$, $\theta \in \Theta$, weil $\hat{\theta}$ erwartungstreu ist.
3. Nach Lemma 3.5.4 genügt es zu zeigen, dass $\mathbb{E}_\theta(\theta^* \varphi) = 0$ für $\theta \in \Theta$, falls $\mathbb{E}_\theta \varphi = 0$, $\theta \in \Theta$.

$$\begin{aligned} \mathbb{E}_\theta(\theta^* \varphi) &= \mathbb{E}_\theta \left(\underbrace{\mathbb{E}(\hat{\theta} | \tilde{\theta})}_{=g(\tilde{\theta}), \tilde{\theta} \text{ suf.}} \varphi \right) = \mathbb{E}_\theta(g(\tilde{\theta}) \varphi) = \mathbb{E}_\theta(\mathbb{E}(g(\tilde{\theta}) \varphi | \tilde{\theta})) \\ &\stackrel{g(\tilde{\theta}) \tilde{\theta}\text{-messbar}}{=} \mathbb{E}_\theta(g(\tilde{\theta}) \cdot \underbrace{\mathbb{E}(\varphi | \tilde{\theta})}_{=g_1(\tilde{\theta})}) = 0, \end{aligned}$$

falls $g_1(\tilde{\theta}) \stackrel{\text{f.s.}}{=} 0, \theta \in \Theta$. Zeigen wir, dass $\mathbb{E}_\theta g_1(\tilde{\theta}) = 0$. Es gilt $\mathbb{E}_\theta g_1(\tilde{\theta}) = \mathbb{E}_\theta(\mathbb{E}(\varphi | \tilde{\theta})) = \mathbb{E}_\theta \varphi = 0$ nach Voraussetzung. Daraus und aus der Vollständigkeit von $\tilde{\theta}$ folgt genauso wie im Beweis des Satzes 3.5.4, dass $g_1(\tilde{\theta}) = 0$ fast sicher. $\square$

Lemma 3.5.5 (Ungleichung von Blackwell-Rao):

Sei $\hat{\theta}$ ein erwartungstreuer Schätzer für θ , $\mathbb{E}_\theta \hat{\theta}^2 < \infty, \theta \in \Theta$. Sei $\tilde{\theta}$ ein suffizienter Schätzer für θ . Dann besitzt der erwartungstreu Schätzer $\theta^* := \mathbb{E}(\hat{\theta} | \tilde{\theta})$ eine Varianz, die kleiner oder gleich als $\text{Var}_\theta \hat{\theta}$ ist.

Beweis Siehe Beweis des Satzes 3.5.5. Dabei folgt die Erwartungstreue von θ^* aus Beweispunkt 2) des Satzes 3.5.5 und $\text{Var}_\theta \theta^* = \mathbb{E}_\theta \theta^{*2} - \theta^2 \leq \mathbb{E}_\theta \hat{\theta}^2 - \theta^2 = \text{Var}_\theta \hat{\theta}$ aus Beweispunkt 1) des Satzes 3.5.5. $\square$

Bemerkung 3.5.6

Die Suffizienz $\tilde{\theta}$ kommt im Beweis des Lemmas 3.5.5 explizit nicht vor. Dennoch ist sie notwendig, damit der Schätzer $\theta^* = \mathbb{E}(\hat{\theta} | \tilde{\theta}) = g(\tilde{\theta})$ nicht von θ abhängt.

Folgerung 3.5.2

Falls $\hat{\theta}$ ein vollständiger und suffizienter Schätzer für θ ist und falls eine Funktion $g : \mathbb{R} \rightarrow \mathbb{R}$ existiert, dass $\mathbb{E}_\theta g(\hat{\theta}) = \theta \quad \forall \theta \in \Theta$, dann ist $g(\hat{\theta})$ der beste erwartungstreue Schätzer für θ .

Beweis $g(\hat{\theta}) = \mathbb{E}(g(\hat{\theta}) | \hat{\theta})$, welcher nach Satz 3.5.5 der beste erwartungstreue Schätzer ist. $\square$

4 Konfidenzintervalle

4.1 Einführung

Konfidenz- oder Vertrauensintervalle wurden bereits in Kapitel 3 exemplarisch behandelt (vgl. Folgerung 3.3.2 und Bemerkung 3.3.4). In diesem Kapitel werden wir eine formale Definition eines Konfidenzintervalles angeben, um Vertrauensintervalle in größerer Tiefe studieren zu können. Dabei werden sowohl *Ein-* als auch *Zweistichprobenprobleme* behandelt.

Rufen wir uns die Annahmen eines parametrischen Modells in Erinnerung: es sei eine Stichprobe $(X_1, \dots, X_n)$ von unabhängigen, identisch verteilten Zufallsvariablen mit $X_i \sim F_\theta$ gegeben, wobei F_θ eine Verteilungsfunktion aus einer parametrischen Familie von Verteilungen $\{F_\theta : \theta \in \Theta\}$, $\Theta \subset \mathbb{R}$ ist.

Die Punktschätzer von θ liefern jeweils einen Wert für den Parametervektor. Es wäre allerdings auch vorteilhaft, die Genauigkeit solcher Schätzansätze zu nennen, das heißt, einen Bereich anzugeben, in dem θ mit hoher Wahrscheinlichkeit $1 - \alpha$ liegt. Dabei heißt α *Irrtumswahrscheinlichkeit*; übliche Werte für α sind $\alpha = 0,01; 0,05; 0,1$. Die Wahrscheinlichkeit $1 - \alpha$, daß θ für $m = 1$ im vorgegebenen *Konfidenzintervall* liegt, heißt dann *Überdeckungswahrscheinlichkeit* oder *Konfidenzniveau* und soll dann entsprechend hoch ausfallen, z.B. $0,99; 0,95; 0,9$.

Definition 4.1.1

Es sei $1 - \alpha$ ein Konfidenzniveau und $\underline{\theta} : \mathbb{R}^n \rightarrow \overline{\mathbb{R}} = \mathbb{R} \cup \{\pm\infty\}$, $\bar{\theta} : \mathbb{R}^n \rightarrow \overline{\mathbb{R}}$ zwei Stichprobenfunktionen mit der Eigenschaft

$$\underline{\theta}(x_1, \dots, x_n) \leq \bar{\theta}(x_1, \dots, x_n) \quad \forall (x_1, \dots, x_n) \in \mathbb{R}^n.$$

Falls

1. $P_\theta \left(\theta \in [\underline{\theta}(X_1, \dots, X_n), \bar{\theta}(X_1, \dots, X_n)] \right) \geq 1 - \alpha, \quad \theta \in \Theta$
2. $\inf_{\theta \in \Theta} P_\theta \left(\theta \in [\underline{\theta}(X_1, \dots, X_n), \bar{\theta}(X_1, \dots, X_n)] \right) = 1 - \alpha$
3. $\lim_{n \rightarrow \infty} P_\theta \left(\theta \in [\underline{\theta}(X_1, \dots, X_n), \bar{\theta}(X_1, \dots, X_n)] \right) = 1 - \alpha, \quad \theta \in \Theta$

dann heißt $I = [\underline{\theta}(X_1, \dots, X_n), \bar{\theta}(X_1, \dots, X_n)]$ ein

1. *Konfidenzintervall*
2. *minimales Konfidenzintervall*
3. *asymptotisches Konfidenzintervall*

zum Konfidenzniveau $1 - \alpha$. Dabei heißt $l_\theta(X_1, \dots, X_n) = \bar{\theta}(X_1, \dots, X_n) - \underline{\theta}(X_1, \dots, X_n)$ die *Länge* des Konfidenzintervalls. Es ist erwünscht, möglichst kleine Konfidenzintervalle (mit minimaler Länge) bei großem Konfidenzniveau für θ zu konstruieren.

Wie bereits bei den Beispielen von Kapitel 3 ersichtlich ist, folgt die Konstruktion eines Konfidenzintervalls einem bestimmten Muster, das wir jetzt genauer studieren werden:

1. Finde eine Statistik $T(X_1, \dots, X_n, \theta)$, die
 - vom Parameter θ abhängt und
 - eine bekannte (Prüf-) Verteilung F besitzt (möglicherweise asymptotisch für $n \rightarrow \infty$).
2. Bestimme die Quantile $F^{-1}(\alpha_1)$ und $F^{-1}(1 - \alpha_2)$ von der Verteilung F für Niveaus α_1 und $1 - \alpha_2$, sodaß $\alpha_1 + \alpha_2 = \alpha$.
3. Löse (falls möglich) die Ungleichung $F^{-1}(\alpha_1) \leq T(X_1, \dots, X_n, \theta) \leq F^{-1}(1 - \alpha_2)$ bzgl. θ auf. Das entsprechende Ergebnis $I = [T^{-1}(F^{-1}(\alpha_1)), T^{-1}(F^{-1}(1 - \alpha_2))]$ (im Falle einer monoton in θ steigenden Statistik T) ist ein Konfidenzintervall für θ zum Niveau $1 - \alpha$, denn es gilt

$$\begin{aligned}
 \mathbb{P}_\theta(\theta \in I) &= \mathbb{P}_\theta\left(T_\theta^{-1}(F^{-1}(\alpha_1)) \leq \theta \leq T^{-1}(F^{-1}(1 - \alpha_2))\right) \\
 &= \mathbb{P}_\theta\left(F^{-1}(\alpha_1) \leq T_\theta(X_1, \dots, X_n, \theta) \leq F^{-1}(1 - \alpha_2)\right) \\
 &= F(F^{-1}(1 - \alpha_2)) - F(F^{-1}(\alpha_1)) \\
 &= 1 - \alpha_2 - \alpha_1 \\
 &= 1 - \alpha \text{ für alle } \theta \in \Theta.
 \end{aligned}$$

Für asymptotische Konfidenzintervalle soll überall noch $\lim_{n \rightarrow \infty}$ geschrieben werden:

$\lim_{n \rightarrow \infty} \mathbb{P}_\theta(\theta \in I) = \dots = 1 - \alpha$. Hierbei ist T_θ^{-1} die Inverse von $T(X_1, \dots, X_n, \theta)$ bezüglich θ . Grafisch kann dies auf Abb. 4.1 veranschaulicht werden.

Definition 4.1.2

1. Falls $\alpha_1 = \alpha_2 = \alpha/2$, dann heißt das Konfidenzintervall $I = [T^{-1}(F^{-1}(\frac{\alpha}{2})), T^{-1}(F^{-1}(1 - \frac{\alpha}{2}))]$ *symmetrisch*.
2. Falls $\alpha_1 = 0$ (bzw. $\underline{\theta}(X_1, \dots, X_n) = -\infty$), dann heißt das Konfidenzintervall $[-\infty, \bar{\theta}(X_1, \dots, X_n)]$ *einseitig*. Das selbe gilt für $\alpha_2 = 0$ (bzw. $\bar{\theta}(X_1, \dots, X_n) = +\infty$) und das Vertrauensintervall $[\underline{\theta}(X_1, \dots, X_n), +\infty)$.

In der Zukunft werden wir oft, ohne Beschränkung der Allgemeinheit, symmetrische Konfidenzintervalle konstruieren, obwohl man auch ein allgemeineres, nicht-symmetrisches Intervall leicht angeben kann.

Bemerkung 4.1.1

Man sieht leicht, daß der Algorithmus zur Konstruktion eines Vertrauensbereiches sich sehr dem eines statistischen Tests ähnelt. Im letzten Fall heißt $T(X_1, \dots, X_n)$ *Teststatistik*. Im Allgemeinen kann man für jedes Konfidenzintervall einen entsprechenden statistischen Test angeben, aber nicht umgekehrt. In der Vorlesung Stochastik III werden wir einige Beispiele dieser Übertragung „Konfidenzintervall $\mapsto$ Test“ sehen.

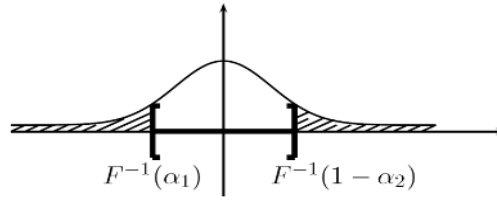


Abb. 4.1: asymptotisches Konfidenzintervall

4.2 Ein-Stichproben-Probleme

In diesem Abschnitt werden wir einige Beispiele von Vertrauensbereichen für Parameter einiger bekannter Verteilungen nach dem oben genannten Schema konstruieren. Dabei werden wir immer mit einer Stichprobe $(X_1, \dots, X_n)$ wie in Abschnitt 4.1 arbeiten.

4.2.1 Normalverteilung

Es seien $X_1, \dots, X_n$ unabhängig, identisch verteilt, mit $X_i \sim N(\mu, \sigma^2)$.

Konfidenzintervalle für den Erwartungswert μ

- **bei bekannter Varianz σ^2** Wenn wir annehmen, daß σ^2 bekannt ist, so ermöglicht uns der Satz 3.3.1, 4., ein exaktes Konfidenzintervall für μ zum Niveau $1 - \alpha$ zu berechnen. Denn es gilt $\bar{X}_n \sim N(\mu, \sigma^2/n)$ und somit

$$T(X_1, \dots, X_n, \mu) = \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \sim N(0, 1)$$

Es seien z_{α_1} und $z_{1-\alpha_2}$ Quantile der $N(0, 1)$ -Verteilung, $\alpha_1 + \alpha_2 = \alpha$ und $1 - \alpha$ das vorgegebene Konfidenzniveau.

Dann gilt

$$\begin{aligned} 1 - \alpha &= \mathbb{P}(z_{\alpha_1} \leq T(X_1, \dots, X_n, \mu) \leq z_{1-\alpha_2}) \\ &= \mathbb{P}\left(z_{\alpha_1} \leq \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \leq z_{1-\alpha_2}\right) \\ &\stackrel{(-z_{\alpha_1} = z_{1-\alpha_1})}{=} \mathbb{P}\left(\bar{X}_n - \frac{z_{1-\alpha_2}\sigma}{\sqrt{n}} \leq \mu \leq \bar{X}_n + \frac{z_{1-\alpha_1}\sigma}{\sqrt{n}}\right). \end{aligned}$$

Somit ist $[\underline{\theta}(X_1, \dots, X_n), \bar{\theta}(X_1, \dots, X_n)]$ mit $\underline{\theta}(X_1, \dots, X_n) = \bar{X}_n - z_{1-\alpha_2} \frac{\sigma}{\sqrt{n}}$ und $\bar{\theta}(X_1, \dots, X_n) = \bar{X}_n + z_{1-\alpha_1} \frac{\sigma}{\sqrt{n}}$ ein exaktes Konfidenzintervall für μ zum Niveau $1 - \alpha$.

Es hat die Länge $l_\mu(X_1, \dots, X_n) = \frac{\sigma}{\sqrt{n}} (z_{1-\alpha_2} + z_{1-\alpha_1})$. Es gilt $l_\mu(X_1, \dots, X_n) \rightarrow 0$ für $n \rightarrow \infty$ was bedeutet, daß bei wachsendem Informationsumfang ($n \rightarrow \infty$) die Präzision der Schätzung immer besser wird.

Im Symmetriefall ($\alpha_1 = \alpha_2 = \alpha/2$) gilt $\underline{\theta}(X_1, \dots, X_n) = \bar{X}_n - z_{1-\alpha/2} \frac{\sigma}{\sqrt{n}}$, $\bar{\theta}(X_1, \dots, X_n) = \bar{X}_n + z_{1-\alpha/2} \frac{\sigma}{\sqrt{n}}$ und $l_\mu(X_1, \dots, X_n) = \frac{2\sigma}{\sqrt{n}} z_{1-\alpha/2}$.

Daraus folgt, daß man bei vorgegebener Länge $\varepsilon > 0$ die Anzahl der Beobachtungen n bestimmen kann, die dann notwendig sind, um die vorgegebene Präzision zu erreichen:

$$\frac{2\sigma}{\sqrt{n}} z_{1-\alpha/2} \leq \varepsilon \iff n \geq \left(\frac{2\sigma z_{1-\alpha/2}}{\varepsilon} \right)^2 \quad (4.2.1)$$

Für $\alpha_1 = 0$ bzw. $\alpha_2 = 0$ kann man einseitige Intervalle $\left(-\infty, \bar{X}_n + z_{1-\alpha} \frac{\sigma}{\sqrt{n}}\right]$ und $\left[\bar{X}_n - z_{1-\alpha} \frac{\sigma}{\sqrt{n}}, +\infty\right)$ genauso angeben.

- **bei unbekannter Varianz σ^2 :** siehe Bemerkung 3.3.4.

Dort wurde das Konfidenzintervall $\left[\bar{X}_n - \frac{t_{n-1,1-\alpha/2}}{\sqrt{n}} S_n, \bar{X}_n + \frac{t_{n-1,1-\alpha/2}}{\sqrt{n}} S_n\right]$ für μ zum Konfidenzniveau $1 - \alpha$ konstruiert, wobei $t_{n-1,1-\alpha/2}$ das $(1 - \frac{\alpha}{2})$ -Quantil der t_{n-1} -Verteilung ist.

Wie man sieht, ist die Länge des Konfidenzintervalls zufällig: $l_\mu(X_1, \dots, X_n) = \frac{2S_n}{\sqrt{n}} t_{n-1,1-\alpha/2}$, somit macht es Sinn, mit erwarteter Länge

$$\mathbb{E} l_\mu(X_1, \dots, X_n) = \frac{2}{\sqrt{n}} \mathbb{E} S_n t_{n-1,1-\alpha/2}$$

zu arbeiten, um zum Beispiel die Frage nach der notwendigen Anzahl n von Beobachtungen bei vorgegebener Genauigkeit $\varepsilon > 0$ (vergleiche Gleichung (4.2.1)) zu beantworten.

Konfidenzintervalle für die Varianz σ^2

- **bei bekanntem Erwartungswert μ :**

Betrachten wir den Schätzer $\tilde{S}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2$ für σ^2 . Aus Satz 3.3.5, 2. folgt $\frac{n\tilde{S}_n^2}{\sigma^2} \sim \chi_n^2$.

Wir setzen $T(X_1, \dots, X_n, \sigma^2) = \frac{n\tilde{S}_n^2}{\sigma^2}$ und bekommen

$$\mathbb{P} \left(\chi_{n,\alpha_2}^2 \leq \frac{n\tilde{S}_n^2}{\sigma^2} \leq \chi_{n,1-\alpha_1}^2 \right) = \mathbb{P} \left(\frac{n\tilde{S}_n^2}{\chi_{n,1-\alpha_1}^2} \leq \sigma^2 \leq \frac{n\tilde{S}_n^2}{\chi_{n,\alpha_2}^2} \right) = 1 - \alpha.$$

Somit ist $\left[\frac{n\tilde{S}_n^2}{\chi_{n,1-\alpha_1}^2}, \frac{n\tilde{S}_n^2}{\chi_{n,\alpha_2}^2} \right]$ ein Konfidenzintervall für σ^2 zum Niveau $1 - \alpha, \alpha = \alpha_1 + \alpha_2$

mit der mittleren Länge $\mathbb{E} l_{\sigma^2} = n\sigma^2 \left(\frac{1}{\chi_{n,\alpha_2}^2} - \frac{1}{\chi_{n,1-\alpha_1}^2} \right)$.

- **bei unbekanntem Erwartungswert μ :**

Ähnlich wie oben beschrieben folgt das Konfidenzintervall $\left[\frac{(n-1)S_n^2}{\chi_{n-1,1-\alpha_1}^2}, \frac{(n-1)S_n^2}{\chi_{n-1,\alpha_2}^2} \right]$ zum Niveau $1 - \alpha, \alpha = \alpha_1 + \alpha_2$ aus Satz 3.3.5, 1., weil $\frac{(n-1)S_n^2}{\sigma^2} \sim \chi_{n-1}^2$ für die Stichprobenvarianz $S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2$. Die erwartete Länge ist $\mathbb{E} l_{\sigma^2} = (n-1)\sigma^2 \left(\frac{1}{\chi_{n-1,\alpha_2}^2} - \frac{1}{\chi_{n-1,1-\alpha_1}^2} \right)$.

4.2.2 Konfidenzintervalle aus stochastischen Ungleichungen

Eine alternative Methode zur Gewinnung von Konfidenzintervallen besteht in der Anwendung stochastischer Ungleichungen. So kann man zum Beispiel bei einer Stichprobe $(X_1, \dots, X_n)$ von unabhängigen und identisch verteilten Zufallsvariablen mit $\mathbb{E} X_i = \mu$, $\text{Var } X_i = \sigma^2 \in (0, \infty)$ die Ungleichung von Tschebyschew benutzen, um ein einfaches, aber grobes Konfidenzintervall für μ zu konstruieren:

$$\begin{aligned} \mathbb{P}\left(|\bar{X}_n - \mu| > \varepsilon\right) &\leq \frac{\text{Var } \bar{X}_n}{\varepsilon^2} = \frac{\sigma^2}{n\varepsilon^2} = \alpha \\ &\Rightarrow \text{für } \varepsilon = \frac{\sigma}{\sqrt{n\alpha}} \text{ gilt: } 1 - \alpha \leq \mathbb{P}\left(|\bar{X}_n - \mu| \leq \varepsilon\right) \\ &= \mathbb{P}\left(-\frac{\sigma}{\sqrt{n\alpha}} \leq -\bar{X}_n + \mu \leq \frac{\sigma}{\sqrt{n\alpha}}\right) \\ &= \mathbb{P}\left(\bar{X}_n - \frac{\sigma}{\sqrt{n\alpha}} \leq \mu \leq \bar{X}_n + \frac{\sigma}{\sqrt{n\alpha}}\right). \end{aligned}$$

Das Konfidenzintervall $\left[\bar{X}_n - \frac{\sigma}{\sqrt{n\alpha}}, \bar{X}_n + \frac{\sigma}{\sqrt{n\alpha}}\right]$ für μ bei bekannter Varianz σ^2 ist verteilungsunabhängig, da keinerlei Annahmen über die Verteilung von X_i gemacht wurden.

Präzisere Konfidenzintervalle können bei der Verwendung folgender *Ungleichung von Hoeffding* konstruiert werden:

Satz 4.2.1 (Ungleichung von Hoeffding):

Es seien $Y_1, \dots, Y_n$ unabhängige Zufallsvariablen mit $\mathbb{E} Y_i = 0$, $a_i \leq Y_i \leq b_i$ fast sicher, $i = 1, \dots, n$. Für alle $\varepsilon > 0$ gilt

$$\mathbb{P}\left(\sum_{i=1}^n Y_i \geq \varepsilon\right) \leq \exp\left(-\frac{2\varepsilon^2}{\sum_{i=1}^n (b_i - a_i)^2}\right)$$

(ohne Beweis).

Nehmen wir z.B. an, daß $X_1, \dots, X_n$ unabhängige, identisch verteilte Zufallsvariablen sind, $X_i \sim \text{Bernoulli}(p)$, $p \in (0, 1)$. Wir wollen ein Konfidenzintervall für p bestimmen.

Folgerung 4.2.1

Es seien $X_1, \dots, X_n$ unabhängige Bernoulli(p)-verteilte Zufallsvariablen. Dann gilt $\mathbb{P}\left(|\bar{X}_n - p| > \varepsilon\right) \leq 2e^{-2n\varepsilon^2}$, $\varepsilon > 0$.

Beweis Es gilt

$$\bar{X}_n - p = \frac{1}{n} \sum_{i=1}^n \underbrace{(X_i - p)}_{Y_i}, \quad Y_i \in [-p, 1 - p],$$

das heißt $a_i = -p$, $b_i = 1 - p$, $b_i - a_i = 1$, $i = 1, \dots, n$, $\mathbb{E} Y_i = p - p = 0$. Dann gilt:

$$\begin{aligned} \mathbb{P}_p(|\bar{X}_n - p| > \varepsilon) &\leq \mathbb{P}_p\left(\left|\sum_{i=1}^n Y_i\right| \geq \varepsilon n\right) \\ &= \mathbb{P}_p\left(\sum_{i=1}^n Y_i \geq \varepsilon n\right) + \mathbb{P}_p\left(\sum_{i=1}^n (-Y_i) \geq \varepsilon n\right) \\ &\stackrel{(\text{Satz 4.2.1})}{\leq} 2e^{-\frac{2\varepsilon^2 n^2}{n}} = 2e^{-2\varepsilon^2 n}, \end{aligned}$$

wobei man den Satz 4.2.1 sowohl für die Folge $\{Y_i\}$ als auch $\{-Y_i\}$ anwendet. Damit ist die Behauptung bewiesen. $\square$

Bemerkung 4.2.1

Die Form der Ungleichung von Hoeffding ähnelt sehr der von Dvoretzky-Kiefer-Wolfowitz, Satz 3.3.10.

Nun fixieren wir $\alpha > 0$ und wählen $\varepsilon_n = \sqrt{\frac{1}{2n} \log \frac{2}{\alpha}}$. Durch Anwendung von Folgerung 4.2.1 mit diesem ε_n erhalten wir $\mathbb{P}_p(|\bar{X}_n - p| > \varepsilon_n) \leq \alpha$, somit $\mathbb{P}_p(|\bar{X}_n - p| \leq \varepsilon_n) \geq 1 - \alpha$ und darum ist $[\bar{X}_n - \sqrt{\frac{1}{2n} \log \frac{2}{\alpha}}, \bar{X}_n + \sqrt{\frac{1}{2n} \log \frac{2}{\alpha}}]$ ein Konfidenzintervall für p zum Niveau $1 - \alpha$.

4.2.3 Asymptotische Konfidenzintervalle

Die Philosophie der Konstruktion von asymptotischen Konfidenzintervallen ist relativ einfach: Wir erläutern sie am Beispiel eines asymptotisch normalverteilten Schätzers $\hat{\theta}$ für einen Parameter θ .

Sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen, $X_i \sim F_\theta$, $\theta \in \Theta \subseteq \mathbb{R}$. Sei $\hat{\theta}_n = \hat{\theta}(X_1, \dots, X_n)$ ein Schätzer für θ , der asymptotisch normalverteilt ist. Dann gilt für erwartungstreue $\hat{\theta}_n$

$$\frac{\hat{\theta}_n - \theta}{\hat{\sigma}_n} \xrightarrow{d} Y \sim N(0, 1),$$

wobei $\hat{\sigma}_n$ ein konsistenter Schätzer der asymptotischen Varianz von $\hat{\theta}_n$ ist.

$$\begin{aligned} &\lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(z_{\alpha/2} \leq \frac{\hat{\theta}_n - \theta}{\hat{\sigma}_n} \leq z_{1-\alpha/2} \right) \\ &= \lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(\theta \in [\hat{\theta}_n - z_{1-\alpha/2} \hat{\sigma}_n, \hat{\theta}_n + z_{1-\alpha/2} \hat{\sigma}_n] \right) = 1 - \alpha. \end{aligned}$$

Somit ist $[\hat{\theta}_n - z_{1-\alpha/2} \hat{\sigma}_n, \hat{\theta}_n + z_{1-\alpha/2} \hat{\sigma}_n]$ ein asymptotisches Konfidenzintervall für θ zum Niveau $1 - \alpha$.

Diese Vorgehensweise werden wir jetzt anhand von zwei Beispielen klar machen:

- **Bernoulli-Verteilung:**

Seien $X_i \sim \text{Bernoulli}(p)$ -verteilt, $i = 1, \dots, n$. Dann gilt $\theta = p$, $\hat{\theta}_n = \hat{p}_n = \bar{X}_n$. $\mathbb{E}_p \hat{p}_n = p$, $\text{Var}_p \hat{p}_n = \frac{p(1-p)}{n}$. Wir wählen $\hat{\sigma}^2 = \frac{1}{n} \hat{p}_n(1 - \hat{p}_n) = \frac{\bar{X}_n}{n}(1 - \bar{X}_n)$ als Plug-In-Schätzer

für σ^2 . Dann gilt nach dem zentralen Grenzwertsatz (Satz 7.2.1, WR) und dem Satz von Slutsky (Satz 6.4.2, 3. WR):

$$\sqrt{n} \frac{\bar{X}_n - p}{\sqrt{\bar{X}_n(1 - \bar{X}_n)}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1),$$

das heißt $p \in \left[\bar{X}_n - z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n(1-\bar{X}_n)}{n}}, \bar{X}_n + z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n(1-\bar{X}_n)}{n}} \right]$ stellt ein asymptotisches Konfidenzintervall für p zum Niveau $1 - \alpha$ dar. Da aber $p \in [0, 1]$ sein soll, betrachtet man

$$\underline{p}(X_1, \dots, X_n) = \max \left\{ 0, \bar{X}_n - z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n(1 - \bar{X}_n)}{n}} \right\}$$

und

$$\bar{p}(X_1, \dots, X_n) = \min \left\{ 1, \bar{X}_n + z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n(1 - \bar{X}_n)}{n}} \right\}.$$

Bemerkung 4.2.2

Ein anderes asymptotisches Konfidenzintervall für den Parameter p der Bernoulli-Verteilung bekommt man, wenn man die Aussage des zentralen Grenzwertsatzes

$\lim_{n \rightarrow \infty} \mathbb{P}_p \left(-z_{1-\alpha/2} \leq \sqrt{n} \frac{\bar{X}_n - p}{\sqrt{p(1-p)}} \leq z_{1-\alpha/2} \right) = 1 - \alpha$ nimmt und die quadratische Ungleichung dann bezüglich p auflöst.

Übungsaufgabe 4.2.1

Lösen Sie die Ungleichung auf!

• Poissonverteilung:

Es seien $X_i \sim \text{Poisson}(\lambda)$, $i = 1, \dots, n$, dann gilt $\theta = \lambda$, $\hat{\theta}_n = \hat{\lambda} = \bar{X}_n$. Da $\mathbb{E}_\lambda X_i = \text{Var}_\lambda X_i = \lambda$, kann man den zentralen Grenzwertsatz (Satz 7.2.1, WR) anwenden

$$\sqrt{n} \frac{\bar{X}_n - \lambda}{\sqrt{\lambda}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1),$$

Da $\bar{X}_n$ stark konsistent für λ ist, gilt nach dem Satz von Slutsky (Satz 6.4.2, 4, WR)

$$\sqrt{n} \frac{\bar{X}_n - \lambda}{\sqrt{\bar{X}_n}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1).$$

Daraus folgt ein asymptotisches Konfidenzintervall

$$\left[\bar{X}_n - z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n}{n}}, \bar{X}_n + z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n}{n}} \right]$$

für den Parameter λ zum Konfidenzniveau $1 - \alpha$.

Bemerkung 4.2.3 1. Ähnlich wie in Bemerkung 4.2.2 angegeben, kann man durch Auflösen der quadratischen Ungleichung in

$$\lim_{n \rightarrow \infty} \mathbb{P}_\lambda \left(\sqrt{n} \frac{\bar{X}_n - \lambda}{\sqrt{\lambda}} \in [-z_{1-\alpha/2}, z_{1-\alpha/2}] \right) = 1 - \alpha$$

bezüglich λ ein alternatives asymptotisches Konfidenzintervall für λ angeben.

Übungsaufgabe 4.2.2

Bitte führen Sie diese Berechnungen durch.

2. Da $\lambda > 0$ ist, kann man die untere Schranke diesbezüglich korrigieren:

$$\underline{\lambda}(X_1, \dots, X_n) = \max \left\{ 0, \bar{X}_n - z_{1-\alpha/2} \sqrt{\frac{\bar{X}_n}{n}} \right\}$$

4.3 Zwei-Stichproben-Probleme

In diesem Abschnitt werden Charakteristiken bzw. Parameter von zwei unterschiedlichen Stichproben miteinander verglichen, indem man Konfidenzintervalle für einfache Funktionen dieser Parameter konstruiert.

Betrachten wir zwei Zufallsstichproben $Y_1 = (X_{11}, \dots, X_{1n_1})$, $Y_2 = (X_{21}, \dots, X_{2n_2})$ von Zufallsvariablen $X_{i1}, \dots, X_{in_i}$, $i = 1, 2$, die innerhalb der Stichprobe Y_i jeweils unabhängig und identisch verteilt sind, $X_{ij} \stackrel{d}{=} X_i$, $j = 1, \dots, n_i$, $i = 1, 2$ und die Prototyp-Zufallsvariable $X_i \sim F_{\theta_i}$, $\theta_i \in \Theta \subset \mathbb{R}^m$. Es wird im Allgemeinen nicht gefordert, daß Y_1 und Y_2 unabhängig sind. Falls sie voneinander abhängen, spricht man von *verbundenen Stichproben* Y_1 und Y_2 . Betrachten wir eine Funktion $g : \mathbb{R}^{2m} \rightarrow \mathbb{R}$ von den Parametervektoren θ_1 und θ_2 . In diesem Skript werden dabei meistens die Fälle $m = 1, 2$, $g(\theta_1, \theta_2) = \theta_{1j} - \theta_{2j}$, $g(\theta_1, \theta_2) = \frac{\theta_{1j}}{\theta_{2j}}$ untersucht, wobei $\theta_i = (\theta_{i1}, \dots, \theta_{im})$, $i = 1, 2$.

Unsere Zielstellung wird sein, ein (möglicherweise asymptotisches) Konfidenzintervall für $g(\theta_1, \theta_2)$ mit Hilfe der Stichprobe (Y_1, Y_2) zu gewinnen.

Dabei wird die selbe Philosophie wie in Abschnitt 4.1 beschrieben verfolgt. Es wird eine Statistik $T(Y_1, Y_2, g(\theta_1, \theta_2))$ gesucht, die eine (möglicherweise asymptotische) Prüfverteilung F besitzt und von $g(\theta_1, \theta_2)$ explizit abhängt.

Durch das Auflösen der Ungleichung $F_{\alpha_1}^{-1} \leq T(Y_1, Y_2, g(\theta_1, \theta_2)) \leq F_{1-\alpha_2}^{-1}$ bzgl. $g(\theta_1, \theta_2)$ bekommt man dann ein (möglicherweise asymptotisches) Konfidenzintervall zum Niveau $1 - \alpha$, $\alpha = \alpha_1 + \alpha_2$.

4.3.1 Normalverteilte Stichproben

Hier wird angenommen, daß $X_i \sim N(\mu_i, \sigma_i^2)$, $i = 1, 2$.

Konfidenzintervall für die Differenz $\mu_1 - \mu_2$ bei bekannten Varianzen σ_1^2 und σ_2^2 und unabhängigen Stichproben

Seien Y_1 und Y_2 voneinander unabhängig und σ_1^2, σ_2^2 bekannt. Wir betrachten die Parameterfunktion $g(\mu_1, \mu_2) = \mu_1 - \mu_2$. Es seien $\bar{X}_{in_i} = \frac{1}{n_i} \sum_{j=1}^{n_i} X_{ij}$, $i = 1, 2$ die Stichprobenmittel der Stichproben Y_1 und Y_2 . Es gilt $\bar{X}_{in_i} \sim N(\mu_i, \frac{\sigma_i^2}{n_i})$, $i = 1, 2$. Nach Satz 3.3.3, 4)

sind $\bar{X}_{1n_1}$ und $\bar{X}_{2n_2}$ unabhängig. Dann ist wegen der Faltungsstabilität der Normalverteilung $\bar{X}_{1n_1} - \bar{X}_{2n_2} \sim N\left(\mu_1 - \mu_2, \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}\right)$. Nach dem Normieren erhält man die Statistik $T(Y_1, Y_2, \mu_1 - \mu_2) = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2} - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \sim N(0, 1)$. Daraus bekommt man das Konfidenzintervall

$$\left[\bar{X}_{1n_1} - \bar{X}_{2n_2} - z_{1-\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}, \bar{X}_{1n_1} - \bar{X}_{2n_2} + z_{1-\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} \right]$$

für $\mu_1 - \mu_2$ zum Niveau $1 - \alpha$.

Konfidenzintervall für den Quotienten $\frac{\sigma_1^2}{\sigma_2^2}$ bei unbekannten Erwartungswerten μ_1 und μ_2 und unabhängigen Stichproben

Seien Y_1 und Y_2 voneinander unabhängig. Sei $g(\sigma_1, \sigma_2) = \frac{\sigma_1^2}{\sigma_2^2}$. Wir konstruieren die Statistik $T(Y_1, Y_2, \frac{\sigma_1^2}{\sigma_2^2})$ folgendermaßen: Seien $S_{in_i}^2 = \frac{1}{n_i-1} \sum_{j=1}^{n_i} (X_{ij} - \bar{X}_{in_i})^2$, $i = 1, 2$ die Stichprobenvarianzen der Stichproben Y_1 und Y_2 . Dann gilt $\frac{(n_i-1)S_{in_i}^2}{\sigma_i^2} \sim \chi_{n_i-1}^2$, $i = 1, 2$ nach Satz 3.3.5.

Da die $S_{in_i}^2$ voneinander unabhängig sind, gilt

$$T\left(Y_1, Y_2, \frac{\sigma_1^2}{\sigma_2^2}\right) = \frac{\frac{(n_2-1)S_{2n_2}^2}{(n_2-1)\sigma_2^2}}{\frac{(n_1-1)S_{1n_1}^2}{(n_1-1)\sigma_1^2}} = \frac{S_{2n_2}^2}{S_{1n_1}^2} \cdot \frac{\sigma_1^2}{\sigma_2^2} \sim F_{n_2-1, n_1-1}$$

nach der Definition der F -Verteilung. Daraus ergibt sich das Konfidenzintervall

$$\left[\frac{S_{1n_1}^2}{S_{2n_2}^2} F_{n_2-1, n_1-1, \alpha_1}, \frac{S_{1n_1}^2}{S_{2n_2}^2} F_{n_2-1, n_1-1, 1-\alpha_2} \right]$$

für $\frac{\sigma_1^2}{\sigma_2^2}$ zum Niveau $1 - \alpha$.

Konfidenzintervall für die Differenz $\mu_1 - \mu_2$ der Erwartungswerte bei verbundenen Stichproben

Dieses Mal seien Y_1 und Y_2 verbunden, $X_1 - X_2 \sim N(\mu_1 - \mu_2, \sigma^2)$ für ein unbekanntes $\sigma^2 > 0$, $n_1 = n_2 = n$. Da X_{ij} , $j = 1, \dots, n$ unabhängig und identisch verteilt sind, gilt $Z_j = X_{1j} - X_{2j} \sim N(\mu_1 - \mu_2, \sigma^2)$, $j = 1, \dots, n$.

Unser Ziel ist es, ein Konfidenzintervall für $\mu_1 - \mu_2$ zu bekommen. Wenn wir die Stichprobe $(Z_1, \dots, Z_n)$ betrachten, und Ergebnisse des Abschnittes 4.2.1, 2. anwenden, so erhalten wir sofort folgendes Konfidenzintervall:

$$\left[\bar{Z}_n - t_{n-1, 1-\frac{\alpha}{2}} \frac{S_n}{\sqrt{n}}, \bar{Z}_n + t_{n-1, 1-\frac{\alpha}{2}} \frac{S_n}{\sqrt{n}} \right]$$

für $\mu_1 - \mu_2$ zum Niveau $1 - \frac{\alpha}{2}$, wobei $\bar{Z}_n = \frac{1}{n} \sum_{j=1}^n Z_j = \frac{1}{n} \sum_{j=1}^n (X_{1j} - X_{2j}) = \bar{X}_{1n} - \bar{X}_{2n}$,

$$S_n^2 = \frac{1}{n-1} \sum_{j=1}^n (Z_j - \bar{Z}_n)^2 = \frac{1}{n-1} \sum_{j=1}^n (X_{1j} - X_{2j} - \bar{X}_{1n} + \bar{X}_{2n})^2.$$

4.3.2 Poissonverteilte Stichproben

Wir nehmen jetzt an, daß die Stichproben Y_1 und Y_2 unabhängig sind, und $X_i \sim \text{Poisson}(\lambda_i)$, $i = 1, 2$. Konstruieren wir asymptotische Konfidenzintervalle für $g(\lambda_1, \lambda_2) = \lambda_1 - \lambda_2$ und $g(\lambda_1, \lambda_2) = \frac{n_2 \lambda_2}{n_1 \lambda_1 + n_2 \lambda_2} = \frac{\lambda_2}{\rho \lambda_1 + \lambda_2}$, $\rho = \frac{n_1}{n_2} = \text{const}$, wobei $n_1, n_2 \rightarrow \infty$.

Asymptotisches Konfidenzintervall für $\lambda_1 - \lambda_2$

Um zu einer Statistik $T(Y_1, Y_2, \lambda_1 - \lambda_2)$ zu kommen, die asymptotisch (für $n_1, n_2 \rightarrow \infty$) $N(0, 1)$ -verteilt ist, verwenden wir den zentralen Grenzwertsatz von Ljapunow (vergleiche Satz 7.2.6, WR).

Lemma 4.3.1

Für $n_1, n_2 \rightarrow \infty$ mit $0 < c_1 \leq n_1/n_2 \leq c_2 < \infty$ gilt

$$\frac{\bar{X}_{1n_1} - \bar{X}_{2n_2} - \lambda_1 + \lambda_2}{\sqrt{\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2}}} \xrightarrow[n_1 \rightarrow \infty, n_2 \rightarrow \infty]{d} Y \sim N(0, 1)$$

Beweis Führen wir die Zufallsvariable

$$Z_{nk} = \begin{cases} \frac{X_{1k} - \lambda_1}{n_1 \sqrt{\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2}}}, & k = 1, \dots, n_1 \\ -\frac{X_{2k-n_1} - \lambda_2}{n_2 \sqrt{\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2}}}, & k = n_1 + 1, \dots, n_1 + n_2 \end{cases}$$

ein, wobei $n = n_1 + n_2$. Es gilt: $\mathbb{E} Z_{nk} = 0$ für alle $k = 1, \dots, n$, und

$$0 < \sigma_{nk}^2 = \text{Var } Z_{nk} = \begin{cases} \frac{\text{Var } X_{1k}}{n_1^2 \left(\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2} \right)} = \frac{\lambda_1}{n_1^2 \left(\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2} \right)}, & k = 1, \dots, n_1, \\ \frac{\lambda_2}{n_2^2 \left(\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2} \right)}, & k = n_1 + 1, \dots, n, \end{cases}$$

somit

$$\sum_{k=1}^n \sigma_{nk}^2 = \left(\frac{\lambda_1}{n_1^2} n_1 + \frac{\lambda_2}{n_2^2} n_2 \right) \frac{1}{\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2}} = 1.$$

Außerdem gilt für $\delta > 0$ und $n_1, n_2 \rightarrow \infty$:

$$\lim_{n \rightarrow \infty} \sum_{k=1}^n \mathbb{E} (|Z_{nk}|)^{2+\delta} = \lim_{n_1, n_2 \rightarrow \infty} \left(\frac{\mathbb{E} (|X_{11} - \lambda_1|^{2+\delta})}{n_1^{1+\delta} \left(\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2} \right)^{(2+\delta)/2}} + \frac{\mathbb{E} (|X_{21} - \lambda_2|^{2+\delta})}{n_2^{1+\delta} \left(\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2} \right)^{(2+\delta)/2}} \right) = 0$$

Somit ist die Ljapunow-Bedingung erfüllt und nach Satz 7.2.6 (WR) gilt

$$\sum_{k=1}^n Z_{nk} \xrightarrow[n_1 \rightarrow \infty, n_2 \rightarrow \infty]{d} Y \sim N(0, 1).$$

Es gilt aber auch $\sum_{n=1}^n Z_{nk} = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2} - \lambda_1 + \lambda_2}{\sqrt{\frac{\lambda_1}{n_1} + \frac{\lambda_2}{n_2}}}$, somit ist das Lemma bewiesen. $\square$

Da $\bar{X}_{in_i} \xrightarrow{f.s.} \lambda_i$, $i = 1, 2$ nach dem starken Gesetz der großen Zahlen, gilt mit Hilfe des Satzes von Slutsky

$$T(Y_1, Y_2, \lambda_1 - \lambda_2) = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2} - \lambda_1 + \lambda_2}{\sqrt{\bar{X}_{1n_1}/n_1 + \bar{X}_{2n_2}/n_2}} \xrightarrow[n_1, n_2 \rightarrow \infty]{d} Y \sim N(0, 1)$$

Daraus läßt sich sofort das asymptotische Konfidenzintervall für $\lambda_1 - \lambda_2$ zum Niveau $1 - \alpha$ ableiten:

$$\left[\bar{X}_{1n_1} - \bar{X}_{2n_2} - z_{1-\alpha/2} \sqrt{\frac{\bar{X}_{1n_1}}{n_1} + \frac{\bar{X}_{2n_2}}{n_2}}, \bar{X}_{1n_1} - \bar{X}_{2n_2} + z_{1-\alpha/2} \sqrt{\frac{\bar{X}_{1n_1}}{n_1} + \frac{\bar{X}_{2n_2}}{n_2}} \right]$$

Asymptotisches Konfidenzintervall für $\frac{n_2 \lambda_2}{n_1 \lambda_1 + n_2 \lambda_2}$

Es sei $n_1/n_2 = \rho = \text{const}$ und $g(\lambda_1, \lambda_2) = \frac{n_2 \lambda_2}{n_1 \lambda_1 + n_2 \lambda_2} = \frac{\lambda_2}{\rho \lambda_1 + \lambda_2} \stackrel{\text{Def.}}{=} p$. Es wird ein asymptotisches Konfidenzintervall für p gesucht. Wir führen die Statistik

$$T(Y_1, Y_2, p) = \frac{S_{2n_2} - p(S_{1n_1} + S_{2n_2})}{\sqrt{\hat{p}(1 - \hat{p})(S_{1n_1} + S_{2n_2})}}$$

ein, wobei $S_{in_i} = \sum_{j=1}^{n_i} X_{ij}$, $i = 1, 2$ und

$$\hat{p} = \frac{S_{2n_2}}{S_{1n_1} + S_{2n_2}} = \frac{n_2 \bar{X}_{2n_2}}{n_1 \bar{X}_{1n_1} + n_2 \bar{X}_{2n_2}} \xrightarrow[n_1, n_2 \rightarrow \infty]{f.s.} p$$

ein konsistenter Schätzer für p (wegen des starken Gesetzes der großen Zahlen) ist. Falls wir zeigen können, daß $T(Y_1, Y_2, p) \xrightarrow[n_1, n_2 \rightarrow \infty]{d} Y \sim N(0, 1)$, so wird daraus folgendes Konfidenzintervall ableitbar: Aus

$$\lim_{\substack{n_1 \rightarrow \infty \\ n_2 \rightarrow \infty}} \mathbb{P} \left(-z_{1-\alpha/2} \leq \frac{\frac{S_{2n_2}}{S_{1n_1} + S_{2n_2}} - p}{\sqrt{S_{1n_1} \cdot S_{2n_2}}} \cdot (S_{1n_1} + S_{2n_2})^{3/2} \leq z_{1-\alpha/2} \right) = 1 - \alpha$$

folgt, daß

$$[\underline{\theta}(Y_1, Y_2), \bar{\theta}(Y_1, Y_2)]$$

mit

$$\begin{aligned} \underline{\theta}(\lambda_1, \lambda_2) &= \frac{S_{2n_2}}{S_{1n_1} + S_{2n_2}} - z_{1-\alpha/2} \cdot \sqrt{\frac{S_{1n_1} \cdot S_{2n_2}}{(S_{1n_1} + S_{2n_2})^3}} \\ \bar{\theta}(\lambda_1, \lambda_2) &= \frac{S_{2n_2}}{S_{1n_1} + S_{2n_2}} + z_{1-\alpha/2} \cdot \sqrt{\frac{S_{1n_1} \cdot S_{2n_2}}{(S_{1n_1} + S_{2n_2})^3}} \end{aligned}$$

ein asymptotisches Konfidenzintervall für p zum Niveau $1 - \alpha$ ist.

Da $0 < p < 1$ sein soll, können die Schranken des Intervalls diesbezüglich korrigiert werden:

$$\begin{aligned}\underline{\theta}^*(Y_1, Y_2) &= \max\{0, \underline{\theta}(Y_1, Y_2)\}, \\ \bar{\theta}^*(Y_1, Y_2) &= \min\{1, \bar{\theta}(Y_1, Y_2)\}.\end{aligned}$$

Nun soll die asymptotische Normalverteiltheit von $T(Y_1, Y_2, p)$ gezeigt werden. Sie folgt aus dem Satz von Slutsky und folgendem Lemma:

Lemma 4.3.2

Es gilt:

$$\frac{S_{2n_2} - p(S_{1n_1} + S_{2n_2})}{\sqrt{p(1-p)(S_{1n_1} + S_{2n_2})}} \xrightarrow[n_1 \rightarrow \infty]{d} Y \sim N(0, 1)$$

Beweis Um die Aussage des Lemmas zu zeigen, verwenden wir einen zentralen Grenzwertsatz für Summen von Zufallsvariablen in zufälliger Anzahl (vgl. Satz 7.2.2 (WR)). Führen wir die Folge $N_n = S_{1n_1} + S_{2n_2}$ von nichtnegativen Zufallsvariablen ein. Die Summe ist monoton wachsend. Gleichzeitig setzen wir $a_{n_2} = n_1\lambda_1 + n_2\lambda_2$. Offensichtlich gilt

$$\begin{aligned}\frac{N_n}{a_{n_2}} &= \frac{S_{1n_1}}{n_1\lambda_1 + n_2\lambda_2} + \frac{S_{2n_2}}{n_1\lambda_1 + n_2\lambda_2} \\ &= \frac{\bar{X}_{1n_1}}{\lambda_1 + \rho^{-1}\lambda_2} + \frac{\bar{X}_{2n_2}}{\rho\lambda_1 + \lambda_2} \\ &\xrightarrow[n_1, n_2 \rightarrow \infty]{f.s.} \frac{\lambda_1}{\lambda_1 + \rho^{-1}\lambda_2} + \frac{\lambda_2}{\rho\lambda_1 + \lambda_2} \\ &= \frac{\rho\lambda_1}{\rho\lambda_1 + \lambda_2} + \frac{\lambda_2}{\rho\lambda_1 + \lambda_2} = 1\end{aligned}$$

Außerdem gilt:

$$\begin{aligned}\mathbb{P}(S_{2n_2} = k \mid N_n = m) &= \frac{\mathbb{P}(S_{2n_2} = k, S_{1n_1} + S_{2n_2} = m)}{\mathbb{P}(S_{1n_1} + S_{2n_2} = m)} \\ &= \frac{\mathbb{P}(S_{2n_2} = k, S_{1n_1} = m - k)}{\mathbb{P}(S_{1n_1} + S_{2n_2} = m)} \\ &= \frac{e^{-n_2\lambda_2} \frac{(\lambda_2 n_2)^k}{k!} \cdot e^{-n_1\lambda_1} \frac{(n_1\lambda_1)^{m-k}}{(m-k)!}}{e^{-n_1\lambda_1 - n_2\lambda_2} \frac{(n_1\lambda_1 + n_2\lambda_2)^m}{m!}} \\ &= \frac{m!}{(m-k)!k!} \left(\frac{n_2\lambda_2}{n_1\lambda_1 + n_2\lambda_2} \right)^m \left(\frac{n_1\lambda_1}{n_1\lambda_1 + n_2\lambda_2} \right)^{m-k} \\ &= \binom{m}{k} p^k (1-p)^{m-k}\end{aligned}$$

was bedeutet, daß $S_{2n_2} \mid \{N_n = m\} \sim \text{Bin}(m, p)$. Dann gilt $\frac{S_{2n_2} - mp}{\sqrt{mp(1-p)}} \mid \{N_n = m\} \xrightarrow{d} \frac{S_m - mp}{\sqrt{mp(1-p)}}$,

wobei $S_m = \sum_{i=1}^m Z_i$ eine Summe von unabhängigen, identisch verteilten Zufallsvariablen $Z_i \sim \text{Bernoulli}(p)$ ist. Nach Satz 7.2.2 (WR) gilt dann

$$\frac{S_{N_n} - N_n p}{\sqrt{N_n p(1-p)}} \xrightarrow{d} Y \sim N(0, 1) \iff \frac{S_{2n_2} - N_n p}{\sqrt{N_n p(1-p)}} \xrightarrow{d} Y \sim N(0, 1).$$

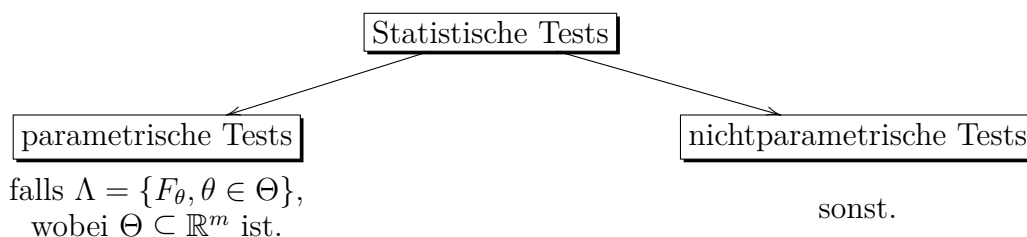
□

5 Tests statistischer Hypothesen

In Kapitel 3 haben wir schon Beispiele von statistischen Tests kennengelernt, wie etwa den Kolmogorow-Smirnow-Test (vergleiche Bemerkung 3.3.8, 3). Jetzt sollen statistische Signifikanztests formal eingeführt und ihre Eigenschaften untersucht werden.

5.1 Allgemeine Philosophie des Testens

Es sei eine Zufallsstichprobe $(X_1, \dots, X_n)$ von unabhängigen, identisch verteilten Zufallsvariablen X_i gegeben, mit Verteilungsfunktion $F \in \Lambda$, wobei Λ eine Klasse von Verteilungsfunktionen ist. Es sei $(x_1, \dots, x_n)$ eine konkrete Stichprobe, die als Realisierung von $(X_1, \dots, X_n)$ interpretiert wird. In der Theorie des statistischen Testens werden Hypothesen über die Beschaffenheit der (unbekannten) Verteilungsfunktion F gestellt und geprüft. Dabei unterscheidet man



Bei parametrischen Tests prüft man, ob der Parameter θ bestimmte Werte annimmt (zum Beispiel $\theta = 0$). Bekannte Beispiele von nichtparametrischen Tests sind Anpassungstests, bei denen man prüft, ob die Verteilungsfunktion F gleich einer vorgegebenen Funktion F_0 ist.

Formalisieren wir zunächst den Begriff *Hypothese*. Die Menge Λ von zulässigen Verteilungsfunktionen F wird in zwei disjunkte Teilmengen Λ_0 und Λ_1 zerlegt, $\Lambda_0 \cup \Lambda_1 = \Lambda$. Die Aussage

„Man testet die *Haupthypothese* $H_0 : F \in \Lambda_0$ gegen die *Alternative* $H_1 : F \in \Lambda_1$,“

bedeutet, daß man an Hand der konkreten Stichprobe $(x_1, \dots, x_n)$ versucht, eine Entscheidung zu fällen, ob die Verteilungsfunktion der Zufallsvariable X_i zu Λ_0 oder zu Λ_1 gehört. Dies passiert auf Grund einer statistischen *Entscheidungsregel*

$$\varphi : \mathbb{R}^n \rightarrow [0, 1],$$

die eine Statistik mit folgender Interpretation ist:

Der Stichprobenraum $\mathbb{R}^n$ wird in drei disjunkte Bereiche K_0, K_{01} und K_1 unterteilt, sodaß $\mathbb{R}^n = K_0 \cup K_{01} \cup K_1$, wobei

$$\begin{aligned} K_0 &= \varphi^{-1}(\{0\}) &= \{x \in \mathbb{R}^n : \varphi(x) = 0\}, \\ K_1 &= \varphi^{-1}(\{1\}) &= \{x \in \mathbb{R}^n : \varphi(x) = 1\}, \\ K_{01} &= \varphi^{-1}((0, 1)) &= \{x \in \mathbb{R}^n : 0 < \varphi(x) < 1\}. \end{aligned}$$

Dementsprechend wird $H_0 : F \in \Lambda_0$

- verworfen, falls $\varphi(x) = 1$, also $x \in K_1$,
- nicht verworfen, falls $\varphi(x) = 0$, also $x \in K_0$;
- falls $\varphi(x) \in (0, 1)$, also $x \in K_{01}$, wird $\varphi(x)$ als Bernoulli-Wahrscheinlichkeit interpretiert, und es wird eine Zufallsvariable $Y \sim \text{Bernoulli}(\varphi(x))$ generiert, für die gilt:

$$Y = \begin{cases} 1 & \implies H_0 \text{ wird verworfen} \\ 0 & \implies H_0 \text{ wird nicht verworfen} \end{cases}$$

Falls $K_{01} \neq \emptyset$, wird eine solche Entscheidungsregel *randomisiert* genannt. Bei $K_{01} = \emptyset$, also $\mathbb{R}^n = K_0 \cup K_1$ spricht man dagegen von *nicht-randomisierten* Tests. Dabei heißt K_0 bzw. K_1 *Annahmehereich* bzw. *Ablehnungsbereich* (*kritischer Bereich*) von H_0 . K_{01} heißt *Randomisierungsbereich*.

Bemerkung 5.1.1. 1. Man sagt absichtlich „ H_0 wird nicht verworfen“, statt „ H_0 wird akzeptiert“, weil die schließende Statistik generell keine positiven, sondern nur negative Entscheidungen treffen kann. Dies ist generell ein philosophisches Problem der Falsifizierbarkeit von Hypothesen oder wissenschaftlichen Theorien, von denen aber keiner behaupten kann, daß sie der Wahrheit entsprechen (vergleiche die wissenschaftliche Erkenntnistheorie von Karl Popper (1902-1994)).

2. Die randomisierten Tests sind hauptsächlich von theoretischem Interesse (vergleiche Abschnitt 2.3). In der Praxis werden meistens nichtrandomisierte Regeln verwendet, bei denen man aus der Stichprobe $(x_1, \dots, x_n)$ allein die Entscheidung über H_0 treffen kann. Hier gilt $\varphi(x) = \mathbb{I}_{K_1}, x = (x_1, \dots, x_n) \in \mathbb{R}^n$.

In diesem und in folgendem Abschnitt betrachten wir ausschließlich nichtrandomisierte Tests, um in Abschnitt 2.3 zu der allgemeinen Situation zurückzukehren.

Definition 5.1.1

Man sagt, daß die nicht-randomisierte Testregel $\varphi : \mathbb{R}^n \rightarrow \{0, 1\}$ einen (*nichtrandomisierten*) statistischen Test zum Signifikanzniveau α angibt, falls für $F \in \Lambda_0$ gilt

$$\mathbb{P}_F(\varphi(X_1, \dots, X_n) = 1) = P(H_0 \text{ verwerfen} \mid H_0 \text{ richtig}) \leq \alpha.$$

Definition 5.1.2 1. Wenn man H_0 verwirft, obwohl H_0 richtig ist, begeht man den sogenannten *Fehler 1. Art*. Die Wahrscheinlichkeit

$$\alpha_n(F) = \mathbb{P}_F(\varphi(x_1, \dots, x_n) = 1), \quad F \in \Lambda_0$$

heißt die *Wahrscheinlichkeit des Fehlers 1. Art* und soll unter dem Niveau α bleiben.

2. Den *Fehler 2. Art* begeht man, wenn man die falsche Hypothese H_0 nicht verwirft. Dabei ist

$$\beta_n(F) = \mathbb{P}_F(\varphi(x_1, \dots, x_n) = 0), \quad F \in \Lambda_1$$

die *Wahrscheinlichkeit des Fehlers 2. Art*.

Eine Zusammenfassung aller Möglichkeiten wird in folgender Tabelle festgehalten:

	H_0 richtig	H_0 falsch
H_0 verwerfen	Fehler 1. Art, Wahrscheinlichkeit $\alpha_n(F) \leq \alpha$	richtige Entscheidung
H_0 nicht verwerfen	richtige Entscheidung	Fehler 2. Art mit Wahrscheinlichkeit $\beta_n(F)$

Dabei sollen α_n und β_n möglichst klein sein, was gegenläufige Tendenzen darstellt, weil beim Kleinwerden von α die Wahrscheinlichkeit des Fehlers 2. Art notwendigerweise wächst.

Definition 5.1.3 1. Die Funktion

$$G_n(F) = \mathbb{P}_F(\varphi(X_1, \dots, X_n) = 1), \quad F \in \Lambda$$

heißt *Gütefunktion* eines Tests φ .

2. Die Einschränkung von G_n auf Λ_1 heißt *Stärke*, *Schärfe* oder *Macht* (englisch *power*) des Tests φ .

Es gilt

$$\begin{cases} G_n(F) = \alpha_n(F) \leq \alpha, & F \in \Lambda_0 \\ G_n(F) = 1 - \beta_n(F), & F \in \Lambda_1 \end{cases}$$

Beispiel 5.1.1

Parametrische Tests. Wie sieht ein parametrischer Test aus? Der Parameterraum Θ wird als $\Theta_0 \cup \Theta_1$ dargestellt, wobei $\Theta_0 \cap \Theta_1 = \emptyset$. Es gilt $\Lambda_0 = \{F_\theta : \theta \in \Theta_0\}$, $\Lambda_1 = \{F_\theta : \theta \in \Theta_1\}$. P_F wird zu P_θ , α_n , G_n und β_n werden statt auf Λ auf Θ definiert.

Welche Hypothesen H_0 und H_1 kommen oft bei parametrischen Tests vor? Zur Einfachheit betrachten wir den Spezialfall $\Theta = \mathbb{R}$.

1. $H_0 : \theta = \theta_0$ vs. $H_1 : \theta \neq \theta_0$
2. $H_0 : \theta \geq \theta_0$ vs. $H_1 : \theta < \theta_0$
3. $H_0 : \theta \leq \theta_0$ vs. $H_1 : \theta > \theta_0$
4. $H_0 : \theta \in [a, b]$ vs. $H_1 : \theta \notin [a, b]$

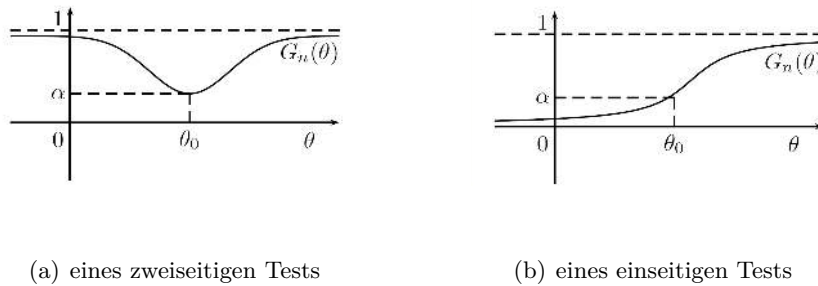
Im Fall (1) heißt der parametrische Test *zweiseitig*, in den Fällen (2) und (3) *einseitig* (*rechts-* bzw. *linksseitig*). In Fall (4) spricht man von der *Intervallhypothese* H_0 .

Bei einem zweiseitigen bzw. einseitigen Test kann die Gütefunktion wie in Abbildung 5.1 (a) bzw. 5.1 (b) aussehen,

Bei einem allgemeinen (nicht notwendigerweise parametrischen) Modell kann man die ideale Gütefunktion wie in Abbildung 5.2 schematisch darstellen.

- Man sieht aus Definition 5.1.2, dem Fehler 1. und 2. Art und der Ablehnungsregel, daß die Hypothesen H_0 und H_1 nicht symmetrisch behandelt werden, denn nur die Wahrscheinlichkeit des Fehlers 1. Art wird kontrolliert. Dies ist der Grund dafür, daß Statistiker die eigentlich interessierende Hypothese nicht als H_0 , sondern als H_1 formulieren, damit, wenn man sich für H_1 entscheidet, man mit Sicherheit sagen kann, daß die Wahrscheinlichkeit der Fehlentscheidung unter dem Niveau α liegt.

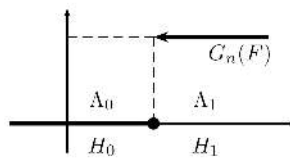
Abb. 5.1: Gütefunktion



(a) eines zweiseitigen Tests

(b) eines einseitigen Tests

Abb. 5.2: Schematische Darstellung der idealen Gütefunktion



- Wie wird ein statistischer, nicht randomisierter Test praktisch konstruiert? Die Konstruktion der Ablehnungsregel φ ähnelt sich sehr der von Konfidenzintervallen:
 1. Finde eine Teststatistik $T : \mathbb{R}^n \rightarrow \mathbb{R}$, die unter H_0 eine (möglicherweise asymptotisch für $n \rightarrow \infty$) bestimmte Prüfverteilung hat.
 2. Definiere $B_0 = [t_{\alpha_1}, t_{1-\alpha_2}]$, wobei t_{α_1} und $t_{1-\alpha_2}$ Quantile der Prüfverteilung von T sind, $\alpha_1 + \alpha_2 = \alpha \in [0, 1]$.
 3. Falls $T(X_1, \dots, X_n) \in \mathbb{R} \setminus B_0 = B_1$, setze $\varphi(X_1, \dots, X_n) = 1$. H_0 wird verworfen. Ansonsten setze $\varphi(X_1, \dots, X_n) = 0$.
- Falls die Verteilung von T nur asymptotisch bestimmt werden kann, so heißt φ *asymptotischer Test*.
- Sehr oft aber ist auch die asymptotische Verteilung von T nicht bekannt. Dann verwendet man sogenannte *Monte-Carlo Tests*, in denen dann Quantile t_α näherungsweise aus sehr vielen Monte-Carlo-Simulationen von T (unter H_0) bestimmt werden: Falls $t^i, i = 1, \dots, m$ die Werte von T in m unabhängigen Simulationsvorgängen sind, das heißt $t^i = T(x_1^i, \dots, x_n^i)$, x_j^i sind unabhängige Realisierungen von $X_j \sim F \in \Lambda_0$, $j =$

$1, \dots, n, i = 1, \dots, m$ dann bildet man ihre Ordnungsstatistiken $t^{(1)}, \dots, t^{(m)}$ und setzt $t_\alpha \approx t^{(\lfloor \alpha \cdot m \rfloor)}, \alpha \in [0, 1]$, wobei $t^{(0)} = -\infty$.

Bemerkung 5.1.2. Man sieht deutlich, daß aus einem beliebigen Konfidenzintervall

$$I_\theta = [I_1^\theta(X_1, \dots, X_n), I_2^\theta(X_1, \dots, X_n)]$$

zum Niveau $1 - \alpha$ für einen Parameter $\theta \in \mathbb{R}$ ein Test für θ konstruierbar ist. Die Hypothese $H_0 : \theta = \theta_0$ vs. $H_1 : \theta \neq \theta_0$ wird mit folgender Entscheidungsregel getestet:

$$\varphi(X_1, \dots, X_n) = 1, \text{ falls } \theta_0 \notin [I_1^{\theta_0}(X_1, \dots, X_n), I_2^{\theta_0}(X_1, \dots, X_n)].$$

Das Signifikanzniveau des Tests ist α .

Beispiel 5.1.2

Normalverteilung, Test des Erwartungswertes bei bekannter Varianz. Es seien

$$X_1, \dots, X_n \sim N(\mu, \sigma^2)$$

mit bekannter Varianz σ^2 . Ein Konfidenzintervall für μ ist

$$I^\mu = [I_1^\mu(X_1, \dots, X_n), I_2^\mu(X_1, \dots, X_n)] = \left[\bar{X}_n - \frac{z_{1-\alpha/2} \cdot \sigma}{\sqrt{n}}, \bar{X}_n + \frac{z_{1-\alpha/2} \cdot \sigma}{\sqrt{n}} \right]$$

(vergleiche Abschnitt 4.2.1) $H_0 : \mu = \mu_0$ (gegen die Alternative $H_1 : \mu \neq \mu_0$), wird verworfen, falls $|\mu_0 - \bar{X}_n| > \frac{z_{1-\alpha/2} \cdot \sigma}{\sqrt{n}}$. In der Testsprache bedeutet es, dass

$$\varphi(x_1, \dots, x_n) = \mathbb{I}((x_1, \dots, x_n) \in K_1),$$

wobei

$$K_1 = \left\{ (x_1, \dots, x_n) \in \mathbb{R}^n : |\mu_0 - \bar{x}_n| > \frac{\sigma z_{1-\alpha/2}}{\sqrt{n}} \right\}$$

der Ablehnungsbereich ist. Für die Teststatistik $T(X_1, \dots, X_n)$ gilt:

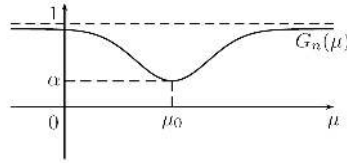
$$T(X_1, \dots, X_n) = \frac{\bar{X}_n - \mu_0}{\sigma} \sqrt{n} \sim N(0, 1) \mid \text{unter } H_0,$$

$$\alpha_n(\mu) = \alpha.$$

Berechnen wir nun die Gütefunktion (vergleiche Abbildung 5.3).

$$\begin{aligned} G_n(\mu) &= \mathbb{P}_\mu \left(|\mu_0 - \bar{X}_n| > \frac{z_{1-\alpha/2}}{\sqrt{n}} \right) = 1 - \mathbb{P}_\mu \left(\left| \bar{X}_n - \mu_0 \right| \leq \frac{\sigma z_{1-\alpha/2}}{\sqrt{n}} \right) \\ &= 1 - \mathbb{P}_\mu \left(\left| \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} + \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right| \leq z_{1-\alpha/2} \right) \\ &= 1 - \mathbb{P}_\mu \left(-z_{1-\alpha/2} - \frac{\mu - \mu_0}{\sigma} \sqrt{n} \leq \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \leq z_{1-\alpha/2} - \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right) \\ &= 1 - \Phi \left(z_{1-\alpha/2} - \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right) + \Phi \left(-z_{1-\alpha/2} - \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right) \\ &= \Phi \left(-z_{1-\alpha/2} + \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right) + \Phi \left(-z_{1-\alpha/2} - \frac{\mu - \mu_0}{\sigma} \sqrt{n} \right). \end{aligned}$$

Abb. 5.3: Gütefunktion für den zweiseitigen Test des Erwartungswertes einer Normalverteilung bei bekannter Varianz



Die „Ja-Nein“-Entscheidung des Testens wird oft als zu grob empfunden. Deswegen versucht man, ein feineres Maß der Verträglichkeit der Daten mit den Hypothesen H_0 und H_1 zu bestimmen. Dies ist der sogenannte p -Wert, der von den meisten Statistik-Softwarepaketen ausgegeben wird.

Definition 5.1.4

Es sei $(x_1, \dots, x_n)$ die konkrete Stichprobe von Daten, die als Realisierung von $(X_1, \dots, X_n)$ interpretiert wird und $T(X_1, \dots, X_n)$ die Teststatistik, mit deren Hilfe die Entscheidungsregel φ konstruiert wurde. Der p -Wert des statistischen Tests φ ist das kleinste Signifikanzniveau, zu dem der Wert $t = T(x_1, \dots, x_n)$ zur Verwerfung der Hypothese H_0 führt.

Im Beispiel eines einseitigen Tests mit dem Ablehnungsbereich $B_1 = (t, \infty)$ sagt man grob, daß

$$p = „\mathbb{P}(T(X_1, \dots, X_n) \geq t \mid H_0)“,$$

wobei die Anführungszeichen bedeuten, daß dies keine klassische, sondern eine bedingte Wahrscheinlichkeit ist, die später präzise angegeben wird.

Bei der Verwendung des p -Wertes verändert sich die Ablehnungsregel: die Hypothese H_0 wird zum Signifikanzniveau α abgelehnt, falls $\alpha \geq p$. Früher hat man die Signifikanz der Testentscheidung (Ablehnung von H_0) an Hand folgender Tabelle festgesetzt:

p -Wert	Interpretation
$p \leq 0,001$	sehr stark signifikant
$0,001 < p \leq 0,01$	stark signifikant
$0,01 < p \leq 0,05$	schwach signifikant
$0,05 < p$	nicht signifikant

Da aber heute der p -Wert an sich verwendet werden kann, kann der Anwender der Tests bei vorgegebenem p -Wert selbst entscheiden, zu welchem Niveau er seine Tests durchführen will.

Bemerkung 5.1.3. 1. Das Signifikanzniveau darf nicht in Abhängigkeit von p festgelegt werden. Dies würde die allgemeine Testphilosophie zerstören!

2. Der p -Wert ist keine Wahrscheinlichkeit, sondern eine Zufallsvariable, denn er hängt von $(X_1, \dots, X_n)$ ab. Der Ausdruck $p = \mathbb{P}(T(X_1, \dots, X_n) \geq t \mid H_0)$, der in Definition 5.1.4 für den p -Wert eines einseitigen Tests mit Teststatistik T gegeben wurde, soll demnach

als Überschreitungswahrscheinlichkeit interpretiert werden, daß bei Wiederholung des Zufallsexperiments unter H_0 der Wert $t = T(x_1, \dots, x_n)$ oder extremere Werte in Richtung der Hypothese H_1 beobachtet werden:

$$p = \mathbb{P}(T(X'_1, \dots, X'_n) \geq T(x_1, \dots, x_n) \mid H_0),$$

wobei $(X'_1, \dots, X'_n) \stackrel{d}{=} (X_1, \dots, X_n)$. Falls wir von einer konkreten Realisierung $(x_1, \dots, x_n)$ zur Zufallsstichprobe $(X_1, \dots, X_n)$ übergehen, erhalten wir

$$p = p(X_1, \dots, X_n) = \mathbb{P}(T(X'_1, \dots, X'_n) \geq T(X_1, \dots, X_n) \mid H_0)$$

3. Für andere Hypothesen H_0 wird der p -Wert auch eine andere Form haben. Zum Beispiel für

a) einen symmetrischen zweiseitigen Test ist

$$B_0 = [-t_{1-\alpha/2}, t_{1-\alpha/2}]$$

der Akzeptanzbereich für H_0 .

$$\Rightarrow p = P(|T(X'_1, \dots, X'_n)| \geq t \mid H_0), \quad t = T(X_1, \dots, X_n)$$

b) einen linksseitigen Test mit $B_0 = [t_\alpha, \infty]$ gilt

$$p = P(T(X'_1, \dots, X'_n) \leq t \mid H_0), \quad t = T(X_1, \dots, X_n)$$

c) Das Verhalten des p -Wertes kann folgendermaßen untersucht werden:

Lemma 5.1.1

Falls die Verteilungsfunktion F von T stetig und streng monoton steigend ist (die Verteilung von T ist absolut stetig mit zum Beispiel stetiger Dichte), dann ist $p \sim U[0, 1]$.

Beweis Wir zeigen es am speziellen Beispiel des rechtsseitigen Tests.

$$\begin{aligned} \mathbb{P}(p \leq \alpha \mid H_0) &= \mathbb{P}(\bar{F}_T(T(X_1, \dots, X_n)) \leq \alpha \mid H_0) \\ &= \mathbb{P}(F_T(T(X_1, \dots, X_n)) \geq 1 - \alpha \mid H_0) \\ &= \mathbb{P}(U \geq 1 - \alpha) = 1 - (1 - \alpha) = \alpha, \quad \alpha \in [0, 1], \end{aligned}$$

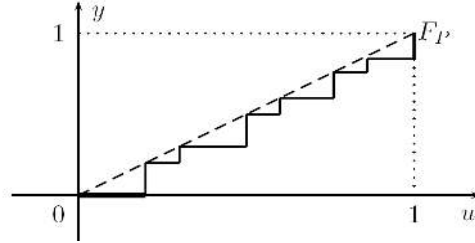
da $F_T(T(X_1, \dots, X_n)) \stackrel{d}{=} U \sim U[0, 1]$ und F_T absolut stetig ist. $\square$

Übung 5.1.1. Zeigen Sie, daß für eine beliebige Zufallsvariable X mit absolut stetiger Verteilung und streng monoton steigender Verteilungsfunktion F_X gilt:

$$F_X(X) \sim U[0, 1]$$

Falls die Verteilung von T diskret ist mit dem Wertebereich $\{t_1, \dots, t_n\}$, $t_i < t_j$ für $i < j$, so ist auch die Verteilung von p diskret, somit gilt nicht $p \sim U[0, 1]$. In diesem Fall ist $F_p(x)$ eine Treppenfunktion, die die Gerade $y = u$ in den Punkten

$$u = \sum_{i=1}^k \mathbb{P}(T(X_1, \dots, X_n) = t_i), \quad k = 1, \dots, n \text{ berührt (vgl. Abbildung 5.4).}$$

Abb. 5.4: Verteilung von p für diskrete T 

Definition 5.1.5 1. Falls die Macht $G_n(\cdot)$ eines Tests φ zum Niveau α die Ungleichung

$$G_n(F) \geq \alpha, \quad F \in \Lambda_1$$

erfüllt, dann heißt der Test *unverfälscht*.

2. Es seien φ und φ^* zwei Tests zum Niveau α mit Gütefunktionen $G_n(\cdot)$ und $G_n^*(\cdot)$. Man sagt, daß der Test φ *besser* als φ^* ist, falls er eine größere Macht besitzt:

$$G_n(F) \geq G_n^*(F) \quad \forall F \in \Lambda_1.$$

3. Der Test φ heißt *konsistent*, falls $G_n(F) \xrightarrow{n \rightarrow \infty} 1$ für alle $F \in \Lambda_1$.

Bemerkung 5.1.4. 1. Die einseitigen Tests haben oft eine größere Macht als ihre zweiseitigen Versionen.

Beispiel 5.1.3

Betrachten wir zum Beispiel den Gauß-Test des Erwartungswertes der Normalverteilung bei bekannter Varianz. Beim zweiseitigen Test

$$H_0 : \mu = \mu_0 \text{ vs. } H_1 : \mu \neq \mu_0.$$

erhalten wir die Gütefunktion

$$G_n(\mu) = \Phi\left(-z_{1-\alpha/2} + \sqrt{n} \frac{\mu - \mu_0}{\sigma}\right) + \Phi\left(-z_{1-\alpha/2} - \sqrt{n} \frac{\mu - \mu_0}{\sigma}\right).$$

Beim einseitigen Test φ^* der Hypothesen

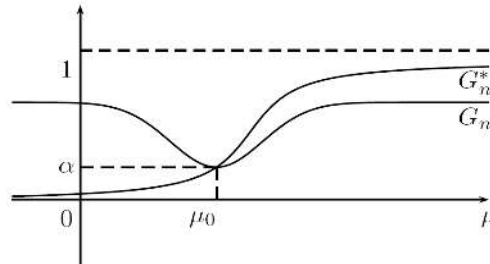
$$H_0^* : \mu \leq \mu_0 \text{ vs. } H_1^* : \mu > \mu_0$$

ist seine Gütefunktion gleich

$$G_n^*(\mu) = \Phi\left(-z_{1-\alpha} + \sqrt{n} \frac{\mu - \mu_0}{\sigma}\right).$$

Beide Tests sind offensichtlich konsistent, denn $G_n(\mu) \xrightarrow{n \rightarrow \infty} 1$, $G_n^*(\mu) \xrightarrow{n \rightarrow \infty} 1$. Dabei ist φ^* besser als φ . Beide Tests sind unverfälscht (vergleiche Abbildung 5.5).

Abb. 5.5: Gütefunktionen eines ein- bzw. zweiseitigen Tests der Erwartungswertes einer Normalverteilung



2. Beim Testen einer Intervallhypothese $H_0 : \theta \in [a, b]$ vs. $H_1 : \theta \notin [a, b]$ zum Niveau α kann man wie folgt vorgehen: *Teste*

a) $H_0^a : \theta \geq a$ vs. $H_1^a : \theta < a$ zum Niveau $\alpha/2$.

b) $H_0^b : \theta \leq b$ vs. $H_1^b : \theta > b$ zum Niveau $\alpha/2$.

H_0 wird nicht abgelehnt, falls H_0^a und H_0^b nicht abgelehnt werden. Die Wahrscheinlichkeit des Fehlers 1. Art ist hier α . Die Macht dieses Tests ist im Allgemeinen schlecht.

3. Je mehr Parameter für den Aufbau der Teststatistik T geschätzt werden müssen, desto kleiner wird in der Regel die Macht.

5.2 Nichtrandomisierte Tests

5.2.1 Parametrische Signifikanztests

In diesem Abschnitt geben wir Beispiele einiger Tests, die meistens aus den entsprechenden Konfidenzintervallen für die Parameter von Verteilungen entstehen. Deshalb werden wir sie nur kurz behandeln.

1. Tests für die Parameter der Normalverteilung $N(\mu, \sigma^2)$

a) Test von μ bei unbekannter Varianz

- Hypothesen: $H_0 : \mu = \mu_0$ vs. $H_1 : \mu \neq \mu_0$.
- Teststatistik:

$$T(X_1, \dots, X_n) = \frac{\bar{X}_n - \mu_0}{S_n} \sim t_{n-1} \quad | \quad H_0$$

- Entscheidungsregel:

$$\varphi(X_1, \dots, X_n) = 1, \text{ falls } |T(X_1, \dots, X_n)| > t_{n-1, 1-\alpha/2}.$$

b) **Test von σ^2 bei bekanntem μ**

- Hypothesen: $H_0 : \sigma^2 = \sigma_0^2$ vs. $H_1 : \sigma^2 \neq \sigma_0^2$.
- Teststatistik:

$$T(X_1, \dots, X_n) = \frac{n\tilde{S}_n^2}{\sigma_0^2} \sim \chi_n^2 \quad | H_0$$

$$\text{mit } \tilde{S}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2.$$

- Entscheidungsregel:

$$\varphi(X_1, \dots, X_n) = 1, \text{ falls } T(X_1, \dots, X_n) \notin [\chi_{n,\alpha/2}^2, \chi_{n,1-\alpha/2}^2].$$

- Gütefunktion:

$$\begin{aligned} G_n(\sigma^2) &= 1 - \mathbb{P}_{\sigma^2} \left(\chi_{n,\alpha/2}^2 \leq \frac{n\tilde{S}_n^2}{\sigma_0^2} \leq \chi_{n,1-\alpha/2}^2 \right) \\ &= 1 - \mathbb{P}_{\sigma^2} \left(\frac{\chi_{n,\alpha/2}^2 \sigma_0^2}{\sigma^2} \leq \frac{n\tilde{S}_n^2}{\sigma^2} \leq \frac{\chi_{n,1-\alpha/2}^2 \sigma_0^2}{\sigma^2} \right) \\ &= 1 - F_{\chi_n^2} \left(\chi_{n,1-\alpha/2}^2 \frac{\sigma_0^2}{\sigma^2} \right) + F_{\chi_n^2} \left(\chi_{n,\alpha/2}^2 \frac{\sigma_0^2}{\sigma^2} \right) \end{aligned}$$

c) **Test von σ^2 bei unbekanntem μ**

- Hypothesen: $H_0 : \sigma^2 = \sigma_0^2$ vs. $H_1 : \sigma^2 \neq \sigma_0^2$.
- Teststatistik:

$$T(X_1, \dots, X_n) = \frac{(n-1)S_n^2}{\sigma_0^2} \sim \chi_{n-1}^2 \quad | H_0,$$

$$\text{wobei } S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2.$$

- Entscheidungsregel:

$$\varphi(X_1, \dots, X_n) = 1, \text{ falls } T(X_1, \dots, X_n) \notin [\chi_{n-1,\alpha/2}^2, \chi_{n-1,1-\alpha/2}^2].$$

Übung 5.2.1. i. Finden Sie $G_n(\cdot)$ für die einseitige Version der obigen Tests.
 ii. Zeigen Sie, daß diese einseitigen Tests unverfälscht sind, die zweiseitigen aber nicht.

2. Asymptotische Tests

Bei asymptotischen Tests ist die Verteilung der Teststatistik nur näherungsweise (für große n) bekannt. Ebenso asymptotisch wird das Konfidenzniveau α erreicht. Ihre Konstruktion basiert meistens auf Verwendung der Grenzwertsätze.

Die allgemeine Vorgehensweise wird im sogenannten *Wald-Test* (genannt nach dem Statistiker Abraham Wald (1902-1980)) fixiert:

- Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe, X_i seien unabhängig und identisch verteilt für $i = 1, \dots, n$, mit $X_i \sim F_\theta$, $\theta \in \Theta \subseteq \mathbb{R}$.
- Wir testen $H_0 : \theta = \theta_0$ vs. $H_1 : \theta \neq \theta_0$. Es sei $\hat{\theta}_n = \hat{\theta}(X_1, \dots, X_n)$ ein erwartungstreuer, asymptotisch normalverteilter Schätzer für θ .

$$\frac{\hat{\theta}_n - \theta_0}{\hat{\sigma}_n} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1) \quad | H_0,$$

wobei $\hat{\sigma}_n^2$ ein konsistenter Schätzer für die Varianz von $\hat{\theta}_n$ sei.

Die Teststatistik ist

$$T(X_1, \dots, X_n) = \frac{\hat{\theta}_n(X_1, \dots, X_n) - \theta_0}{\hat{\sigma}_n}.$$

- Die Entscheidungsregel lautet: H_0 wird abgelehnt, wenn $|T(X_1, \dots, X_n)| > z_{1-\alpha/2}$, wobei $z_{1-\alpha/2} = \Phi^{-1}(1 - \alpha/2)$. Diese Entscheidungsregel soll nur bei großen n verwendet werden. Die Wahrscheinlichkeit des Fehlers 1. Art ist asymptotisch gleich α , denn $\mathbb{P}(|T(X_1, \dots, X_n)| > z_{1-\alpha/2} | H_0) \xrightarrow[n \rightarrow \infty]{} \alpha$ wegen der asymptotischen Normalverteilung von T .

Die Gütefunktion des Tests ist asymptotisch gleich

$$\lim_{n \rightarrow \infty} G_n(\theta) = 1 - \Phi\left(z_{1-\alpha/2} + \frac{\theta_0 - \theta}{\sigma}\right) + \Phi\left(-z_{1-\alpha/2} + \frac{\theta_0 - \theta}{\sigma}\right),$$

wobei $\hat{\sigma}_n^2 \xrightarrow[n \rightarrow \infty]{P} \sigma^2$.

Spezialfälle des Wald-Tests sind asymptotische Tests der Erwartungswerte bei einer Poisson- oder Bernoulliverteilten Stichprobe.

Beispiel 5.2.1 a) Bernoulliverteilung

Es seien $X_i \sim \text{Bernoulli}(p)$, $p \in [0, 1]$ unabhängige, identisch verteilte Zufallsvariablen.

- Hypothesen: $H_0 : p = p_0$ vs. $H_1 : p \neq p_0$.
- Teststatistik:

$$T(X_1, \dots, X_n) = \begin{cases} \sqrt{n} \frac{\bar{X}_n - p_0}{\sqrt{\bar{X}_n(1 - \bar{X}_n)}}, & \text{falls } \bar{X}_n \neq 0, 1, \\ 0, & \text{sonst.} \end{cases}$$

Unter H_0 gilt: $T(X_1, \dots, X_n) \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1)$.

b) Poissonverteilung

Es seien $X_i \sim \text{Poisson}(\lambda)$, $\lambda > 0$ unabhängige, identisch verteilte Zufallsvariablen.

- Hypothesen: $H_0 : \lambda = \lambda_0$ vs. $H_1 : \lambda \neq \lambda_0$
- Teststatistik:

$$T(X_1, \dots, X_n) = \begin{cases} \sqrt{n} \frac{\bar{X}_n - \lambda_0}{\sqrt{\bar{X}_n}}, & \text{falls } \bar{X}_n > 0, \\ 0, & \text{sonst.} \end{cases}$$

Unter H_0 gilt: $T(X_1, \dots, X_n) \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1)$

3. Zwei-Stichproben-Probleme

Gegeben seien zwei Zufallsstichproben

$$Y_1 = (X_{11}, \dots, X_{1n_1}), Y_2 = (X_{21}, \dots, X_{2n_2}), n = \max\{n_1, n_2\}.$$

X_{ij} seien unabhängig für $j = 1, \dots, n_i$, $X_{ij} \sim F_{\theta_i}$, $i = 1, 2$.

a) Test der Gleichheit zweier Erwartungswerte bei normalverteilten Stichproben

• bei bekannten Varianzen

Es seien $X_{ij} \sim N(\mu_i, \sigma_i^2)$, $i = 1, 2$, $j = 1, \dots, n$. Dabei seien σ_1^2, σ_2^2 bekannt, X_{ij} seien unabhängig voneinander für alle i, j .

Die Hypothesen sind $H_0 : \mu_1 = \mu_2$ vs. $H_1 : \mu_1 \neq \mu_2$. Wir betrachten die Teststatistik:

$$T(Y_1, Y_2) = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2}}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$$

Unter H_0 gilt: $T(Y_1, Y_2) \sim N(0, 1)$. Als Entscheidungsregel gilt: H_0 wird abgelehnt, falls $|T(Y_1, Y_2)| > z_{1-\alpha/2}$.

• bei unbekannten (jedoch gleichen) Varianzen

Es seien $X_{ij} \sim N(\mu_i, \sigma_i^2)$, $i = 1, 2$, $j = 1, \dots, n$. Dabei seien σ_1^2, σ_2^2 unbekannt, $\sigma_1^2 = \sigma_2^2$ und X_{ij} seien unabhängig voneinander für alle i, j .

Die Hypothesen sind: $H_0 : \mu_1 = \mu_2$ vs. $H_1 : \mu_1 \neq \mu_2$. Wir betrachten die Teststatistik

$$T(Y_1, Y_2) = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2}}{S_{n_1 n_2}} \sqrt{\frac{n_1 n_2}{n_1 + n_2}},$$

wobei

$$S_{n_1 n_2}^2 = \frac{1}{n_1 + n_2 - 2} \cdot \left(\sum_{j=1}^{n_1} (X_{1j} - \bar{X}_{1n_1})^2 + \sum_{j=1}^{n_2} (X_{2j} - \bar{X}_{2n_2})^2 \right).$$

Man kann zeigen, daß unter H_0 gilt: $T(Y_1, Y_2) \sim t_{n_1+n_2-2}$. Die Entscheidungsregel lautet: H_0 ablehnen, falls $|T(Y_1, Y_2)| > t_{n_1+n_2-2, 1-\alpha/2}$.

b) Test der Gleichheit von Erwartungswerten bei verbundenen Stichproben

Es seien $Y_1 = (X_{11}, \dots, X_{1n})$ und $Y_2 = (X_{21}, \dots, X_{2n})$, $n_1 = n_2 = n$,

$$Z_j = X_{1j} - X_{2j} \sim N(\mu_1 - \mu_2, \sigma^2), j = 1, \dots, n$$

unabhängig und identisch verteilt mit $\mu_i = \mathbb{E} X_{ij}$, $i = 1, 2$. Die Hypothesen sind: $H_0 : \mu_1 = \mu_2$ vs. $H_1 : \mu_1 \neq \mu_2$ bei unbekannter Varianz σ^2 . Als Teststatistik verwenden wir

$$T(Z_1, \dots, Z_n) = \sqrt{n} \frac{\bar{Z}_n}{S_n},$$

wobei

$$S_n^2 = \frac{1}{n-1} \sum_{j=1}^n (Z_j - \bar{Z}_n)^2.$$

Unter H_0 gilt dann: $T(Z_1, \dots, Z_n) \sim t_{n-1}$. Die Entscheidungsregel lautet: H_0 wird abgelehnt, falls $|T(z_1, \dots, z_n)| > t_{n-1, 1-\alpha/2}$.

c) **Test der Gleichheit von Varianzen bei unabhängigen Gaußschen Stichproben**

Es seien $Y_1 = (X_{11}, \dots, X_{1n_1})$ und $Y_2 = (X_{21}, \dots, X_{2n_2})$ unabhängig und identisch verteilt mit $X_{ij} \sim N(\mu_i, \sigma_i^2)$, wobei μ_i und σ_i^2 beide unbekannt sind. Die Hypothesen sind: $H_0 : \sigma_1^2 = \sigma_2^2$ vs. $H_1 : \sigma_1^2 \neq \sigma_2^2$. Als Teststatistik verwenden wir

$$T(Y_1, Y_2) = \frac{S_{2n_2}^2}{S_{1n_1}^2},$$

wobei

$$S_{in_i}^2 = \frac{1}{n_i - 1} \sum_{j=1}^{n_i} (X_{ij} - \bar{X}_{in_i})^2, \quad i = 1, 2.$$

Unter H_0 gilt: $T(Y_1, Y_2) \sim F_{n_2-1, n_1-1}$. Die Entscheidungsregel lautet: H_0 wird abgelehnt, falls $T(Y_1, Y_2) \notin [F_{n_2-1, n_1-1, \alpha/2}, F_{n_2-1, n_1-1, 1-\alpha/2}]$.

d) **Asymptotische Zwei-Stichproben-Tests**

• **bei Bernoulli-verteilten Stichproben**

Es gilt $X_{ij} \sim \text{Bernoulli}(p_i)$, $j = 1, \dots, n_i$, $i = 1, 2$. Die Hypothesen sind $H_0 : p_1 = p_2$ vs. $H_1 : p_1 \neq p_2$. Als Teststatistik verwenden wir

$$T(Y_1, Y_2) = \frac{(\bar{X}_{1n_1} - \bar{X}_{2n_2})(1 - \mathbb{I}(\bar{X}_{1n_1} = o, \bar{X}_{2n_2} = o))}{\sqrt{\frac{\bar{X}_{1n_1}(1-\bar{X}_{1n_1})}{n_1} + \frac{\bar{X}_{2n_2}(1-\bar{X}_{2n_2})}{n_2}}}$$

Unter H_0 gilt: $T(Y_1, Y_2) \xrightarrow[n_1, n_2 \rightarrow \infty]{d} Y \sim N(0, 1)$. Die Entscheidungsregel lautet: H_0 wird verworfen, falls $|T(Y_1, Y_2)| > z_{1-\alpha/2}$. Dies ist ein Test zum asymptotischen Signifikanzniveau α .

• **bei Poisson-verteilten Stichproben**

Es seien X_{ij} unabhängig, $X_{ij} \sim \text{Poisson}(\lambda_i)$, $i = 1, 2$. Die Hypothesen sind: $H_0 : \lambda_1 = \lambda_2$ vs. $H_1 : \lambda_1 \neq \lambda_2$. Als Teststatistik verwenden wir:

$$T(Y_1, Y_2) = \frac{\bar{X}_{1n_1} - \bar{X}_{2n_2}}{\sqrt{\frac{\bar{X}_{1n_1}}{n_1} + \frac{\bar{X}_{2n_2}}{n_2}}}$$

Die Entscheidungsregel lautet: H_0 ablehnen, falls $|T(Y_1, Y_2)| > z_{1-\alpha/2}$. Dies ist ein Test zum asymptotischen Niveau α .

Bemerkung 5.2.1. *Asymptotische Tests dürfen nur für große Stichprobenumfänge verwendet werden. Bei ihrer Verwendung für kleine Stichproben kann das asymptotische Signifikanzniveau nicht garantiert werden.*

5.3 Randomisierte Tests

In diesem Abschnitt werden wir klassische Ergebnisse von Neyman-Pearson über die besten Tests präsentieren. Dabei werden randomisierte Tests eine wichtige Rolle spielen.

5.3.1 Grundlagen

Gegeben sei eine Zufallsstichprobe $(X_1, \dots, X_n)$ von unabhängigen und identisch verteilten Zufallsvariablen X_i mit konkreter Ausprägung $(x_1, \dots, x_n)$. Sei unser Stichprobenraum $(B, \mathcal{B})$ entweder $(\mathbb{R}^n, \mathcal{B}_{\mathbb{R}^n})$ oder $(\mathbb{N}_0^n, \mathcal{B}_{\mathbb{N}_0^n})$, je nachdem, ob die Stichprobenvariablen X_i , $i = 1, \dots, n$ absolut stetig oder diskret verteilt sind.

Hier wird zur Einfachheit im Falle einer diskret verteilten Zufallsvariable X_i ihr diskreter Wertebereich mit $\mathbb{N}_0 = \mathbb{N} \cup \{0\}$ gleichgesetzt. Der Wertebereich sei mit einem Maß μ versehen, wobei

$$\mu = \begin{cases} \text{Lebesgue-Maß auf } \mathbb{R}, & \text{falls } B = \mathbb{R}^n, \\ \text{Zählmaß auf } \mathbb{N}_0, & \text{falls } B = \mathbb{N}_0^n. \end{cases}$$

Dementsprechend gilt

$$\int g(x) \mu(dx) = \begin{cases} \int_{\mathbb{R}} g(x) dx, & \text{im absolut stetigen Fall,} \\ \sum_{x \in \mathbb{N}_0} g(x), & \text{im diskreten Fall.} \end{cases}$$

Es sei zusätzlich $X_i \sim F_\theta$, $\theta \in \Theta \subseteq \mathbb{R}^m$, $i = 1, \dots, n$ (parametrisches Modell). Für $\Theta = \Theta_0 \cup \Theta_1$, $\Theta_0 \cap \Theta_1 = \emptyset$ formulieren wir die Hypothesen $H_0 : \theta \in \Theta_0$ vs. $H_1 : \theta \in \Theta_1$, die mit Hilfe eines randomisierten Tests

$$\varphi(x) = \begin{cases} 1, & x \in K_1, \\ \gamma \in (0, 1), & x \in K_{01} \\ 0, & x \in K_0 \end{cases} \quad x = (x_1, \dots, x_n),$$

getestet werden.

Im Falle $x \in K_{01}$ wird mit Hilfe einer Zufallsvariable $Y \sim \text{Bernoulli}(\varphi(x))$ entschieden, ob H_0 verworfen wird ($Y = 1$) oder nicht ($Y = 0$).

Definition 5.3.1 1. Die *Gütefunktion* eines randomisierten Tests φ sei

$$G_n(\theta) = G_n(\varphi, \theta) = \mathbb{E}_\theta \varphi(X_1, \dots, X_n), \quad \theta \in \Theta.$$

2. Der Test φ hat das *Signifikanzniveau* $\alpha \in [0, 1]$, falls $G_n(\varphi, \theta) \leq \alpha$, $\forall \theta \in \Theta_0$ ist. Die Zahl

$$\sup_{\theta \in \Theta_0} G_n(\varphi, \theta)$$

wird *Umfang* des Tests φ genannt. Offensichtlich ist der Umfang eines Niveau- α -Tests kleiner gleich α .

3. Sei $\Psi(\alpha)$ die Menge aller Tests zum Niveau α . Der Test $\varphi_1 \in \Psi(\alpha)$ ist (*gleichmäßig*) *besser* als Test $\varphi_2 \in \Psi(\alpha)$, falls $G_n(\varphi_1, \theta) \geq G_n(\varphi_2, \theta)$, $\theta \in \Theta_1$, also falls φ_1 eine größere Macht besitzt.

4. Ein Test $\varphi^* \in \Psi(\alpha)$ ist (gleichmäßig) bester Test in $\Psi(\alpha)$, falls

$$G_n(\varphi^*, \theta) \geq G_n(\varphi, \theta), \text{ für alle Tests } \varphi \in \Psi(\alpha), \theta \in \Theta_1.$$

Bemerkung 5.3.1. 1. Definition 5.3.1 1) ist eine offensichtliche Verallgemeinerung der Definition 5.1.3 der Gütefunktion eines nicht-randomisierten Tests φ . Nämlich, für $\varphi(x) = \mathbb{I}(x \in K_1)$ gilt:

$$\begin{aligned} G_n(\varphi, \theta) &= \mathbb{E}_\theta \varphi(X_1, \dots, X_n) \\ &= \mathbb{P}_\theta((X_1, \dots, X_n) \in K_1) \\ &= \mathbb{P}_\theta(H_0 \text{ ablehnen}), \theta \in \Theta. \end{aligned}$$

2. Ein bester Test φ^* in $\Psi(\alpha)$ existiert nicht immer, sondern nur unter gewissen Voraussetzungen an $\mathbb{P}_\theta, \Theta_0, \Theta_1$ und $\Psi(\alpha)$.

5.3.2 Neyman-Pearson-Tests bei einfachen Hypothesen

In diesem Abschnitt betrachten wir einfache Hypothesen

$$H_0 : \theta = \theta_0 \quad \text{vs.} \quad H_1 : \theta = \theta_1 \quad (5.3.1)$$

wobei $\theta_0, \theta_1 \in \Theta, \theta_1 \neq \theta_0$.

Dementsprechend sind $\Theta_0 = \{\theta_0\}, \Theta_1 = \{\theta_1\}$. Wir setzen voraus, daß F_{θ_i} eine Dichte $g_i(x)$ bezüglich μ besitzt, $i = 0, 1$. Führen wir einige abkürzende Bezeichnungen $\mathbb{P}_0 = \mathbb{P}_{\theta_0}, \mathbb{P}_1 = \mathbb{P}_{\theta_1}, \mathbb{E}_0 = \mathbb{E}_{\theta_0}, \mathbb{E}_1 = \mathbb{E}_{\theta_1}$ ein. Sei $f_i(x) = \prod_{j=1}^n g_i(x_j), x = (x_1, \dots, x_n), i = 0, 1$ die Dichte der Stichprobe unter H_0 bzw. H_1 .

Definition 5.3.2

Ein *Neyman-Pearson-Test* (NP-Test) der einfachen Hypothesen in (5.3.1) ist gegeben durch die Regel

$$\varphi(x) = \varphi_K(x) = \begin{cases} 1, & \text{falls } f_1(x) > K f_0(x), \\ \gamma, & \text{falls } f_1(x) = K f_0(x), \\ 0, & \text{falls } f_1(x) < K f_0(x) \end{cases} \quad (5.3.2)$$

für Konstanten $K > 0$ und $\gamma \in [0, 1]$.

Bemerkung 5.3.2. 1. Manchmal werden $K = K(x)$ und $\gamma = \gamma(x)$ als Funktionen von x und nicht als Konstanten betrachtet.

2. Der Ablehnungsbereich des Neyman-Pearson-Tests φ_K ist

$$K_1 = \{x \in B : f_1(x) > K f_0(x)\}.$$

3. Der Umfang des Neyman-Pearson-Tests φ_K ist

$$\begin{aligned} \mathbb{E}_0 \varphi_K(X_1, \dots, X_n) &= \mathbb{P}_0(f_1(X_1, \dots, X_n) > K f_0(X_1, \dots, X_n)) \\ &\quad + \gamma \mathbb{P}_0(f_1(X_1, \dots, X_n) = K f_0(X_1, \dots, X_n)) \end{aligned}$$

4. Die Definition 5.3.2 kann man äquivalent folgendermaßen geben: Wir definieren eine Teststatistik

$$T(x) = \begin{cases} \frac{f_1(x)}{f_0(x)}, & x \in B : f_0(x) > 0, \\ \infty, & x \in B : f_0(x) = 0. \end{cases}$$

Dann wird der neue Test

$$\tilde{\varphi}_K(x) = \begin{cases} 1, & \text{falls } T(x) > K, \\ \gamma, & \text{falls } T(x) = K, \\ 0, & \text{falls } T(x) < K \end{cases}$$

eingeführt, der für P_0 - und P_1 - fast alle $x \in B$ äquivalent zu φ_K ist. In der Tat gilt $\varphi_K(x) = \tilde{\varphi}_K(x) \forall x \in B \setminus C$, wobei $C = \{x \in B : f_0(x) = f_1(x) = 0\}$ das $\mathbb{P}_0$ - bzw. $\mathbb{P}_1$ -Maß Null besitzt.

In der neuen Formulierung ist der Umfang von φ bzw. $\tilde{\varphi}_K$ gleich

$$\mathbb{E}_0 \tilde{\varphi}_K = \mathbb{P}_0(T(X_1, \dots, X_n) > K) + \gamma \cdot \mathbb{P}_0(T(X_1, \dots, X_n) = K).$$

Satz 5.3.1

Optimalitätssatz

Es sei φ_K ein Neyman-Pearson-Test für ein $K > 0$ und $\gamma \in [0, 1]$. Dann ist φ_K der beste Test zum Niveau $\alpha = \mathbb{E}_0 \varphi_K$ seines Umfangs.

Beweis Sei $\varphi \in \Psi(\alpha)$, also $\mathbb{E}_0(\varphi(X_1, \dots, X_n)) \leq \alpha$. Um zu zeigen, daß φ_K besser als φ ist, genügt es bei einfachen Hypothesen H_0 und H_1 zu zeigen, daß $\mathbb{E}_1 \varphi_K(X_1, \dots, X_n) \geq \mathbb{E}_1 \varphi(X_1, \dots, X_n)$. Wir führen dazu die folgenden Mengen ein:

$$M^+ = \{x \in B : \varphi_K(x) > \varphi(x)\}$$

$$M^- = \{x \in B : \varphi_K(x) < \varphi(x)\}$$

$$M^= = \{x \in B : \varphi_K(x) = \varphi(x)\}$$

Es gilt offensichtlich $x \in M^+ \Rightarrow \varphi_K(x) > 0 \Rightarrow f_1(x) \geq K f_0(x)$,

$$x \in M^- \Rightarrow \varphi_K(x) < 1 \Rightarrow f_1(x) \leq K f_0(x) \text{ und } B = M^+ \cup M^- \cup M^=.$$

Als Folgerung erhalten wir

$$\begin{aligned} \mathbb{E}_1(\varphi_K(X_1, \dots, X_n) - \varphi(X_1, \dots, X_n)) &= \int_B (\varphi_K(x) - \varphi(x)) f_1(x) \mu(dx) \\ &= \left(\int_{M^+} + \int_{M^-} + \int_{M^=} \right) (\varphi_K(x) - \varphi(x)) f_1(x) \mu(dx) \\ &\geq \int_{M^+} (\varphi_K(x) - \varphi(x)) K f_0(x) \mu(dx) \\ &\quad + \int_{M^-} (\varphi_K(x) - \varphi(x)) K f_0(x) \mu(dx) \\ &= \int_B (\varphi_K(x) - \varphi(x)) K f_0(x) \mu(dx) \\ &= K [\mathbb{E}_0 \varphi_K(X_1, \dots, X_n) - \mathbb{E}_0 \varphi(X_1, \dots, X_n)] \\ &\geq K(\alpha - \alpha) = 0, \end{aligned}$$

weil beide Tests das Niveau α haben. Damit ist die Behauptung bewiesen. $\square$

Bemerkung 5.3.3. 1. Da im Beweis γ nicht vorkommt, wird derselbe Beweis im Falle von $\gamma(x) \neq \text{const}$ gelten.

2. Aus dem Beweis folgt die Gültigkeit der Ungleichung

$$\int_B (\varphi_K(x) - \varphi(x)) (f_1(x) - K f_0(x)) \mu(dx) \geq 0$$

im Falle des konstanten K , bzw.

$$\mathbb{E}_1 (\varphi_K(X_1, \dots, X_n) - \varphi(X_1, \dots, X_n)) \geq \int_B (\varphi_K(x) - \varphi(x)) K(x) f_0(x) \mu(dx)$$

im allgemeinen Fall.

Satz 5.3.2

(Fundamentallema von Neyman-Pearson)

1. Zu einem beliebigen $\alpha \in (0, 1)$ gibt es einen Neyman-Pearson-Test φ_K mit Umfang α , der dann nach Satz 5.3.1 der beste Niveau- α -Test ist.
2. Ist φ ebenfalls bester Test zum Niveau α , so gilt $\varphi(x) = \varphi_K(x)$ für μ -fast alle $x \in K_0 \cup K_1 = \{x \in B : f_1(x) \neq K f_0(x)\}$ und φ_K aus Teil 1).

Beweis 1. Für $\varphi_K(x)$ gilt

$$\varphi_K(x) = \begin{cases} 1, & \text{falls } x \in K_1 = \{x : f_1(x) > K \cdot f_0(x)\}, \\ \gamma, & \text{falls } x \in K_{01} = \{x : f_1(x) = K \cdot f_0(x)\}, \\ 0, & \text{falls } x \in K_0 = \{x : f_1(x) < K \cdot f_0(x)\}. \end{cases}$$

Der Umfang von φ_K ist

$$\mathbb{P}_0(T(X_1, \dots, X_n) > K) + \gamma \mathbb{P}_0(T(X_1, \dots, X_n) = K) = \alpha, \quad (5.3.3)$$

wobei

$$T(x_1, \dots, x_n) = \begin{cases} \frac{f_1(x_1, \dots, x_n)}{f_0(x_1, \dots, x_n)}, & \text{falls } f_0(x_1, \dots, x_n) > 0, \\ \infty, & \text{sonst.} \end{cases}$$

Nun suchen wir ein $K > 0$ und ein $\gamma \in [0, 1]$, sodaß Gleichung (5.3.3) stimmt. Es sei $\tilde{F}_0(x) = \mathbb{P}_0(T(X_1, \dots, X_n) \leq x)$, $x \in \mathbb{R}$ die Verteilungsfunktion von T . Da $T \geq 0$ ist, gilt $\tilde{F}_0(x) = 0$, falls $x < 0$. Außerdem ist $\mathbb{P}_0(T(X_1, \dots, X_n) < \infty) = 1$, das heißt $\tilde{F}_0^{-1}(\alpha) \in [0, \infty)$, $\alpha \in (0, 1)$. Die Gleichung (5.3.3) kann dann folgendermaßen umgeschrieben werden:

$$1 - \tilde{F}_0(K) + \gamma (\tilde{F}_0(K) - \tilde{F}_0(K-)) = \alpha, \quad (5.3.4)$$

wobei $\tilde{F}_0(K-) = \lim_{x \rightarrow K-0} \tilde{F}_0(x)$.

Sei $K = \tilde{F}_0^{-1}(1 - \alpha)$, dann gilt:

- a) Falls K ein Stetigkeitspunkt von $\tilde{F}_0$ ist, ist Gleichung (5.3.4) erfüllt für alle $\gamma \in [0, 1]$, zum Beispiel $\gamma = 0$.
- b) Falls K kein Stetigkeitspunkt von $\tilde{F}_0$ ist, dann ist $\tilde{F}_0(K) - \tilde{F}_0(K-) > 0$, woraus folgt

$$\gamma = \frac{\alpha - 1 + \tilde{F}_0(K)}{\tilde{F}_0(K) - \tilde{F}_0(K-)}$$

$\Rightarrow$ es gibt einen Neyman-Pearson-Test zum Niveau α .

2. Wir definieren $M^\neq = \{x \in B : \varphi(x) \neq \varphi_K(x)\}$. Es muss gezeigt werden, daß

$$\mu((K_0 \cup K_1) \cap M^\neq) = 0.$$

Dazu betrachten wir

$$\begin{aligned} \mathbb{E}_1 \varphi(X_1, \dots, X_n) - \mathbb{E}_1 \varphi_K(X_1, \dots, X_n) &= 0 \quad (\varphi \text{ und } \varphi_K \text{ sind beste Tests}) \\ \mathbb{E}_0 \varphi(X_1, \dots, X_n) - \mathbb{E}_0 \varphi_K(X_1, \dots, X_n) &\leq 0 \quad (\varphi \text{ und } \varphi_K \text{ sind } \alpha\text{-Tests} \\ &\quad \text{mit Umfang von } \varphi_K = \alpha) \end{aligned}$$

$$\Rightarrow \int_B (\varphi - \varphi_K) \cdot (f_1 - K \cdot f_0) \mu(dx) \geq 0.$$

In Bemerkung 5.3.3 wurde bewiesen, daß

$$\begin{aligned} \int_B (\varphi - \varphi_K)(f_1 - K \cdot f_0) d\mu &\leq 0 \\ \Rightarrow \int_B (\varphi - \varphi_K)(f_1 - K \cdot f_0) d\mu &= 0 = \int_{M^\neq \cap (K_0 \cup K_1)} (\varphi - \varphi_K)(f_1 - K f_0) d\mu. \end{aligned}$$

Es gilt $\mu(M^\neq \cap (K_0 \cup K_1)) = 0$, falls der Integrand $(\varphi_K - \varphi)(f_1 - K f_0) > 0$ auf $M^\neq$ ist. Wir zeigen, daß

$$(\varphi_K - \varphi)(f_1 - K f_0) > 0 \text{ für } x \in M^\neq \quad (5.3.5)$$

ist. Es gilt

$$\begin{aligned} f_1 - K f_0 > 0 &\Rightarrow \varphi_K - \varphi > 0, \\ f_1 - K f_0 < 0 &\Rightarrow \varphi_K - \varphi < 0, \end{aligned}$$

weil

$$\begin{aligned} f_1(x) > K f_0(x) &\Rightarrow \varphi_K(x) = 1 \\ &\quad \text{und mit } \varphi(x) < 1 \Rightarrow \varphi_K(x) - \varphi(x) > 0 \text{ auf } M^\neq. \\ f_1(x) < K f_0(x) &\Rightarrow \varphi_K(x) = 0 \\ &\quad \text{und mit } \varphi(x) > 0 \Rightarrow \varphi_K(x) - \varphi(x) < 0 \text{ auf } M^\neq. \end{aligned}$$

Daraus folgt die Gültigkeit der Ungleichung (5.3.5) und somit

$$\mu((K_0 \cup K_1) \cap M^\neq) = 0.$$

□

Bemerkung 5.3.4. Falls φ und φ_K beste α -Tests sind, dann sind sie P_0 - bzw. P_1 - fast sicher gleich.

Beispiel 5.3.1 (Neyman-Pearson-Test für den Parameter der Poissonverteilung):

Es sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe mit $X_i \sim \text{Poisson}(\lambda)$, $\lambda > 0$, wobei X_i unabhängig und identisch verteilt sind für $i = 1, \dots, n$. Wir testen die Hypothesen $H_0 : \lambda = \lambda_0$ vs. $H_1 : \lambda = \lambda_1$. Dabei ist

$$g_i(x) = e^{-\lambda_i} \frac{\lambda_i^x}{x!}, \quad x \in \mathbb{N}_0, \quad i = 0, 1,$$

$$f_i(x) = f_i(x_1, \dots, x_n) = \prod_{j=1}^n g_i(x_j) = \prod_{j=1}^n e^{-\lambda_i} \frac{\lambda_i^{x_j}}{x_j!} = e^{-n\lambda_i} \cdot \frac{\lambda_i^{\sum_{j=1}^n x_j}}{(x_1! \cdot \dots \cdot x_n!)}$$

für $i = 0, 1$. Die Neyman-Pearson-Teststatistik ist

$$T(x_1, \dots, x_n) = \begin{cases} \frac{f_1(x)}{f_0(x)} = e^{-n(\lambda_1 - \lambda_0)} \cdot (\lambda_1 / \lambda_0)^{\sum_{j=1}^n x_j}, & \text{falls } x_1, \dots, x_n \in \mathbb{N}_0, \\ \infty, & \text{sonst.} \end{cases}$$

Die Neyman-Pearson-Entscheidungsregel lautet

$$\varphi_K(x_1, \dots, x_n) = \begin{cases} 1, & \text{falls } T(x_1, \dots, x_n) > K, \\ \gamma, & \text{falls } T(x_1, \dots, x_n) = K, \\ 0, & \text{falls } T(x_1, \dots, x_n) < K. \end{cases}$$

Wir wählen $K > 0$, $\gamma \in [0, 1]$, sodaß φ_K den Umfang α hat. Dazu lösen wir

$$\alpha = \mathbb{P}_0(T(X_1, \dots, X_n) > K) + \gamma \mathbb{P}_0(T(X_1, \dots, X_n) = K)$$

bezüglich γ und K auf.

$$\begin{aligned} \mathbb{P}_0(T(X_1, \dots, X_n) > K) &= \mathbb{P}_0(\log T(X_1, \dots, X_n) > \log K) \\ &= \mathbb{P}_0\left(-n(\lambda_1 - \lambda_0) + \sum_{j=1}^n X_j \cdot \log\left(\frac{\lambda_1}{\lambda_0}\right) > \log K\right) = \mathbb{P}_0\left(\sum_{j=1}^n X_j > A\right) \\ \text{wobei } A &:= \left\lfloor \frac{\log K + n \cdot (\lambda_1 - \lambda_0)}{\log \frac{\lambda_1}{\lambda_0}} \right\rfloor, \end{aligned}$$

falls zum Beispiel $\lambda_1 > \lambda_0$. Im Falle $\lambda_1 < \lambda_0$ ändert sich das $>$ auf $<$ in der Wahrscheinlichkeit.

Wegen der Faltungstabilität der Poissonverteilung ist unter H_0

$$\sum_{j=1}^n X_j \sim \text{Poisson}(n\lambda_0),$$

also wählen wir K als minimale, nichtnegative Zahl, für die gilt: $\mathbb{P}_0\left(\sum_{j=1}^n X_j > A\right) \leq \alpha$, und setzen

$$\gamma = \frac{\alpha - \mathbb{P}_0(\sum_{j=1}^n X_j > A)}{\mathbb{P}_0(\sum_{j=1}^n X_j = A)},$$

wobei

$$\mathbb{P}_0 \left(\sum_{j=1}^n X_j > A \right) = 1 - \sum_{j=0}^A e^{-\lambda_0 n} \frac{(\lambda_0 n)^j}{j!},$$

$$\mathbb{P}_0 \left(\sum_{j=1}^n X_j = A \right) = e^{-\lambda_0 n} \frac{(\lambda_0 n)^A}{A!}.$$

Somit haben wir die Parameter K und γ gefunden und damit einen Neyman-Pearson-Test φ_K konstruiert.

5.3.3 Einseitige Neyman-Pearson-Tests

Bisher betrachteten wir Neyman-Pearson-Tests für einfache Hypothesen der Form $H_i : \theta = \theta_i$, $i = 0, 1$. In diesem Abschnitt wollen wir einseitige Neyman-Pearson-Tests einführen, für Hypothesen der Form $H_0 : \theta \leq \theta_0$ vs. $H_1 : \theta > \theta_0$.

Zunächst konstruieren wir einen Test für diese Hypothesen: Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe, X_i seien unabhängig und identisch verteilt mit

$$X_i \sim F_\theta \in \Lambda = \{F_\theta : \theta \in \Theta\},$$

wobei $\Theta \subset \mathbb{R}$ offen ist und Λ eindeutig parametrisiert, das heißt

$$\theta \neq \theta' \Rightarrow F_\theta \neq F_{\theta'}.$$

Ferner besitze F_θ eine Dichte g_θ bezüglich des Lebesgue-Maßes (bzw. Zählmaßes) auf $\mathbb{R}$ (bzw. $\mathbb{N}_0$). Dann ist

$$f_\theta(x) = \prod_{j=1}^n g_\theta(x_j), \quad x = (x_1, \dots, x_n)$$

eine Dichte von $(X_1, \dots, X_n)$ bezüglich μ auf B .

Definition 5.3.3

Eine Verteilung auf B mit Dichte f_θ gehört zur Klasse von *Verteilungen mit monotonen Dichtekoeffizienten* in T , falls es für alle $\theta < \theta'$ eine Funktion $h : \mathbb{R} \times \Theta^2 \rightarrow \mathbb{R} \cup \infty$, die monoton wachsend in $t \in \mathbb{R}$ ist und eine Statistik $T : B \rightarrow \mathbb{R}$ gibt, mit der Eigenschaft

$$\frac{f_{\theta'}(x)}{f_\theta(x)} = h(T(x), \theta, \theta'),$$

wobei

$$h(T(x), \theta, \theta') = \infty \quad \text{für alle } x \in B : f_\theta(x) = 0, f_{\theta'}(x) > 0.$$

Der Fall $f_\theta(x) = f_{\theta'}(x) = 0$ tritt mit $\mathbb{P}_\Theta$ - bzw. $\mathbb{P}_{\Theta'}$ -Wahrscheinlichkeit 0 auf.

Definition 5.3.4

Es sei Q_θ eine Verteilung auf $(B, \mathcal{B})$ mit der Dichte f_θ bzgl. μ . Q_θ gehört zur *einparametrischen Exponentialklasse* ($\theta \in \Theta \subset \mathbb{R}$ offen), falls die Dichte folgende Form hat:

$$f_\theta(x) = \exp \{c(\theta) \cdot T(x) + a(\theta)\} \cdot l(x), \quad x = (x_1, \dots, x_n) \in B,$$

wobei $c(\theta)$ eine monoton steigende Funktion ist, und $\text{Var}_\theta T(X_1, \dots, X_n) > 0$, $\theta \in \Theta$.

Lemma 5.3.1

Verteilungen aus der einparametrischen Exponentialfamilie besitzen einen monotonen Dichtekoeffizienten.

Beweis Es sei Q_θ aus der einparametrischen Exponentialfamilie mit der Dichte

$$f_\theta(x) = \exp \{c(\theta) \cdot T(x) + a(\theta)\} \cdot l(x).$$

Für $\theta < \theta'$ ist dann

$$\frac{f_{\theta'}(x)}{f_\theta(x)} = \exp \{(c(\theta') - c(\theta)) \cdot T(x) + a(\theta') - a(\theta)\}$$

monoton bezüglich T , weil $c(\theta') - c(\theta) > 0$ wegen der Monotonie von $c(\theta)$. Also besitzt f_θ einen monotonen Dichtekoeffizienten. $\square$

Beispiel 5.3.2 1. *Normalverteilte Stichprobenvariablen*

Es seien $X_i \sim N(\mu, \sigma_0^2)$, $i = 1, \dots, n$, unabhängige, identisch verteilte Zufallsvariablen, mit unbekanntem Parameter μ und bekannter Varianz σ_0^2 (Hier wird μ für die Bezeichnung des Erwartungswertes von X_i und nicht des Maßes auf $\mathbb{R}^n$ verwendet. (wie früher)). Die Dichte des Zufallsvektors $X = (X_1, \dots, X_n)^\top$ ist gleich

$$\begin{aligned} f_\mu(x) &= \prod_{i=1}^n g_\mu(x_i) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma_0^2}} e^{-\frac{(x_i-\mu)^2}{2\sigma_0^2}} \\ &= \frac{1}{(2\pi\sigma_0^2)^{n/2}} \exp \left\{ -\frac{1}{2\sigma_0^2} \sum_{i=1}^n (x_i - \mu)^2 \right\} \\ &= \frac{1}{(2\pi\sigma_0^2)^{n/2}} \exp \left\{ -\frac{1}{2\sigma_0^2} \left(\sum_{i=1}^n x_i^2 - 2\mu \sum_{i=1}^n x_i + \mu^2 n \right) \right\} \\ &= \exp \left(\underbrace{\frac{\mu}{\sigma_0^2}}_{c(\mu)} \cdot \underbrace{\sum_{i=1}^n x_i}_{T(x)} - \underbrace{\frac{\mu^2 n}{2\sigma_0^2}}_{a(\mu)} \right) \cdot \underbrace{\frac{1}{(2\pi\sigma_0^2)^{n/2}} \exp \left(-\frac{\sum_{i=1}^n x_i^2}{2\sigma_0^2} \right)}_{l(x)}. \end{aligned}$$

Also gehört $N(\mu, \sigma_0^2)$ zur einparametrischen Exponentialklasse mit $c(\mu) = \frac{\mu}{\sigma_0^2}$ und $T(x) = \sum_{i=1}^n x_i$.

2. *Binomialverteilte Stichprobenvariablen*

Es seien $X_i \sim \text{Bin}(k, p)$ unabhängig und identisch verteilt, $i = 1, \dots, n$. Der Parameter p

sei unbekannt. Die Zähldichte des Zufallsvektors $X = (X_1, \dots, X_n)^\top$ ist

$$\begin{aligned} f_p(x) &= \mathbb{P}_p(X_i = x_i, i = 1, \dots, n) \\ &= \prod_{i=1}^n \binom{k}{x_i} p^{x_i} (1-p)^{k-x_i} = p^{\sum_{i=1}^n x_i} \cdot \frac{(1-p)^{nk}}{(1-p)^{\sum_{i=1}^n x_i}} \cdot \prod_{i=1}^n \binom{k}{x_i} \\ &= \exp \left\{ \underbrace{\left(\sum_{i=1}^n x_i \right)}_{T(x)} \cdot \underbrace{\log \left(\frac{p}{1-p} \right)}_{c(p)} + \underbrace{nk \cdot \log(1-p)}_{a(p)} \right\} \cdot \underbrace{\prod_{i=1}^n \binom{k}{x_i}}_{l(x)}, \end{aligned}$$

also gehört $\text{Bin}(n, p)$ zur einparametrischen Exponentialklasse mit

$$c(p) = \log \left(\frac{p}{1-p} \right)$$

und

$$T(x) = \sum_{i=1}^n x_i.$$

Lemma 5.3.2

Falls φ_K der Neyman-Pearson-Test der Hypothesen $H_0 : \theta = \theta_0$ vs. $H_1 : \theta = \theta_1$ ist, dann gilt:

$$\mu(\underbrace{\{x \in B : f_1(x) \neq K f_0(x)\}}_{K_0 \cup K_1}) > 0.$$

Beweis Wegen $\theta_0 \neq \theta_1$ und der eindeutigen Parametrisierung gilt $f_0 \neq f_1$ auf einer Menge mit μ -Maß > 0 .

Nun sei $\mu(K_0 \cup K_1) = 0$. Daraus folgt, daß $f_1(x) = K \cdot f_0(x)$ μ -fast sicher. Das heißt

$$1 = \int_B f_1(x) dx = K \cdot \int_B f_0(x) dx,$$

woraus folgt, daß $K = 1$ und $f_1(x) = f_0(x)$ μ -fast sicher, was aber ein Widerspruch zur eindeutigen Parametrisierung ist. $\square$

Im Folgenden sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen, identisch verteilten Zufallsvariablen mit $X_i \sim$ Dichte g_θ , $i = 1, \dots, n$ und

$$(X_1, \dots, X_n) \sim \text{Dichte } f_\theta(x) = \prod_{i=1}^n g_\theta(x_i)$$

aus der Klasse der Verteilungen mit monotonen Dichtekoeffizienten und einer Statistik $T(X_1, \dots, X_n)$.

Wir betrachten die Hypothesen $H_0 : \theta \leq \theta_0$ vs. $H_1 : \theta > \theta_0$ und den Neyman-Pearson-Test:

$$\varphi_{K^*}^*(x) = \begin{cases} 1, & \text{falls } T(x) > K^*, \\ \gamma^*, & \text{falls } T(x) = K^*, \\ 0, & \text{falls } T(x) < K^* \end{cases} \quad (5.3.6)$$

für $K^* \in \mathbb{R}$ und $\gamma^* \in [0, 1]$. Die Gütefunktion von $\varphi_{K^*}^*$ bei θ_0 ist

$$G_n(\theta_0) = \mathbb{E}_0 \varphi_{K^*}^* = \mathbb{P}_0(T(X_1, \dots, X_n) > K^*) + \gamma^* \cdot \mathbb{P}_0(T(X_1, \dots, X_n) = K^*)$$

- Satz 5.3.3** 1. Falls $\alpha = \mathbb{E}_0 \varphi_{K^*}^* > 0$, dann ist der soeben definierte Neyman-Pearson-Test ein bester Test der einseitigen Hypothesen H_0 vs. H_1 zum Niveau α .
2. Zu jedem Konfidenzniveau $\alpha \in (0, 1)$ gibt es ein $K^* \in \mathbb{R}$ und $\gamma^* \in [0, 1]$, sodaß $\varphi_{K^*}^*$ ein bester Test zum Umfang α ist.
3. Die Gütefunktion $G_n(\theta)$ von $\varphi_{K^*}^*(\theta)$ ist monoton wachsend in θ . Falls $0 < G_n(\theta) < 1$, dann ist sie sogar streng monoton wachsend.

Beweis 1. Wähle $\theta_1 > \theta_0$ und betrachte die einfachen Hypothesen $H'_0 : \theta = \theta_0$ und $H'_1 : \theta = \theta_1$. Sei

$$\varphi_K(x) = \begin{cases} 1, & f_1(x) > K f_0(x), \\ \gamma, & f_1(x) = K f_0(x), \\ 0, & f_1(x) < K f_0(x) \end{cases}$$

der Neyman-Pearson-Test für H'_0, H'_1 mit $K > 0$. Da f_θ den monotonen Dichtekoeffizienten mit Statistik T besitzt,

$$\frac{f_1(x)}{f_0(x)} = h(T(x), \theta_0, \theta_1),$$

existiert ein $K > 0$, so dass

$$\left\{ x : \frac{f_1(x)}{f_0(x)} \geq K \right\} \subset \left\{ T(x) \geq K^* \right\} \quad \text{mit } K = h(K^*, \theta_0, \theta_1).$$

φ_K ist ein bester Neyman-Pearson-Test zum Niveau $\alpha = \mathbb{E}_0 \varphi_K = \mathbb{E}_0 \varphi_{K^*}^*$. Aus $\alpha > 0$ folgt $K < \infty$, denn aus $K = \infty$ würde folgen

$$\begin{aligned} 0 < \alpha = \mathbb{E}_0 \varphi_K &\leq \mathbb{P}_0(T(X_1, \dots, X_n) \geq K^*) \leq \mathbb{P}_0\left(\frac{f_1(X_1, \dots, X_n)}{f_0(X_1, \dots, X_n)} = \infty\right) \\ &= \mathbb{P}_0(f_1(X_1, \dots, X_n) > 0, f_0(X_1, \dots, X_n) = 0) \\ &= \int_B \mathbb{I}(f_1(x) > 0, f_0(x) = 0) \cdot f_0(x) \mu(dx) = 0. \end{aligned}$$

Für den Test $\varphi_{K^*}^*$ aus (5.3.6) gilt dann

$$\varphi^*(x) = \begin{cases} 1, & \text{falls } f_1(x)/f_0(x) > K, \\ \gamma^*(x), & \text{falls } f_1(x)/f_0(x) = K, \\ 0, & \text{falls } f_1(x)/f_0(x) < K, \end{cases}$$

wobei $\gamma^*(x) \in \{\gamma^*, 0, 1\}$. Daraus folgt, daß $\varphi_{K^*}^*$ ein bester Neyman-Pearson-Test ist für H'_0 vs. H'_1 (vergleiche Bemerkung 5.3.2, 1.) und Bemerkung 5.3.3 für beliebige $\theta_1 > \theta_0$. Deshalb ist $\varphi_{K^*}^*$ ein bester Neyman-Pearson-Test für $H''_0 : \theta = \theta_0$ vs. $H''_1 : \theta > \theta_0$ ist.

Die selbe Behauptung erhalten wir aus dem Teil 3. des Satzes für $H_0 : \theta \leq \theta_0$ vs. $H_1 : \theta > \theta_0$, weil dann $G_n(\theta) \leq G_n(\theta_0) = \alpha$ für alle $\theta < \theta_0$.

2. Siehe Beweis zu Satz 5.3.2, 1.).

3. Wir müssen zeigen, daß $G_n(\theta)$ monoton ist. Dazu wählen wir $\theta_1 < \theta_2$ und zeigen, daß $\alpha_1 = G_n(\theta_1) \leq G_n(\theta_2)$. Wir betrachten die neuen, einfachen Hypothesen $H_0'' : \theta = \theta_1$ vs. $H_1'' : \theta = \theta_2$. Der Test $\varphi_{K^*}^*$ kann genauso wie in 1. als Neyman-Pearson-Test dargestellt werden (für die Hypothesen H_0'' und H_1''), der ein bester Test zum Niveau α_1 ist. Betrachten wir einen weiteren konstanten Test $\varphi(x) = \alpha_1$. Dann ist $\alpha_1 = \mathbb{E}_{\theta_2} \varphi \leq \mathbb{E}_{\theta_2} \varphi_{K^*}^* = G_n(\theta_2)$. Daraus folgt, daß $G_n(\theta_1) \leq G_n(\theta_2)$.

Nun zeigen wir, daß für $G_n(\theta) \in (0, 1)$ gilt: $G_n(\theta_1) < G_n(\theta_2)$. Wir nehmen an, daß $\alpha_1 = G_n(\theta_1) = G_n(\theta_2)$ und $\theta_1 < \theta_2$ für $\alpha \in (0, 1)$. Es folgt, daß $\varphi(x) = \alpha_1$ auch ein bester Test für H_0'' und H_1'' ist. Aus Satz 5.3.2, 2.) folgt

$$\mu(\{x \in B : \underbrace{\varphi(x)}_{=\alpha_1} \neq \varphi_{K^*}^*(x)\}) = 0 \text{ auf } K_0 \cup K_1 = \{f_1(x) \neq Kf_0(x)\},$$

was ein Widerspruch zur Bauart des Tests $\varphi_{K^*}^*$ ist, der auf $K_0 \cup K_1$ nicht gleich $\alpha_1 \in (0, 1)$ sein kann. □

Bemerkung 5.3.5. 1. Der Satz 5.3.3 ist genauso auf Neyman-Pearson-Tests der einseitigen Hypothesen

$$H_0 : \theta \geq \theta_0 \text{ vs. } H_1 : \theta < \theta_0$$

anwendbar, mit dem entsprechenden Unterschied

$$\begin{aligned} \theta &\mapsto -\theta \\ T &\mapsto -T \end{aligned}$$

Somit existiert der beste α -Test auch in diesem Fall.

2. Man kann zeigen, daß die Gütefunktion $G_n(\varphi_{K^*}^*, \theta)$ des besten Neyman-Pearson-Tests auf $\Theta_0 = (-\infty, \theta_0)$ folgende Minimalitätseigenschaft besitzt:

$$G_n(\varphi_{K^*}^*, \theta) \leq G_n(\varphi, \theta) \quad \forall \varphi \in \Psi(\alpha), \theta \leq \theta_0$$

Beispiel 5.3.3

Wir betrachten eine normalverteilte Stichprobe $(X_1, \dots, X_n)$ von unabhängigen und identisch verteilten Zufallsvariablen X_i , wobei $X_i \sim N(\mu, \sigma_0^2)$ und σ_0^2 sei bekannt. Es werden die Hypothesen

$$H_0 : \mu \leq \mu_0 \text{ vs. } H_1 : \mu > \mu_0,$$

getestet. Aus Beispiel 5.1.2 kennen wir die Testgröße

$$T(X_1, \dots, X_n) = \sqrt{n} \frac{\bar{X}_n - \mu_0}{\sigma_0},$$

wobei unter H_0 gilt: $T(X_1, \dots, X_n) \sim N(0, 1)$. H_0 wird verworfen, falls

$$T(X_1, \dots, X_n) > z_{1-\alpha}, \quad \text{wobei } \alpha \in (0, 1).$$

Wir zeigen jetzt, daß dieser Test der beste Neyman-Pearson-Test zum Niveau α ist. Aus Beispiel 5.3.2 ist bekannt, daß die Dichte f_n von $(X_1, \dots, X_n)$ aus der einparametrischen Exponentialklasse ist, mit

$$\tilde{T}(X_1, \dots, X_n) = \sum_{i=1}^n X_i.$$

Dann gehört f_μ von $(x_1, \dots, x_n)$ zur einparametrischen Exponentialklasse auch bezüglich der Statistik

$$T(X_1, \dots, X_n) = \sqrt{n} \frac{\bar{X}_n - \mu}{\sigma_0}$$

Es gilt nämlich

$$\begin{aligned} f_\mu(x) &= \exp \left(\underbrace{\frac{\mu}{\sigma_0^2}}_{\tilde{c}(\mu)} \cdot \underbrace{\sum_{i=1}^n x_i}_{\tilde{T}} - \underbrace{\frac{\mu^2 n}{2\sigma_0^2}}_{\tilde{a}(\mu)} \right) \cdot l(x) \\ &= \exp \left(\underbrace{\frac{\mu \sqrt{n}}{\sigma_0}}_{c(\mu)} \cdot \underbrace{\sqrt{n} \frac{\bar{x}_n - \mu}{\sigma_0}}_T + \underbrace{\frac{\mu^2 n}{2\sigma_0^2}}_{a(\mu)} \right) \cdot l(x). \end{aligned}$$

Die Statistik T kann also in der Konstruktion des Neyman-Pearson-Tests (Gleichung (5.3.6)) verwendet werden:

$$\varphi_{K^*}(x) = \begin{cases} 1, & \text{falls } T(x) > z_{1-\alpha}, \\ 0, & \text{falls } T(x) = z_{1-\alpha}, \\ 0, & \text{falls } T(x) < z_{1-\alpha} \end{cases}$$

(mit $K^* = z_{1-\alpha}$ und $\gamma^* = 0$). Nach Satz 5.3.3 ist dieser Test der beste Neyman-Pearson-Test zum Niveau α für unsere Hypothesen:

$$\begin{aligned} G_n(\varphi_{K^*}, \mu_0) &= \mathbb{P}_0(T(X_1, \dots, X_n) > z_{1-\alpha}) + 0 \cdot \mathbb{P}_0(T(X_1, \dots, X_n) \leq z_{1-\alpha}) \\ &= 1 - \Phi(z_{1-\alpha}) = 1 - (1 - \alpha) = \alpha. \end{aligned}$$

5.3.4 Unverfälschte zweiseitige Tests

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen mit der Dichte

$$f_\theta(x) = \prod_{i=1}^n g_\theta(x_i).$$

Es wird ein zweiseitiger Test der Hypothesen

$$H_0 : \theta = \theta_0 \text{ vs. } H_1 : \theta \neq \theta_0$$

betrachtet. Für alle $\alpha \in (0, 1)$ kann es jedoch keinen besten Test φ zum Niveau α für H_0 vs. H_1 geben. Denn, nehmen wir an, φ wäre der beste Test zum Niveau α für H_0 vs. H_1 , dann wäre φ der beste Test für die Hypothesen

1. $H'_0 : \theta = \theta_0$ vs. $H'_1 : \theta > \theta_0$
2. $H''_0 : \theta = \theta_0$ vs. $H''_1 : \theta < \theta_0$.

Dann ist nach Satz 5.3.3, 3. die Gütefunktion

1. $G_n(\varphi, \theta) < \alpha$ auf $\theta < \theta_0$, bzw.
2. $G_n(\varphi, \theta) > \alpha$ auf $\theta < \theta_0$,

was ein Widerspruch ist!

Darum werden wir die Klasse aller möglichen Tests auf unverfälschte Tests (Definition 5.1.5) eingrenzen. Der Test φ ist unverfälscht genau dann, wenn

$$\begin{aligned} G_n(\varphi, \theta) &\leq \alpha \text{ für } \theta \in \Theta_0 \\ G_n(\varphi, \theta) &\geq \alpha \text{ für } \theta \in \Theta_1 \end{aligned}$$

Beispiel 5.3.4 1. $\varphi(x) \equiv \alpha$ ist unverfälscht.

2. Der zweiseitige Gauß-Test ist unverfälscht, vergleiche Beispiel 5.1.2: $G_n(\varphi, \mu) \geq \alpha$ für alle $\mu \in \mathbb{R}$.

Im Folgenden seien X_i unabhängig und identisch verteilt. Die Dichte f_θ von $(X_1, \dots, X_n)$ gehöre zur einparametrischen Exponentialklasse:

$$f_\theta(x) = \exp \{c(\theta) \cdot T(x) + a(\theta)\} \cdot l(x), \quad (5.3.7)$$

wobei $c(\theta)$ und $a(\theta)$ stetig differenzierbar auf Θ sein sollen, mit

$$c'(\theta) > 0 \quad \text{und} \quad \text{Var}_\theta T(X_1, \dots, X_n) > 0$$

für alle $\theta \in \Theta$. Sei $f_\Phi(x)$ stetig in (x, Θ) auf $B \times \Theta$.

Übungsaufgabe 5.3.1

Zeigen Sie, daß folgende Relation gilt:

$$a'(\theta) = -c'(\theta) \mathbb{E}_\theta T(X_1, \dots, X_n).$$

Lemma 5.3.3

Es sei φ ein unverfälschter Test zum Niveau α für

$$H_0 : \theta = \theta_0 \text{ vs. } H_1 : \theta \neq \theta_0.$$

Dann gilt:

1. $\alpha = \mathbb{E}_0 \varphi(X_1, \dots, X_n) = G_n(\varphi, \theta_0)$
2. $\mathbb{E}_0 [T(X_1, \dots, X_n) \varphi(X_1, \dots, X_n)] = \alpha \cdot \mathbb{E}_0 T(X_1, \dots, X_n)$

Beweis 1. Die Gütefunktion von φ ist

$$G_n(\varphi, \theta) = \int_B \varphi(x) f_\theta(x) \mu(dx)$$

Da f_θ aus der einparametrischen Exponentialklasse ist, ist $G_n(\varphi, \theta)$ differenzierbar (unter dem Integral) bezüglich θ . Wegen der Unverfälschtheit von φ gilt

$$G_n(\varphi, \theta_0) \leq \alpha, \quad G_n(\varphi, \theta) \geq \alpha, \quad \theta \neq \theta_0$$

und daraus folgt $G_n(\varphi, \theta_0) = \alpha$ und θ_0 ist ein Minimumpunkt von G_n . Somit ist 1) bewiesen.

2. Da θ_0 der Minimumpunkt von G_n ist, gilt

$$\begin{aligned} 0 &= G'_n(\varphi, \theta_0) = \int_B \varphi(x) (c'(\theta_0)T(x) + a'(\theta_0)) f_0(x) \mu(dx) \\ &= c'(\theta_0) \cdot \mathbb{E}_0 [\varphi(X_1, \dots, X_n) T(X_1, \dots, X_n)] + a'(\theta_0) \cdot G_n(\varphi, \theta_0) \\ &= c'(\theta_0) \cdot \mathbb{E}_0 [\varphi(X_1, \dots, X_n) T(X_1, \dots, X_n)] + \alpha a'(\theta_0) \\ &\stackrel{\text{(Übung 5.3.1)}}{=} c'(\theta_0) (\mathbb{E}_0 (\varphi \cdot T) - \alpha \mathbb{E}_0 T) \end{aligned}$$

Daraus folgt $\mathbb{E}_0 (\varphi T) = \alpha \mathbb{E}_0 T$ und damit ist das Lemma bewiesen. $\square$

Wir definieren jetzt die modifizierten Neyman-Pearson-Tests für einfache Hypothesen

$$H_0 : \theta = \theta_0 \text{ vs. } H'_1 : \theta = \theta_1, \quad \theta_1 \neq \theta_0.$$

Für $\lambda, K \in \mathbb{R}$, $\gamma : B \rightarrow [0, 1]$ definieren wir

$$\varphi_{K,\lambda}(x) = \begin{cases} 1, & \text{falls } f_1(x) > (K + \lambda T(x)) f_0(x), \\ \gamma(x), & \text{falls } f_1(x) = (K + \lambda T(x)) f_0(x), \\ 0, & \text{falls } f_1(x) < (K + \lambda T(x)) f_0(x), \end{cases} \quad (5.3.8)$$

wobei $T(x)$ die Statistik aus der Darstellung (5.3.7) ist.

Es sei $\tilde{\Psi}(\alpha)$ die Klasse aller Tests, die Aussagen 1) und 2) des Lemmas 5.3.3 erfüllen. Aus Lemma 5.3.3 folgt dann, daß die Menge der unverfälschten Tests zum Niveau α eine Teilmenge von $\tilde{\Psi}(\alpha)$ ist.

Satz 5.3.4

Der modifizierte Neyman-Pearson-Test $\varphi_{K,\lambda}$ ist der beste α -Test in $\tilde{\Psi}(\alpha)$ für Hypothesen H_0 vs. H'_1 zum Niveau $\alpha = \mathbb{E}_0 \varphi_{K,\lambda}$, falls $\varphi_{K,\lambda} \in \tilde{\Psi}(\alpha)$.

Beweis Es ist zu zeigen, daß $\mathbb{E}_1 \varphi_{K,\lambda} \geq \mathbb{E}_1 \varphi$ für alle $\varphi \in \tilde{\Psi}(\alpha)$, bzw. $\mathbb{E}_1 (\varphi_{K,\lambda} - \varphi) \geq 0$. Es gilt

$$\begin{aligned} \mathbb{E}_1 (\varphi_{K,\lambda} - \varphi) &= \int_B (\varphi_{K,\lambda}(x) - \varphi(x)) f_1(x) \mu(dx) \\ &\stackrel{(\text{Bem. 5.3.3, 2.})}{\geq} \int_B (\varphi_{K,\lambda}(x) - \varphi(x)) (K + \lambda T(x)) f_0(x) \mu(dx) \\ &= K \left(\underbrace{\mathbb{E}_0 \varphi_{K,\lambda}}_{=\alpha} - \underbrace{\mathbb{E}_0 \varphi}_{=\alpha} \right) + \lambda \left(\underbrace{\mathbb{E}_0 (\varphi_{K,\lambda} \cdot T)}_{=\alpha \mathbb{E}_0 T} - \underbrace{\mathbb{E}_0 (\varphi \cdot T)}_{=\alpha \mathbb{E}_0 T} \right) \\ &= 0, \end{aligned}$$

weil $\varphi, \varphi_{K,\lambda} \in \tilde{\Psi}(\alpha)$. $\square$

Wir definieren folgende Entscheidungsregel, die später zum Testen der zweiseitigen Hypothesen

$$H_0 : \theta = \theta_0 \text{ vs. } H_1 : \theta \neq \theta_0$$

verwendet wird:

$$\varphi_c(x) = \begin{cases} 1, & \text{falls } T(x) \notin (c_1, c_2), \\ \gamma_1, & \text{falls } T(x) = c_1, \\ \gamma_2, & \text{falls } T(x) = c_2, \\ 0, & \text{falls } T(x) \in (c_1, c_2), \end{cases} \quad (5.3.9)$$

für $c_1 \leq c_2 \in \mathbb{R}$, $\gamma_1, \gamma_2 \in [0, 1]$ und die Statistik $T(x)$, $x = (x_1, \dots, x_n) \in B$, die in der Dichte (5.3.7) vorkommt. Zeigen wir, daß φ_c sich als modifizierter Neyman-Pearson-Test schreiben lässt.

Für die Dichte

$$f_\theta(x) = \exp\{c(\theta)T(x) + a(\theta)\} \cdot l(x)$$

wird (wie immer) vorausgesetzt, daß $l(x) > 0$, $c'(\theta) > 0$ und $a'(\theta)$ existiert für $\theta \in \Theta$.

Lemma 5.3.4

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen, identisch verteilten Zufallsvariablen mit gemeinsamer Dichte $f_\theta(x)$, $x \in B$, die zur einparametrischen Exponentialfamilie gehört. Sei $T(x)$ die dazugehörige Statistik, die im Exponenten der Dichte f_θ vorkommt. Für beliebige reelle Zahlen $c_1 \leq c_2$, $\gamma_1, \gamma_2 \in [0, 1]$ und Parameterwerte $\theta_0, \theta_1 \in \Theta : \theta_0 \neq \theta_1$ läßt sich der Test φ_c aus (5.3.9) als modifizierter Neyman-Pearson-Test $\varphi_{K,\lambda}$ aus (5.3.8) mit gegebenen $K, \lambda \in \mathbb{R}$, $\gamma(x) \in [0, 1]$ schreiben.

Beweis Falls wir die Bezeichnung

$$f_{\theta_i}(x) = f_i(x), \quad i = 0, 1$$

verwenden, dann gilt

$$\frac{f_1(x)}{f_0(x)} = \exp \left\{ \underbrace{(c(\theta_1) - c(\theta_0))}_{c} T(x) + \underbrace{a(\theta_1) - a(\theta_0)}_a \right\},$$

und somit

$$\{x \in B : f_1(x) > (K + \lambda T(x)) f_0(x)\} = \{x \in B : \exp(cT(x) + a) > K + \lambda T(x)\}.$$

Finden wir solche K und λ aus $\mathbb{R}$, für die die Gerade $K + \lambda t$, $t \in \mathbb{R}$ die konvexe Kurve $\exp\{ct + a\}$ genau an den Stellen c_1 und c_2 schneidet (falls $c_1 \neq c_2$) bzw. an der Stelle $t = c_1$ berührt (falls $c_1 = c_2$). Dies ist immer möglich, siehe Abbildung 5.6.

Ferner setzen wir $\gamma(x) = \gamma_i$ für $\{x \in B : T(x) = c_i\}$. Insgesamt gilt dann

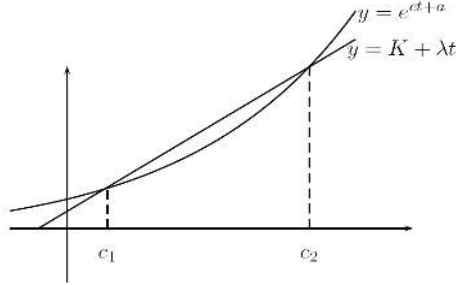
$$\{x : \exp(cT(x) + a) > K + \lambda T(x)\} = \{x : T(x) \notin [c_1, c_2]\}$$

und

$$\{x : \exp(cT(x) + a) < K + \lambda T(x)\} = \{x : T(x) \in (c_1, c_2)\}.$$

Damit ist das Lemma bewiesen. □

Abb. 5.6:



Bemerkung 5.3.6. 1. Die Umkehrung des Lemmas stimmt nicht, denn bei vorgegebenen Kurven $y = K + \lambda t$ und $y = \exp\{ct + a\}$ muss es die Schnittpunkte c_1 und c_2 nicht unbedingt geben. So kann die Gerade vollständig unter der Kurve $y = \exp\{ct + a\}$ liegen.

2. Der Test φ_c macht von den Werten θ_0 und θ_1 nicht explizit Gebrauch. Dies unterscheidet ihn vom Test $\varphi_{K,\lambda}$, für den die Dichten f_0 und f_1 gebraucht werden.

Jetzt sind wir bereit, den Hauptsatz über zweiseitige Tests zum Prüfen der Hypothesen

$$H_0 : \theta = \theta_0 \text{ vs. } H_1 : \theta \neq \theta_0$$

zu formulieren und zu beweisen.

Satz 5.3.5

(Hauptsatz über zweiseitige Tests)

Unter den Voraussetzungen des Lemmas 5.3.4 sei φ_c ein Test aus (5.3.9), für den $\varphi_c \in \tilde{\Psi}(\alpha)$ gilt. Dann ist φ_c bester unverfälschter Test zum Niveau α (und dadurch bester Test in $\tilde{\Psi}(\alpha)$) der Hypothesen

$$H_0 : \theta = \theta_0 \text{ vs. } H_1 : \theta \neq \theta_0.$$

Beweis Wählen wir ein beliebiges $\theta_1 \in \Theta$, $\theta_1 \neq \theta_0$. Nach Lemma 5.3.4 ist φ_c ein modifizierter Neyman-Pearson-Test $\varphi_{K,\lambda}$ für eine spezielle Wahl von K und $\lambda \in \mathbb{R}$. $\varphi_{K,\lambda}$ ist aber nach Satz 5.3.4 bester Test in $\tilde{\Psi}(\alpha)$ für $H_0 : \theta = \theta_0$ vs. $H'_1 : \theta = \theta_1$. Da φ_c nicht von θ_1 abhängt, ist es bester Test in $\tilde{\Psi}(\alpha)$ für $H_1 : \theta \neq \theta_0$. Da unverfälschte Niveau- α -Tests in $\tilde{\Psi}(\alpha)$ liegen, müssen wir nur zeigen, daß φ_c unverfälscht ist. Da φ_c der beste Test ist, ist er nicht schlechter als der konstante unverfälschte Test $\varphi = \alpha$, das heißt

$$G_n(\varphi_c, \theta) \geq G_n(\varphi, \theta) = \alpha, \quad \theta \neq \theta_0.$$

Somit ist auch φ_c unverfälscht. Der Beweis ist beendet. $\square$

Bemerkung 5.3.7. Wir haben gezeigt, daß φ_c der beste Test seines Umfangs ist. Es wäre jedoch noch zu zeigen, daß für beliebiges $\alpha \in (0, 1)$ Konstanten $c_1, c_2, \gamma_1, \gamma_2$ gefunden werden, für die $\mathbb{E}_0 \varphi_c = \alpha$ gilt. Da der Beweis schwierig ist, wird er hier ausgelassen. Im folgenden Beispiel jedoch wird es klar, wie die Parameter $c_1, c_2, \gamma_1, \gamma_2$ zu wählen sind.

Beispiel 5.3.5 (Zweiseitiger Gauß-Test):

Im Beispiel 5.1.2 haben wir folgenden Test des Erwartungswertes einer normalverteilten Stichprobe $(X_1, \dots, X_n)$ mit unabhängigen und identisch verteilten X_i und $X_i \sim N(\mu, \sigma_0^2)$ bei bekannten Varianzen σ_0^2 betrachtet. Getestet werden die Hypothesen

$$H_0 : \mu = \mu_0 \text{ vs. } H_1 : \mu \neq \mu_0.$$

Der Test $\varphi(x)$ lautet

$$\varphi(x) = \mathbb{I} \left(x \in \mathbb{R}^n : |T(x)| > z_{1-\alpha/2} \right),$$

wobei

$$T(x) = \sqrt{n} \frac{\bar{x}_n - \mu_0}{\sigma_0}.$$

Zeigen wir, daß φ der beste Test zum Niveau α in $\tilde{\Psi}(\alpha)$ (und somit bester unverfälschter Test) ist. Nach Satz 5.3.5 müssen wir lediglich prüfen, daß φ als φ_c mit (5.3.9) dargestellt werden kann, weil die n -dimensionale Normalverteilung mit Dichte f_μ (siehe Beispiel 5.3.3) zu der einparametrischen Exponentialfamilie mit Statistik

$$T(x) = \sqrt{n} \frac{\bar{x}_n - \mu}{\sigma_0}$$

gehört. Setzen wir $c_1 = z_{1-\alpha/2}$, $c_2 = -z_{1-\alpha/2}$, $\gamma_1 = \gamma_2 = 0$. Damit ist

$$\varphi(x) = \varphi_c(x) = \begin{cases} 1, & \text{falls } |T(x)| > z_{1-\alpha/2}, \\ 0, & \text{falls } |T(x)| \leq z_{1-\alpha/2}. \end{cases}$$

und die Behauptung ist bewiesen, weil aus der in Beispiel 5.1.2 ermittelten Gütefunktion $G_n(\varphi, \theta)$ von φ ersichtlich ist, daß φ ein unverfälschter Test zum Niveau α ist (und somit $\varphi \in \tilde{\Psi}(\alpha)$).

Bemerkung 5.3.8. Bisher haben wir immer vorausgesetzt, daß nur ein Parameter der Verteilung der Stichprobe $(X_1, \dots, X_n)$ unbekannt ist, um die Theorie des Abschnittes 1.3 über die besten (Neyman-Pearson-) Tests im Fall der einparametrischen Exponentialfamilie aufstellen zu können. Um jedoch den Fall weiterer unbekannten Parameter betrachten zu können (wie im Beispiel der zweiseitigen Tests des Erwartungswertes der normalverteilten Stichprobe bei unbekannter Varianz (der sog. t-Test, vergleiche Abschnitt 1.2.1, 1 (a)), bedarf es einer tiefergehenden Theorie, die aus Zeitgründen in dieser Vorlesung nicht behandelt wird. Der interessierte Leser findet das Material dann im Buch [?].

5.4 Anpassungstests

Sei eine Stichprobe von unabhängigen, identisch verteilten Zufallsvariablen $(X_1, \dots, X_n)$ gegeben mit $X_i \sim F$ (Verteilungsfunktion) für $i = 1, \dots, n$. Bei den Anpassungstests wird die Hypothese

$$H_0 : F = F_0 \text{ vs. } H_1 : F \neq F_0$$

überprüft, wobei F_0 eine vorgegebene Verteilungsfunktion ist.

Einen Test aus dieser Klasse haben wir bereits in der Vorlesung Stochastik I kennengelernt: den Kolmogorow-Smirnov-Test (vergleiche Bemerkung 3.3.8. 3), Vorlesungsskript Stochastik I).

Jetzt werden weitere nichtparametrische Anpassungstests eingeführt. Der erste ist der χ^2 -Anpassungstest von K. Pearson.

5.4.1 χ^2 -Anpassungstest

Der Test von Kolmogorow-Smirnov basiert auf dem Abstand

$$D_n = \sup_{x \in \mathbb{R}} |\hat{F}_n(x) - F_0(x)|$$

zwischen der empirischen Verteilungsfunktion der Stichprobe $(X_1, \dots, X_n)$ und der Verteilungsfunktion F_0 . In der Praxis jedoch erscheint dieser Test zu feinfühlig, denn er ist zu sensibel gegenüber Unregelmäßigkeiten in den Stichproben und verwirft H_0 zu oft. Einen Ausweg aus dieser Situation stellt die Vergrößerung der Haupthypothese H_0 dar, auf welcher der folgende χ^2 -Anpassungstest beruht.

Man zerlegt den Wertebereich der Stichprobenvariablen X_i in r Klassen $(a_j, b_j]$, $j = 1, \dots, r$ mit der Eigenschaft

$$-\infty \leq a_1 < b_1 = a_2 < b_2 = \dots = a_r < b_r \leq \infty.$$

Anstelle von $X_i, i = 1, \dots, n$ betrachten wir die sogenannten *Klassenstärken* $Z_j, j = 1, \dots, r$, wobei

$$Z_j = \#\{i : a_j < X_i \leq b_j, 1 \leq i \leq n\}.$$

Lemma 5.4.1

Der Zufallsvektor $Z = (Z_1, \dots, Z_r)^\top$ ist *multinomialverteilt* mit Parametervektor

$$p = (p_1, \dots, p_{r-1})^\top \in [0, 1]^{r-1},$$

wobei

$$p_j = \mathbb{P}(a_j < X_1 \leq b_j) = F(b_j) - F(a_j), \quad j = 1, \dots, r-1, \quad p_r = 1 - \sum_{j=1}^{r-1} p_j.$$

Schreibweise:

$$Z \sim M_{r-1}(n, p)$$

Beweis Es ist zu zeigen, daß für alle Zahlen $k_1, \dots, k_r \in \mathbb{N}_0$ mit $k_1 + \dots + k_r = n$ gilt:

$$\mathbb{P}(Z_i = k_i, i = 1, \dots, r) = \frac{n!}{k_1! \cdot \dots \cdot k_r!} p_1^{k_1} \cdot \dots \cdot p_r^{k_r}. \quad (5.4.1)$$

Da X_i unabhängig und identisch verteilt sind, gilt

$$\mathbb{P}(X_j \in (a_{i_j}, b_{i_j}], j = 1, \dots, n) = \prod_{j=1}^n \mathbb{P}(a_{i_j} < X_1 \leq b_{i_j}) = p_1^{k_1} \cdot \dots \cdot p_r^{k_r},$$

falls die Folge von Intervallen $(a_{i_j}, b_{i_j}]_{j=1, \dots, n}$ das Intervall $(a_i, b_i]$ k_i Mal enthält, $i = 1, \dots, r$. Die Formel (5.4.1) ergibt sich aus dem Satz der totalen Wahrscheinlichkeit als Summe über die Permutationen von Folgen $(a_{i_j}, b_{i_j}]_{j=1, \dots, n}$ dieser Art. $\square$

Im Sinne des Lemmas 5.4.1 werden neue Hypothesen über die Beschaffenheit von F geprüft.

$$H_0 : p = p_0 \text{ vs. } H_1 : p \neq p_0,$$

wobei $p = (p_1, \dots, p_{r-1})^\top$ der Parametervektor der Multinomialverteilung von Z ist, und $p_0 = (p_{01}, \dots, p_{0,r-1})^\top \in (0, 1)^{r-1}$ mit $\sum_{i=1}^{r-1} p_{0i} < 1$. In diesem Fall ist

$$\Lambda_0 = \{F \in \Lambda : F(b_j) - F(a_j) = p_{0j}, \quad j = 1, \dots, r-1\},$$

$\Lambda_1 = \Lambda \setminus \Lambda_0$, wobei Λ die Menge aller Verteilungsfunktionen ist. Um H_0 vs. H_1 zu testen, führen wir die *Pearson-Teststatistik*

$$T_n(x) = \sum_{j=1}^r \frac{(z_j - np_{0j})^2}{np_{0j}}$$

ein, wobei $x = (x_1, \dots, x_n)$ eine konkrete Stichprobe der Daten ist und $z_j, j = 1, \dots, r$ ihre Klassenstärken sind.

Unter H_0 gilt

$$\mathbb{E} Z_j = np_{0j}, \quad j = 1, \dots, r,$$

somit soll H_0 abgelehnt werden, falls $T_n(X)$ ungewöhnlich große Werte annimmt.

Im nächsten Satz zeigen wir, daß $T(X_1, \dots, X_n)$ asymptotisch (für $n \rightarrow \infty$) χ_{r-1}^2 -verteilt ist, was zu folgendem Anpassungstest (χ^2 -Anpassungstest) führt:

$$H_0 \text{ wird verworfen, falls } T_n(x_1, \dots, x_n) > \chi_{r-1, 1-\alpha}^2.$$

Dieser Test ist nach seinem Entdecker *Karl Pearson* (1857-1936) benannt worden.

Satz 5.4.1

Unter H_0 gilt

$$\lim_{n \rightarrow \infty} \mathbb{P}_{p_0} (T_n(X_1, \dots, X_n) > \chi_{r-1, 1-\alpha}^2) = \alpha, \quad \alpha \in (0, 1),$$

das heißt, der χ^2 -Pearson-Test ist ein asymptotischer Test zum Niveau α .

Beweis Führen wir die Bezeichnung $Z_{nj} = Z_j(X_1, \dots, X_n)$ der Klassenstärken ein, die aus der Stichprobe $(X_1, \dots, X_n)$ entstehen. Nach Lemma 5.4.1 ist

$$Z_n = (Z_{n1}, \dots, Z_{nr}) \sim M_{r-1}(n, p_0) \text{ unter } H_0.$$

Insbesondere soll $\mathbb{E} Z_{nj} = np_{0j}$ und

$$\text{Cov}(Z_{ni}, Z_{nj}) = \begin{cases} np_{0j}(1 - p_{0j}), & i = j, \\ -np_{0i}p_{0j}, & i \neq j \end{cases}$$

für alle $i, j = 1, \dots, r$ gelten. Da

$$Z_{nj} = \sum_{i=1}^n \mathbb{I}(a_j < X_i \leq b_j), \quad j = 1, \dots, r,$$

ist $Z_n = (Z_{n1}, \dots, Z_{n,r-1})$ eine Summe von n unabhängigen und identisch verteilten Zufallsvektoren $Y_i \in \mathbb{R}^{r-1}$ mit Koordinaten $Y_{ij} = \mathbb{I}(a_j < X_i \leq b_j)$, $j = 1, \dots, r-1$. Daher gilt nach dem multivariaten Grenzwertsatz (der in Lemma 5.4.2 bewiesen wird), daß

$$Z'_n = \frac{Z_n - \mathbb{E} Z_n}{\sqrt{n}} = \frac{\sum_{i=1}^n Y_i - n\mathbb{E} Y_1}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, K),$$

mit $N(0, K)$ eine $(r-1)$ -dimensionale multivariate Normalverteilung (vergleiche Vorlesungsskript WR, Beispiel 3.4.1. 3.) mit Erwartungswertvektor Null und Kovarianzmatrix $K = (\sigma_{ij}^2)$, wobei

$$\sigma_{ij}^2 = \begin{cases} -p_{0i}p_{0j}, & i \neq j, \\ p_{0i}(1 - p_{0j}), & i = j \end{cases}$$

für $i, j = 1, \dots, r-1$ ist. Diese Matrix K ist invertierbar mit $K^{-1} = A = (a_{ij})$,

$$a_{ij} = \begin{cases} \frac{1}{p_{0r}}, & i \neq j, \\ \frac{1}{p_{0i}} + \frac{1}{p_{0r}}, & i = j. \end{cases}$$

Außerdem ist K (als Kovarianzmatrix) symmetrisch und positiv definit. Aus der linearen Algebra ist bekannt, daß es eine invertierbare $(r-1) \times (r-1)$ -Matrix $A^{1/2}$ gibt, mit der Eigenschaft $A = A^{1/2}(A^{1/2})^\top$. Daraus folgt,

$$K = A^{-1} = ((A^{1/2})^\top)^{-1} \cdot (A^{1/2})^{-1}.$$

Wenn wir $(A^{1/2})^\top$ auf Z'_n anwenden, so bekommen wir

$$(A^{1/2})^\top \cdot Z'_n \xrightarrow[n \rightarrow \infty]{d} (A^{1/2})^\top \cdot Y,$$

wobei

$$(A^{1/2})^\top \cdot Y \sim N\left(0, (A^{1/2})^\top \cdot K \cdot A^{1/2}\right) = N(0, \mathcal{I}_{r-1})$$

nach der Eigenschaft der multivariaten Normalverteilung, die im Kapitel 2, Satz ?? behandelt wird. Des Weiteren wurde hier der Stetigkeitssatz aus der Wahrscheinlichkeitsrechnung benutzt, daß

$$Y_n \xrightarrow[n \rightarrow \infty]{d} Y \implies \varphi(Y_n) \xrightarrow[n \rightarrow \infty]{d} \varphi(Y)$$

für beliebige Zufallsvektoren $\{Y_n\}$, $Y \in \mathbb{R}^m$ und stetige Abbildungen $\varphi : \mathbb{R} \rightarrow \mathbb{R}$. Diesen Satz haben wir in WR für Zufallsvariablen bewiesen (Satz 6.4.3, Vorlesungsskript WR). Die erneute Anwendung des Stetigkeitssatzes ergibt

$$\left| (A^{1/2})^\top Z'_n \right|^2 \xrightarrow[n \rightarrow \infty]{d} |Y|^2 = R \sim \chi_{r-1}^2.$$

Zeigen wir, daß

$$T_n(X_1, \dots, X_n) = \left| (A^{1/2})^\top Z'_n \right|^2.$$

Es gilt:

$$\begin{aligned} \left| (A^{1/2})^\top Z'_n \right|^2 &= ((A^{1/2})^\top Z'_n)^\top ((A^{1/2})^\top Z'_n) = Z_n'^\top \cdot \underbrace{A^{1/2} \cdot (A^{1/2})^\top}_A Z'_n = Z_n'^\top A Z'_n \\ &= n \sum_{j=1}^{r-1} \frac{1}{p_{0j}} \left(\frac{Z_{nj}}{n} - p_{0j} \right)^2 + \frac{n}{p_{0r}} \sum_{i=1}^{r-1} \sum_{j=1}^{r-1} \left(\frac{Z_{ni}}{n} - p_{0i} \right) \left(\frac{Z_{nj}}{n} - p_{0j} \right) \\ &= \sum_{j=1}^{r-1} \frac{(Z_{nj} - np_{0j})^2}{np_{0j}} + \frac{n}{p_{0r}} \left(\sum_{j=1}^{r-1} \left(\frac{Z_{nj}}{n} - p_{0j} \right) \right)^2 \\ &= \sum_{j=1}^{r-1} \frac{(Z_{nj} - np_{0j})^2}{np_{0j}} + \frac{n}{p_{0r}} \left(\frac{Z_{nr}}{n} - p_{0r} \right)^2 \\ &= \sum_{j=1}^r \frac{(Z_{nj} - np_{0j})^2}{np_{0j}} = T_n(X_1, \dots, X_n), \end{aligned}$$

weil

$$\begin{aligned} \sum_{j=1}^{r-1} Z_{nj} &= n - Z_{nr}, \\ \sum_{j=1}^{r-1} p_{0j} &= 1 - p_{0r}. \end{aligned}$$

□

Lemma 5.4.2 (Multivariater zentraler Grenzwertsatz):

Sei $\{Y_n\}_{n \in \mathbb{N}}$ eine Folge von unabhängigen und identisch verteilten Zufallsvektoren, mit $\mathbb{E} Y_1 = \mu$ und Kovarianzmatrix K . Dann gilt

$$\frac{\sum_{i=1}^n Y_i - n\mu}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, K). \quad (5.4.2)$$

Beweis Sei $Y_j = (Y_{j1}, \dots, Y_{jm})^\top$. Nach dem Stetigkeitssatz für charakteristische Funktionen ist die Konvergenz (5.4.2) äquivalent zu

$$\varphi_n(t) \xrightarrow{n \rightarrow \infty} \varphi(t) \quad t \in \mathbb{R}^m, \quad (5.4.3)$$

wobei

$$\varphi_n(t) = \mathbb{E} e^{itS_n} = \mathbb{E} \exp \left\{ i \sum_{j=1}^m t_j \frac{Y_{1j} + \dots + Y_{nj} - n\mu_j}{\sqrt{n}} \right\}$$

die charakteristische Funktion vom Zufallsvektor

$$S_n = \frac{\sum_{i=1}^n Y_i - n\mu}{\sqrt{n}}$$

und

$$\varphi(t) = e^{-t^\top K t / 2}$$

die charakteristische Funktion der $N(0, K)$ -Verteilung ist. Die Funktion $\varphi_n(t)$ kann in der Form

$$\varphi_n(t) = \mathbb{E} \exp \left\{ i \sum_{i=1}^n \frac{\sum_{j=1}^m t_j (Y_{ij} - \mu_j)}{\sqrt{n}} \right\}, \quad t = (t_1, \dots, t_m)^\top \in \mathbb{R}^m$$

umgeschrieben werden, wobei für die Zufallsvariable

$$L_i := \sum_{j=1}^m t_j (Y_{ij} - \mu_j)$$

gilt:

$$\begin{aligned} \mathbb{E} L_i &= 0, \\ \text{Var } L_i &= \mathbb{E} \left[\sum_{k,j=1}^m t_j (Y_{ij} - \mu_j) (Y_{ik} - \mu_k) t_k \right] = t^\top K t, \quad i \in \mathbb{N}. \end{aligned}$$

Falls $t^\top K t = 0$, dann gilt $L_i = 0$ fast sicher, für alle $i \in \mathbb{N}$. Hieraus folgt $\varphi_n(t) = \varphi(t) = 1$, also gilt die Konvergenz 5.4.2.

Falls jedoch $t^\top K t > 0$, dann kann $\varphi_n(t)$ als charakteristische Funktion der Zufallsvariablen

$$\sum_{i=1}^n L_i / \sqrt{n}$$

an Stelle 1, und $\varphi(t)$ als charakteristische Funktion der eindimensionalen Normalverteilung $N(0, t^\top K t)$ an Stelle 1 interpretiert werden. Aus dem zentralen Grenzwertsatz für eindimensionale Zufallsvariablen (vergleiche Satz 7.2.1, Vorlesungsskript WR) gilt

$$\sum_{i=1}^n \frac{L_i}{\sqrt{n}} \xrightarrow[n \rightarrow \infty]{d} L \sim N(0, t^\top K t)$$

und somit

$$\varphi_n(t) = \varphi\left(\sum_{i=1}^n L_i/\sqrt{n}\right)(1) \xrightarrow{n \rightarrow \infty} \varphi_L(1) = \varphi(t).$$

Somit ist die Konvergenz (5.4.2) bewiesen. $\square$

Bemerkung 5.4.1. 1. Die im letzten Beweis verwendete Methode der Reduktion einer mehrdimensionalen Konvergenz auf den eindimensionalen Fall mit Hilfe von Linearkombinationen von Zufallsvariablen trägt den Namen von Cramér-Wold.

2. Der χ^2 -Pearson-Test ist asymptotisch, also für große Stichprobenumfänge, anzuwenden. Aber welches n ist groß genug? Als „Faustregel“ gilt: np_{0j} soll größer gleich a sein, $a \in (2, \infty)$. Für eine größere Klassenanzahl $r \geq 10$ kann sogar $a = 1$ verwendet werden. Wir zeigen jetzt, daß der χ^2 -Anpassungstest konsistent ist.

Lemma 5.4.3

Der χ^2 -Pearson-Test ist konsistent, das heißt

$$\forall p \in [0, 1]^{r-1}, p \neq p_0 \text{ gilt: } \lim_{n \rightarrow \infty} \mathbb{P}_p \left(T_n(X_1, \dots, X_n) > \chi_{r-1, 1-\alpha}^2 \right) = 1$$

Beweis Unter H_1 gilt

$$Z_{nj}/n = \frac{\sum_{i=1}^n \mathbb{I}(a_j < X_i \leq b_j)}{n} \xrightarrow[n \rightarrow \infty]{f.s.} \underbrace{\mathbb{E} \mathbb{I}(a_j < X_1 \leq b_j)}_{=p_j}$$

nach dem starken Gesetz der großen Zahlen. Wir wählen j so, daß $p_j \neq p_{0j}$. Es gilt

$$T_n(X_1, \dots, X_n) \geq \frac{(Z_{nj} - np_{0j})^2}{np_{0j}} \geq \underbrace{n \left(\frac{Z_{nj}}{n} - p_{0j} \right)^2}_{\sim n(p_j - p_{0j})^2} \xrightarrow[n \rightarrow \infty]{f.s.} \infty.$$

Somit ist auch

$$\mathbb{P}_p \left(T_n(X_1, \dots, X_n) > \chi_{r-1, 1-\alpha}^2 \right) \xrightarrow[n \rightarrow \infty]{f.s.} 1.$$

$\square$

5.4.2 χ^2 -Anpassungstest von Pearson-Fisher

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen $X_i, i = 1, \dots, n$. Wir wollen testen, ob die Verteilungsfunktion F von X_i zu einer parametrischen Familie

$$\Lambda_0 = \{F_\theta : \theta \in \Theta\}, \quad \Theta \subset \mathbb{R}^m$$

gehört. Seien die Zahlen $a_i, b_i, i = 1, \dots, r$ vorgegeben mit der Eigenschaft

$$-\infty \leq a_1 < b_1 = a_2 < b_2 = \dots = a_r < b_r \leq \infty$$

und

$$Z_j = \#\{X_i, i = 1, \dots, n : a_j < X_i \leq b_j\}, \quad j = 1, \dots, r,$$

$$Z = (Z_1, \dots, Z_r)^\top.$$

Nach Lemma 5.4.1 gilt: $Z \sim M_{r-1}(n, p)$, $p = (p_0, \dots, p_{r-1})^\top \in [0, 1]^{r-1}$. Unter der Hypothese $H_0 : F \in \Lambda_0$ gilt: $p = p(\theta)$, $\theta \in \Theta \subset \mathbb{R}^m$. Wir vergrößern die Hypothese H_0 und wollen folgende neue Hypothese testen:

$$H_0 : p \in \{p(\theta) : \theta \in \Theta\} \text{ vs. } H_1 : p \notin \{p(\theta) : \theta \in \Theta\}.$$

Um dieses Hypothesenpaar zu testen, wird der χ^2 -Pearson-Fisher-Test wie folgt aufgebaut:

1. Ein (schwach konsistenter) Maximum-Likelihood-Schätzer $\hat{\theta}_n = \hat{\theta}(X_1, \dots, X_n)$ für θ wird gefunden: $\hat{\theta}_n \xrightarrow[n \rightarrow \infty]{P} \theta$. Dabei muß $\{\hat{\theta}_n\}_{n \in \mathbb{N}}$ asymptotisch normalverteilt sein.
2. Es wird der Plug-In-Schätzer $p(\hat{\theta}_n)$ für $p(\theta)$ gebildet.
3. Die Testgröße

$$\hat{T}_n(X_1, \dots, X_n) = \sum_{j=1}^r \frac{(Z_{nj} - np_j(\hat{\theta}))^2}{np_j(\hat{\theta})} \xrightarrow[n \rightarrow \infty]{P} \eta \sim \chi_{r-m-1}^2$$

unter H_0 und gewissen Voraussetzungen.

4. H_0 wird verworfen, falls $\hat{T}_n(X_1, \dots, X_n) > \chi_{r-m-1, 1-\alpha}^2$. Dies ist ein asymptotischer Test zum Niveau α .

Bemerkung 5.4.2. 1. Bei einem χ^2 -Pearson-Fisher-Test wird vorausgesetzt, daß die Funktion $p(\theta)$ explizit bekannt ist, θ jedoch unbekannt. Das bedeutet, daß für jede Klasse von Verteilungen Λ_0 die Funktion $p(\cdot)$ berechnet werden soll.

2. Warum kann $\hat{T}_n$ die Hypothese H_0 von H_1 unterscheiden? Nach dem Gesetz der großen Zahlen gilt

$$\frac{1}{n} Z_{nj} - p_j(\hat{\theta}_n) = \underbrace{\frac{1}{n} Z_{nj} - p_j(\theta)}_{\xrightarrow{P} 0} - \underbrace{(p_j(\hat{\theta}_n) - p_j(\theta))}_{\xrightarrow{P} 0} \xrightarrow[n \rightarrow \infty]{P} 0,$$

falls $\hat{\theta}_n$ schwach konsistent ist und $p_j(\cdot)$ eine stetige Funktion für alle $j = 1, \dots, r$ ist.

Das heißt, unter H_0 soll $\hat{T}_n(X_1, \dots, X_n)$ relativ kleine Werte annehmen. Eine signifikante Abweichung von diesem Verhalten soll zur Ablehnung von H_0 führen, vergleiche Punkt 4.

Für die Verteilung F_θ von X_i gelten folgende Regularitätsvoraussetzungen (vergleiche Satz 3.4.2, Vorlesungsskript Stochastik I).

1. Die Verteilungsfunktion F_θ ist entweder diskret oder absolut stetig für alle $\theta \in \Theta$.
2. Die Parametrisierung ist eindeutig, das heißt: $\theta \neq \theta_1 \Leftrightarrow F_\theta \neq F_{\theta_1}$.

3. Der Träger der Likelihood-Funktion

$$L(x, \theta) = \begin{cases} P_\theta(X_1 = x), & \text{im Falle von diskreten } F_\theta, \\ f_\theta(x), & \text{im absolut stetigen Fall.} \end{cases}$$

$\text{Supp}L(x, \theta) = \{x \in \mathbb{R} : L(x, \theta) > 0\}$ hängt nicht von θ ab.

4. $L(x, \theta)$ sei 3 Mal stetig differenzierbar, und es gelte für $k = 1, \dots, 3$ und $i_1, \dots, i_k \in \{1 \dots m\}$, daß

$$\left(\sum \right) \int \frac{\partial^k L(x, \theta)}{\partial \theta_{i_1} \dots \partial \theta_{i_k}} dx = \frac{\partial^k}{\partial \theta_{i_1} \dots \partial \theta_{i_k}} \left(\sum \right) \int L(x, \theta) dx = 0.$$

5. Für alle $\theta_0 \in \Theta$ gibt es eine Konstante c_{θ_0} und eine messbare Funktion $g_{\theta_0} : \text{Supp}L \rightarrow \mathbb{R}_+$, sodaß

$$\left| \frac{\partial^3 \log L(x, \theta)}{\partial \theta_{i_1} \partial \theta_{i_2} \partial \theta_{i_3}} \right| \leq g_{\theta_0}(x), \quad |\theta - \theta_0| < c_{\theta_0}$$

und

$$\mathbb{E}_{\theta_0} g_{\theta_0}(X_1) < \infty.$$

Wir definieren die *Informationsmatrix von Fisher* durch

$$I(\theta) = \left(\mathbb{E} \left[\frac{\partial \log L(X_1, \theta)}{\partial \theta_i} \frac{\partial \log L(X_1, \theta)}{\partial \theta_j} \right] \right)_{i,j=1,\dots,m}. \quad (5.4.4)$$

Satz 5.4.2 (asymptotische Normalverteilttheit von konsistenten ML-Schätzern $\hat{\theta}_n$, multivariater Fall $m > 1$):

Es seien $X_1, \dots, X_n$ unabhängig und identisch verteilt mit Likelihood-Funktion L , die den Regularitätsbedingungen 1-5 genügt. Sei $I(\theta)$ positiv definit für alle $\theta \in \Theta \subset \mathbb{R}^m$. Sei $\hat{\theta}_n = \hat{\theta}(X_1, \dots, X_n)$ eine Folge von schwach konsistenten Maximum-Likelihood-Schätzern für θ . Dann gilt:

$$\sqrt{n}(\hat{\theta}_n - \theta) \xrightarrow[n \rightarrow \infty]{d} N(0, I^{-1}(\theta)).$$

Ohne Beweis; siehe den Beweis des Satzes 3.4.2, Vorlesungsskript Stochastik I.

Für unsere vergrößerte Hypothese $H_0 : p \in \{p(\theta), \theta \in \Theta\}$ stellen wir folgende, stückweise konstante, Likelihood-Funktion auf:

$$L(x, \theta) = p_j(\theta), \text{ falls } x \in (a_j, b_j].$$

Dann ist die Likelihood-Funktion der Stichprobe $(x_1, \dots, x_n)$ gleich

$$\begin{aligned} L(x_1, \dots, x_n, \theta) &= \prod_{j=1}^r p_j(\theta)^{Z_j(x_1, \dots, x_n)} \\ \Rightarrow \log L(x_1, \dots, x_n, \theta) &= \sum_{j=1}^r Z_j(x_1, \dots, x_n) \cdot \log p_j(\theta). \\ \hat{\theta}_n = \hat{\theta}(x_1, \dots, x_n) &= \underset{\theta \in \Theta}{\operatorname{argmax}} \log L(x_1, \dots, x_n, \theta) \\ \Rightarrow \sum_{j=1}^r Z_j(x_1, \dots, x_n) \frac{\partial p_j(\theta)}{\partial \theta_i} \cdot \frac{1}{p_j(\theta)} &= 0, \quad i = 1, \dots, m. \end{aligned}$$

Aus $\sum_{j=1}^r p_j(\theta) = 1$ folgt

$$\sum_{j=1}^r \frac{\partial p_j(\theta)}{\partial \theta_i} = 0 \Rightarrow \sum_{j=1}^r \frac{Z_j(x_1, \dots, x_n) - np_j(\theta)}{p_j(\theta)} \cdot \frac{\partial p_j(\theta)}{\partial \theta_i} = 0, \quad i = 1, \dots, m.$$

Lemma 5.4.4

Im obigen Fall gilt $I(\theta) = C^\top(\theta) \cdot C(\theta)$, wobei $C(\theta)$ eine $(r \times m)$ -Matrix mit Elementen

$$c_{ij}(\theta) = \frac{\partial p_i(\theta)}{\partial \theta_j} \cdot \frac{1}{\sqrt{p_i(\theta)}} \quad \text{ist.}$$

Beweis

$$\begin{aligned} \mathbb{E}_0 \left[\frac{\partial \log L(X_1, \theta)}{\partial \theta_i} \cdot \frac{\partial \log L(X_1, \theta)}{\partial \theta_j} \right] &= \sum_{k=1}^r \frac{\partial \log p_k(\theta)}{\partial \theta_i} \cdot \frac{\partial \log p_k(\theta)}{\partial \theta_j} \cdot p_k(\theta) \\ &= \sum_{k=1}^r \frac{\partial p_k(\theta)}{\partial \theta_i} \frac{1}{p_k(\theta)} \cdot \frac{\partial p_k(\theta)}{\partial \theta_j} \cdot \frac{1}{p_k(\theta)} \cdot p_k(\theta) \\ &= \left(C^\top(\theta) \cdot C(\theta) \right)_{ij}, \end{aligned}$$

$$\text{denn } \log L(X_1, \theta) = \sum_{i=1}^r \log p_i(\theta) \cdot \mathbb{I}(x \in (a_i, b_i]).$$

□

Deshalb gilt die Folgerung aus Satz 5.4.2:

Folgerung 5.4.1. Sei $\hat{\theta}_n = \hat{\theta}(X_1, \dots, X_n)$ ein Maximum-Likelihood-Schätzer von θ im vergrößerten Modell, der schwach konsistent ist und den obigen Regularitätsbedingungen genügt. Sei die Informationsmatrix von Fisher $I(\theta) = C^\top(\theta) \cdot C(\theta)$ für alle $\theta \in \Theta$ positiv definit. Dann ist $\hat{\theta}$ asymptotisch normalverteilt:

$$\sqrt{n}(\hat{\theta}_n - \theta) \xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, I^{-1}(\theta))$$

Satz 5.4.3

Es sei $\hat{\theta}_n$ ein Maximum-Likelihood-Schätzer im vergrößerten Modell für θ , für den alle Voraussetzungen der Folgerung 5.4.1 erfüllt sind. Die Teststatistik

$$\hat{T}_n(X_1, \dots, X_n) = \sum_{j=1}^r \frac{(Z_j(X_1, \dots, X_n) - np_j(\hat{\theta}_n))^2}{np_j(\hat{\theta}_n)}$$

ist unter H_0 asymptotisch χ_{r-m-1}^2 -verteilt:

$$\lim_{n \rightarrow \infty} \mathbb{P}_\theta \left(\hat{T}_n(X_1, \dots, X_n) > \chi_{r-m-1, 1-\alpha}^2 \right) = \alpha.$$

ohne Beweis (siehe [23]).

Aus diesem Satz folgt, daß der χ^2 -Pearson-Fisher-Test ein asymptotischer Test zum Niveau α ist.

Beispiel 5.4.1 1. χ^2 -Pearson-Fisher-Test der Normalverteilung

Sei $(X_1, \dots, X_n)$ eine Zufallsstichprobe. Es soll geprüft werden, ob $X_i \sim N(\mu, \sigma^2)$. Es gilt

$$\theta = (\mu, \sigma^2) \in \Theta = \mathbb{R} \times \mathbb{R}_+.$$

Sei $(a_j, b_j]_{j=1, \dots, r}$ eine beliebige Aufteilung von $\mathbb{R}$ in r disjunkte Intervalle. Sei

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

die Dichte der $N(\mu, \sigma^2)$ -Verteilung.

$$p_j(\theta) = \mathbb{P}_0(a_j < X_1 \leq b_j) = \int_{a_j}^{b_j} f_\theta(x) dx, \quad j = 1, \dots, r$$

mit den Klassenstärken

$$Z_j = \# \{i : X_i \in (a_j, b_j]\}.$$

Wir suchen den Maximum-Likelihood-Schätzer im vergrößerten Modell:

$$\begin{aligned} \frac{\partial p_j(\theta)}{\partial \mu} &= \int_{a_j}^{b_j} \frac{\partial}{\partial \mu} f_\theta(x) dx = \frac{1}{\sqrt{2\pi\sigma^2}} \cdot \int_{a_j}^{b_j} \frac{x-\mu}{\sigma^2} \cdot e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} dx \\ \frac{\partial p_j(\theta)}{\partial \sigma^2} &= \int_{a_j}^{b_j} \frac{\partial}{\partial \sigma^2} f_\theta(x) dx \\ &= \frac{1}{\sqrt{2\pi}} \int_{a_j}^{b_j} \left[-\frac{1}{2} \cdot \frac{1}{(\sigma^2)^{3/2}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} + \frac{1}{\sqrt{\sigma^2}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2} \cdot \left(\frac{(x-\mu)^2}{2(\sigma^2)^2} \right) \right] dx \\ &= -\frac{1}{2} \frac{1}{\sigma^2} \int_{a_j}^{b_j} f_\theta(x) dx + \frac{1}{2(\sigma^2)^2} \int_{a_j}^{b_j} (x-\mu)^2 f_\theta(x) dx \end{aligned}$$

Die notwendigen Bedingungen des Maximums sind:

$$\begin{aligned} \sum_{i=1}^r Z_j \frac{\int_{a_j}^{b_j} x f_\theta(x) dx}{\int_{a_j}^{b_j} f_\theta(x) dx} - \mu \underbrace{\sum_{j=1}^r Z_j}_{=n} &= 0, \\ \frac{1}{\sigma^2} \sum_{j=1}^r Z_j \frac{\int_{a_j}^{b_j} (x-\mu)^2 f_\theta(x) dx}{\int_{a_j}^{b_j} f_\theta(x) dx} - \underbrace{\sum_{j=1}^r Z_j}_{=n} &= 0. \end{aligned}$$

Daraus folgen die Maximum-Likelihood-Schätzer $\hat{\mu}$ und $\hat{\sigma}^2$ für μ und σ^2 :

$$\hat{\mu} = \frac{1}{n} \sum_{j=1}^r Z_j \frac{\int_{a_j}^{b_j} x f_{\theta}(x) dx}{\int_{a_j}^{b_j} f_{\theta}(x) dx},$$

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{j=1}^r Z_j \frac{\int_{a_j}^{b_j} (x - \mu)^2 f_{\theta}(x) dx}{\int_{a_j}^{b_j} f_{\theta}(x) dx}.$$

Wir konstruieren eine Näherung zu $\hat{\mu}$ und $\hat{\sigma}^2$ für $r \rightarrow \infty$. Falls $r \rightarrow \infty$ (und somit auch $n \rightarrow \infty$), dann ist $b_j - a_j$ klein und nach der einfachen Quadraturregel gilt:

$$\int_{a_j}^{b_j} x f_{\theta}(x) dx \approx (b_j - a_j) y_j f_{\theta}(y_j),$$

$$\int_{a_j}^{b_j} f_{\theta}(x) dx \approx (b_j - a_j) f_{\theta}(y_j),$$

wobei $y_1 = b_1$, $y_r = b_{r-1} = a_r$,

$$y_j = (b_{j+1} + b_j)/2, \quad j = 2, \dots, r-1.$$

Daraus folgen für die Maximum-Likelihood-Schätzer $\hat{\mu}$ und $\hat{\sigma}^2$:

$$\hat{\mu} \approx \frac{1}{n} \sum_{j=1}^r y_j \cdot Z_j = \tilde{\mu}$$

$$\hat{\sigma}^2 \approx \frac{1}{n} \sum_{j=1}^r (y_j - \tilde{\mu})^2 Z_j = \tilde{\sigma}^2,$$

$$\tilde{\theta} = (\tilde{\mu}, \tilde{\sigma}^2).$$

Der χ^2 -Pearson-Fisher-Test lautet dann: H_0 wird abgelehnt, falls

$$\hat{T}_n = \frac{\sum_{j=1}^r \left(Z_j - np_j(\tilde{\theta}) \right)^2}{np_j(\tilde{\theta})} > \chi_{r-3, 1-\alpha}^2.$$

2. χ^2 -Pearson-Fisher-Test der Poissonverteilung

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen. Wir wollen testen, ob $X_i \sim \text{Poisson}(\lambda)$, $\lambda > 0$. Es gilt $\theta = \lambda$ und $\Theta = (0, +\infty)$. Die Vergrößerung von Θ hat die Form

$$-\infty = a_1 < \underbrace{b_1}_{=0} = a_2 < \underbrace{b_2}_{=1} = a_3 < \dots < \underbrace{b_{r-1}}_{=r-2} = a_r < b_r = +\infty.$$

Dann ist

$$\begin{aligned}
 p_j(\lambda) &= \mathbb{P}_\lambda(X_1 = j - 1) = e^{-\lambda} \frac{\lambda^{j-1}}{(j-1)!}, \quad j = 1, \dots, r-1, \\
 p_r(\lambda) &= \sum_{i=r-1}^{\infty} e^{-\lambda} \frac{\lambda^i}{i!}, \\
 \frac{dp_j(\lambda)}{d\lambda} &= -e^{-\lambda} \frac{\lambda^{j-1}}{(j-1)!} + (j-1) \frac{\lambda^{j-2}}{(j-1)!} e^{-\lambda} = e^{-\lambda} \frac{\lambda^{j-1}}{(j-1)!} \left(\frac{j-1}{\lambda} - 1 \right) \\
 &= p_j(\lambda) \cdot \left(\frac{j-1}{\lambda} - 1 \right), \quad j = 1, \dots, r-1 \\
 \frac{dp_r(\lambda)}{d\lambda} &= \sum_{i \geq r-1} p_i(\lambda) \left(\frac{i-1}{\lambda} - 1 \right).
 \end{aligned}$$

Die Maximum-Likelihood-Gleichung lautet

$$0 = \sum_{j=1}^{r-1} Z_j \cdot \left(\frac{j-1}{\lambda} - 1 \right) + Z_r \frac{\sum_{i \geq r-1} p_i(\lambda) \left(\frac{i-1}{\lambda} - 1 \right)}{p_r(\lambda)}$$

Falls $r \rightarrow \infty$, so findet sich $r(n)$ für jedes n , für das $Z_{r(n)} = 0$. Deshalb gilt für $r > r(n)$:

$$\sum_{j=1}^{r-1} (j-1)Z_j - \lambda \underbrace{\sum_{j=1}^r Z_j}_{=n} = 0,$$

woraus der Maximum-Likelihood-Schätzer

$$\frac{1}{n} \sum_{j=1}^{r-1} (j-1)Z_j = \frac{1}{n} \sum_{j=1}^n X_j = \bar{X}_n$$

folgt. Der χ^2 -Pearson-Fisher-Test lautet: H_0 wird verworfen, falls

$$\hat{T}_n = \sum_{j=1}^r \frac{\left(Z_j - np_j(\bar{X}_n) \right)^2}{\left(np_j(\bar{X}_n) \right)^2} > \chi_{r-2, 1-\alpha}^2.$$

5.4.3 Anpassungstest von Shapiro

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen, identisch verteilten Zufallsvariablen, $X_i \sim F$. Getestet werden soll die Hypothese

$$H_0 : F \in \{N(\mu, \sigma^2) : \mu \in \mathbb{R}, \sigma^2 > 0\} \text{ vs. } H_1 : F \notin \{N(\mu, \sigma^2), \mu \in \mathbb{R}, \sigma^2 > 0\}.$$

Die in den Abschnitten 5.4.1 - 5.4.2 vorgestellten χ^2 -Tests sind asymptotisch; deshalb können sie für relativ kleine Stichprobenumfänge nicht verwendet werden.

Der folgende Test wird diese Lücke füllen und eine Testentscheidung über H_0 selbst bei kleinen Stichproben ermöglichen.

Man bildet Ordnungsstatistiken $X_{(1)}, \dots, X_{(n)}$, $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ und vergleicht ihre Korreliertheit mit den Mittelwerten der entsprechenden Ordnungsstatistiken der $N(0, 1)$ -Verteilung. Sei $(Y_1, \dots, Y_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen, $Y_1 \sim N(0, 1)$. Es sei $a_i = \mathbb{E} Y_{(i)}$, $i = 1, \dots, n$. Falls der empirische Korrelationskoeffizient ρ_{aX} zwischen $(a_1, \dots, a_n)$ und $(X_{(1)}, \dots, X_{(n)})$ bei 1 liegt, dann ist die Stichprobe normalverteilt. Formalisieren wir diese Heuristik:

Es sei b_i der Erwartungswert der i -ten Ordnungsstatistik in einer Stichprobe von $N(\mu, \sigma^2)$ -verteilten, unabhängigen Zufallsvariablen Z_i : $b_i = \mathbb{E} Z_{(i)}$, $i = 1, \dots, n$. Es gilt: $b_i = \mu + \sigma a_i$, $i = 1, \dots, n$. Betrachten wir den Korrelationskoeffizienten

$$\rho_{bX} = \frac{\sum_{i=1}^n (b_i - \bar{b}_n) (X_{(i)} - \bar{X}_n)}{\sqrt{\sum_{i=1}^n (b_i - \bar{b}_n)^2 \sum_{i=1}^n (X_{(i)} - \bar{X}_n)^2}}. \quad (5.4.5)$$

Da ρ invariant bezüglich Lineartransformationen ist und

$$\sum_{i=1}^n a_i = \sum_{i=1}^n \mathbb{E} Y_i = \mathbb{E} \left(\sum_{i=1}^n Y_i \right) = 0, \quad \text{gilt:}$$

$$\begin{aligned} \rho_{bX} &\stackrel{(\text{Stochastik I})}{=} \rho_{aX} = \frac{\sum_{i=1}^n a_i (X_{(i)} - \bar{X}_n)}{\sqrt{\sum_{i=1}^n a_i^2 \sum_{i=1}^n (X_i - \bar{X}_n)^2}} = \frac{\sum_{i=1}^n a_i X_{(i)} - \bar{X}_n \overbrace{\sum_{i=1}^n a_i}^{=0}}{\sqrt{\sum_{i=1}^n a_i^2 \sum_{i=1}^n (X_i - \bar{X}_n)^2}} \\ &= \frac{\sum_{i=1}^n a_i X_{(i)}}{\sqrt{\sum_{i=1}^n a_i^2 \cdot \sum_{i=1}^n (X_i - \bar{X}_n)^2}} \end{aligned}$$

Die Teststatistik lautet:

$$T_n = \frac{\sum_{i=1}^n a_i X_{(i)}}{\sqrt{\sum_{i=1}^n a_i^2 \sum_{i=1}^n (X_i - \bar{X}_n)^2}} \quad (\text{Shapiro-Francia-Test})$$

Die Werte a_i sind bekannt und können den Tabellen bzw. der Statistik-Software entnommen werden. Es gilt: $|T_n| \leq 1$.

H_0 wird abgelehnt, falls $T_n \leq q_{n,\alpha}$, wobei $q_{n,\alpha}$ das α -Quantil der Verteilung von T_n ist. Diese Quantile sind aus den Tabellen bekannt, bzw. können durch Monte-Carlo-Simulationen berechnet werden.

Bemerkung 5.4.3. Einen anderen, weit verbreiteten Test dieser Art bekommt man, wenn man die Lineartransformation $b_i = \mu + \sigma a_i$ durch eine andere Lineartransformation ersetzt:

$$(a'_1, \dots, a'_n)^\top = K^{-1} \cdot (a_1, \dots, a_n),$$

wobei $K = (k_{ij})_{j=1}^n$ die Kovarianzmatrix von $(Y_{(1)}, \dots, Y_{(n)})$ ist:

$$k_{ij} = \mathbb{E} \left(Y_{(i)} - a_i \right) \left(Y_{(j)} - a_j \right), \quad i, j = 1, \dots, n,$$

Der so konstruierte Test trägt den Namen Shapiro-Wilk-Test.

5.5 Weitere, nicht parametrische Tests

5.5.1 Binomialtest

Es sei $(X_1, \dots, X_n)$ eine Stichprobe von unabhängigen, identisch verteilten Zufallsvariablen, wobei $X_i \sim \text{Bernoulli}(p)$. Getestet werden soll:

$$H_0 : p = p_0 \text{ vs. } H_1 : p \neq p_0$$

Die Teststatistik lautet

$$T_n = \sum_{i=1}^n X_i \underset{H_0}{\sim} \text{Bin}(n, p_0),$$

und die Entscheidungsregel ist: H_0 wird verworfen, falls

$$T_n \notin [\text{Bin}(n, p_0)_{\alpha/2}, \text{Bin}(n, p_0)_{1-\alpha/2}],$$

wobei $\text{Bin}(n, p)_\alpha$ das α -Quantil der $\text{Bin}(n, p)$ -Verteilung ist.

Für andere H_0 , wie zum Beispiel $p \leq p_0$ ($p \geq p_0$) muss der Ablehnungsbereich entsprechend angepasst werden.

Die Quantile $\text{Bin}(n, p)_\alpha$ erhält man aus Tabellen oder aus Monte-Carlo-Simulationen. Falls n groß ist, können diese Quantile durch die Normalapproximation berechnet werden:

Nach dem zentralen Grenzwertsatz von DeMoivre-Laplace gilt:

$$\mathbb{P}(T_n \leq x) = \mathbb{P}\left(\frac{T_n - np_0}{\sqrt{np_0(1-p_0)}} \leq \frac{x - np_0}{\sqrt{np_0(1-p_0)}}\right) \underset{n \rightarrow \infty}{\approx} \Phi\left(\frac{x - np_0}{\sqrt{np_0(1-p_0)}}\right).$$

Daraus folgt:

$$\begin{aligned} z_\alpha &\approx \frac{\text{Bin}(n, p_0)_\alpha - np_0}{\sqrt{np_0(1-p_0)}} \\ \Rightarrow \text{Bin}(n, p_0)_\alpha &\approx \sqrt{np_0(1-p_0)} \cdot z_\alpha + np_0 \end{aligned}$$

Nach der Poisson-Approximation (für $n \rightarrow \infty, np_0 \rightarrow \lambda_0$) gilt:

$$\begin{aligned} \text{Bin}(n, p_0)_{\alpha/2} &\approx \text{Poisson}(\lambda_0)_{\alpha/2}, \\ \text{Bin}(n, p_0)_{1-\alpha/2} &\approx \text{Poisson}(\lambda_0)_{1-\alpha/2}, \quad \text{wobei } \lambda_0 = np_0. \end{aligned}$$

Zielstellung: Wie kann mit Hilfe des oben beschriebenen Binomialtests die Symmetrieeigenschaft einer Verteilung getestet werden?

Es sei $(Y_1, \dots, Y_n)$ eine Stichprobe von unabhängigen und identisch verteilten Zufallsvariablen mit Verteilungsfunktion F . Getestet werden soll:

$$H_0 : F \text{ ist symmetrisch vs. } H_1 : F \text{ ist nicht symmetrisch.}$$

Eine symmetrische Verteilung besitzt den Median bei Null. Deswegen vergrößern wir die Hypothese H_0 und testen:

$$H'_0 : F^{-1}(0,5) = 0 \text{ vs. } H'_1 : F^{-1}(0,5) \neq 0.$$

Noch allgemeiner: Für ein $\beta \in [0, 1]$:

$$H''_0 : F^{-1}(\beta) = \gamma_\beta \text{ vs. } H''_1 : F^{-1}(\beta) \neq \gamma_\beta.$$

H''_0 vs. H''_1 wird mit Hilfe des Binomialtests wie folgt getestet: Sei $X_i = \mathbb{I}(Y_i \leq \gamma_\beta)$. Unter H''_0 gilt:

$$X_i \sim \text{Bernoulli}(F(\gamma_\beta)) = \text{Bernoulli}(\beta).$$

Seien $a_1 = -\infty$, $b_1 = \gamma_\alpha$, $a_2 = b_1$, $b_2 = +\infty$ zwei disjunkte Klassen $(a_1, b_1]$, $(a_2, b_2]$ in der Sprache des χ^2 -Pearson-Tests. Die Testgröße ist:

$$T_n = \sum_{i=1}^n X_i = \# \{Y_i : Y_i \leq \gamma_\beta\} \sim \text{Bin}(n, \beta), \quad p = F(\gamma_\beta)$$

Die Hypothese $F^{-1}(\beta) = \gamma_\beta$ ist äquivalent zu $H'''_0 : p = \beta$. Die Entscheidungsregel lautet dann: H'''_0 wird verworfen, falls $T_n \notin [\text{Bin}(n, \beta)_{\alpha/2}, \text{Bin}(n, \beta)_{1-\alpha/2}]$. Dies ist ein Test zum Niveau α .

5.5.2 Iterationstests auf Zufälligkeit

In manchen Fragestellungen der Biologie untersucht man eine Folge von 0 oder 1 auf ihre „Zufälligkeit“ bzw. Vorhandensein von größeren Clustern von 0 oder 1. Diese Hypothesen kann man mit Hilfe der sogenannten *Iterationstests* statistisch überprüfen.

Sei eine Stichprobe $X_i, i = 1, \dots, n$ gegeben, $X_i \in \{0, 1\}$, $\sum_{i=1}^n X_i = n_1$ die Anzahl der Einsen, $n_2 = n - n_1$ die Anzahl der Nullen, n_1, n_2 vorgegeben. Eine Realisierung von $(X_1, \dots, X_n)$ mit $n = 18$, $n_1 = 12$ wäre zum Beispiel

$$x = (0, 0, 0, 1, 1, 1, 0, 1, 1, 0, 1, 1, 0, 1, 1, 1, 1)$$

Es soll getestet werden, ob

$$H_0 : \text{jede Folge } x \text{ ist gleichwahrscheinlich vs.}$$

$$H_1 : \text{Es gibt bevorzugte Folgen (Clusterbildung)}$$

stimmt.

Sei

$$\Omega = \left\{ x = (x_1, \dots, x_n) : x_i = 0 \text{ oder } 1, i = 1, \dots, n, \sum_{i=1}^n x_i = n_1 \right\}$$

der Stichprobenraum. Dann ist der Raum $(\Omega, \mathcal{F}, \mathbb{P})$ mit $\mathcal{F} = \mathcal{P}(\Omega)$,

$$\mathbb{P}(x) = \frac{1}{|\Omega|} = \frac{1}{\binom{n}{n_1}}$$

ein Laplacescher Wahrscheinlichkeitsraum.

Sei

$$\begin{aligned} T_n(X) &= \#\{\text{Iterationen in } X\} = \#\{\text{Teilfolgen der Nullen oder Einsen}\} \\ &= \#\{\text{Wechselstellen von 0 auf 1 oder von 1 auf 0}\} + 1. \end{aligned}$$

Zum Beispiel ist für $x = (0, 1, 1, 1, 0, 0, 0, 1, 0, 1, 1, 1, 0, 0)$, $T_n(x) = 7 = 6 + 1$.

$T_n(X)$ wird folgendermaßen als Teststatistik für H_0 vs. H_1 benutzt. H_0 wird abgelehnt, falls $T(x)$ klein ist, das heißt, falls $T_n(x) < F_{T_n}^{-1}(\alpha)$. Dies ist ein Test zum Niveau α . Wie berechnen wir die Quantile $F_{T_n}^{-1}$?

Satz 5.5.1

Unter H_0 gelten folgende Aussagen:

1.

$$\mathbb{P}(T_n = k) = \begin{cases} \frac{2 \binom{n_1-1}{i-1} \binom{n_2-1}{i-1}}{\binom{n}{n_1}}, & \text{falls } k = 2i, \\ \frac{\binom{n_1-1}{i} \cdot \binom{n_2-1}{i-1} + \binom{n_1-1}{i-1} \cdot \binom{n_2-1}{i}}{\binom{n}{n_1}}, & \text{falls } k = 2i + 1. \end{cases}$$

2.

$$\mathbb{E} T_n = 1 + \frac{2n_1 n_2}{n}$$

3.

$$\text{Var}(T_n) = \frac{2n_1 n_2 (2n_1 n_2 - n)}{n^2 (n - 1)}$$

Beweis 1. Wir nehmen an, daß $k = 2i$ (der ungerade Fall ist analog). Wie können i Klumpen von Einsen gewählt werden? Die Anzahl dieser Möglichkeiten = die Anzahl der Möglichkeiten, wie n_1 Teilchen auf i Klassen verteilt werden.

$$0|00|\dots|0|\binom{n_1}{n_1}$$

Dies ist gleich der Anzahl an Möglichkeiten, wie $i - 1$ Trennwände auf $n_1 - 1$ Positionen verteilt werden können $= \binom{n_1-1}{i-1}$. Das selbe gilt für die Nullen.

2. Sei $Y_j = \mathbb{I}\{X_{j-1} \neq X_j\}_{j=2,\dots,n}$.

$$\Rightarrow \mathbb{E} T_n(X) = 1 + \sum_{j=2}^n \mathbb{E} Y_j = 1 + \sum_{j=2}^n \mathbb{P}(X_{j-1} \neq X_j).$$

$$\begin{aligned} \mathbb{P}(X_{j-1} \neq X_j) &= \frac{2 \binom{n-2}{n_1-1}}{\binom{n}{n_1}} = 2 \cdot \frac{\frac{(n-2)!}{(n-2-(n_1-1))!(n_1-1)!}}{\frac{n!}{(n-n_1)!n_1!}} \\ &= \frac{2n_1(n-n_1)}{(n-1)n} \\ &= \frac{2n_1 n_2}{n(n-1)}. \end{aligned}$$

Daraus folgt

$$\mathbb{E} T_n = 1 + (n-1) \frac{2n_1 n_2}{n(n-1)} = 1 + 2 \frac{n_1 n_2}{n}.$$

3.

Übungsaufgabe 5.5.1

Beweisen Sie Punkt 3. □

Beispiel 5.5.1 (Test von Wald-Wolfowitz):

Seien $Y = (Y_1, \dots, Y_n)$, $Z = (Z_1, \dots, Z_n)$ unabhängige Stichproben von unabhängigen und identisch verteilten Zufallsvariablen, $Y_i \sim F$, $Z_i \sim G$. Getestet werden soll:

$$H_0 : F = G \text{ vs. } H_1 : F \neq G.$$

Sei $(Y, Z) = (Y_1, \dots, Y_n, Z_1, \dots, Z_n)$ und seien X'_i Stichprobenvariablen von (Y, Z) , $i = 1, \dots, n$, $n = n_1 + n_2$. Wir bilden die Ordnungsstatistiken $X'_{(i)}$, $i = 1, \dots, n$ und setzen

$$X_i = \begin{cases} 1, & \text{falls } X'_{(i)} = Y_j \text{ für ein } j = 1, \dots, n_1, \\ 0, & \text{falls } X'_{(i)} = Z_j \text{ für ein } j = 1, \dots, n_2. \end{cases}$$

Unter H_0 sind die Stichprobenwerte in (Y, Z) gut gemischt, das heißt jede Kombination von 0 und 1 in $(X_1, \dots, X_n)$ ist gleichwahrscheinlich. Darum können wir den Iterationstest auf Zufälligkeit anwenden, um H_0 vs. H_1 zu testen: H_0 wird verworfen, falls $T_n(x) \leq F^{-1}(\alpha)$, $x = (x_1, \dots, x_n)$.

Wie können die Quantile von F_{T_n} für große n berechnet werden? Falls

$$\frac{n_1}{n_1 + n_2} \xrightarrow{n \rightarrow \infty} p \in (0, 1),$$

dann ist T_n asymptotisch normalverteilt.

Satz 5.5.2

Unter der obigen Voraussetzung gilt:

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{\mathbb{E} T_n}{n} &= 2p(1-p) \\ \lim_{n \rightarrow \infty} \frac{1}{n} \text{Var } T_n &= 4p^2(1-p)^2 \\ \frac{T_n - 2p(1-p)}{2\sqrt{np(1-p)}} &\xrightarrow[n \rightarrow \infty]{d} Y \sim N(0, 1), \quad \text{falls } \frac{n_1}{n_1 + n_2} \rightarrow p \in (0, 1). \end{aligned}$$

So können Quantile von T_n näherungsweise für große n folgendermaßen berechnet werden:

$$\begin{aligned} \alpha &= \mathbb{P}(T_n \leq F_{T_n}^{-1}(\alpha)) = \mathbb{P}\left(\frac{T_n - 2np(1-p)}{2\sqrt{np(1-p)}} \leq \frac{x - 2np(1-p)}{2\sqrt{np(1-p)}}\right) \Big|_{x=F_{T_n}^{-1}(\alpha)} \\ &\approx \Phi\left(\frac{F_{T_n}^{-1}(\alpha) - 2np(1-p)}{2\sqrt{np(1-p)}}\right) \\ &\Rightarrow z_\alpha \approx \frac{F_{T_n}^{-1}(\alpha) - 2np(1-p)}{2\sqrt{np(1-p)}} \end{aligned}$$

Damit erhalten wir für die Quantile:

$$F_{T_n}^{-1}(\alpha) \approx 2np(1-p) + 2\sqrt{np(1-p)} \cdot z_\alpha$$

In der Praxis setzt man $\hat{p} = \frac{n_1}{n_1+n_2}$ für p ein.

Literaturverzeichnis

- [1] BICKEL, P. ; DOKSUM, K.: *Mathematical Statistics: Basic Ideas and Selected Topics*. London : Prentice Hall, 2001. – 2nd ed., Vol. 1
- [2] BOROVKOV, A. A.: *Mathematical Statistics*. Gordon & Breach, 1998
- [3] BURKSCHAT, M. ; CRAMER, E. ; KAMPS, U.: *Beschreibende Statistik, Grundlegende Methoden*. Berlin : Springer, 2004
- [4] CASELLA, G. ; BERGER, R. L.: *Statistical Inference*. Duxbury : Pacific Grove (CA), 2002
- [5] CRAMER, E. ; CRAMER, K. ; KAMPS, U. ; ZUCKSCHWERDT: *Beschreibende Statistik, Interaktive Grafiken*. Berlin : Springer, 2004
- [6] CRAMER, E. ; KAMPS, U.: *Grundlagen der Wahrscheinlichkeitsrechnung und Statistik*. Berlin : Springer, 2007
- [7] DALGAARD, P.: *Introductory Statistics with R*. Berlin : Springer, 2002
- [8] FAHRMEIR, L. ; KNEIB, T. ; LANG, S.: *Regression. Modelle, Methoden und Anwendungen*. Berlin : Springer, 2007
- [9] FAHRMEIR, L. ; KÜNSTLER, R. ; PIGEOT, I. ; TUTZ, G.: *Statistik. Der Weg zur Datenanalyse*. Berlin : Springer, 2001
- [10] GEORGI, H. O.: *Stochastik: Einführung in die Wahrscheinlichkeitstheorie und Statistik*. Berlin : de Gruyter, 2002
- [11] HARTUNG, J. ; ELPERT, B. ; KLÖSENER, K. H.: *Statistik*. München : R. Oldenbourg Verlag, 1993. – 9. Auflage
- [12] HEYDE, C. C. ; SENETA, E.: *Statisticians of the Centuries*. Berlin : Springer, 2001
- [13] IRLE, A.: *Wahrscheinlichkeitstheorie und Statistik, Grundlagen – Resultate – Anwendungen*. Teubner, 2001
- [14] KAZMIR, L. J.: *Wirtschaftsstatistik*. McGraw-Hill, 1996
- [15] KOCH, K. R.: *Parameter Estimation and Hypothesis Testing in Linear Models*. Berlin : Springer, 1999
- [16] KRAUSE, A. ; OLSON, M.: *The Basics of S-PLUS*. Berlin : Springer, 2002. – Third Ed.
- [17] KRENGEL, U.: *Einführung in die Wahrscheinlichkeitstheorie und Statistik*. Braunschweig : Vieweg, 2002. – 6. Auflage
- [18] LEHMANN, E. L.: *Elements of Large-Sample Theory*. New York : Springer, 1999

- [19] LEHN, J. ; WEGMANN, H.: *Einführung in die Statistik*. Stuttgart : Teubner, 2000. – 3. Auflage
- [20] MAINDONALD, J. ; BRAUN, J.: *Data Analysis and Graphics Using R*. Cambridge University Press, 2003
- [21] OVERBECK-LARISCH, M. ; DOLEJSKY, W.: *Stochastik mit Mathematica*. Braunschweig : Vieweg, 1998
- [22] PRUSCHA, H.: *Angewandte Methoden der Mathematischen Statistik*. Stuttgart : Teubner, 1996
- [23] PRUSCHA, H.: *Vorlesungen über Mathematische Statistik*. Stuttgart : Teubner, 2000
- [24] SACHS, L.: *Angewandte Statistik*. Springer, 1992
- [25] SACHS, L. ; HEDDERICH, J.: *Angewandte Statistik, Methodensammlung mit R*. Berlin : Springer, 2006
- [26] SHIRYAEV, A. N.: *Probability*. New York : Springer, 1996
- [27] SPIEGEL, M. R. ; STEPHENS, L. J.: *Statistik*. McGraw-Hill, 1999
- [28] STAHEL, W. A.: *Statistische Datenanalyse*. Vieweg, 1999
- [29] VENABLES, W. ; RIPLEY, D.: *Modern applied statistics with S-PLUS*. Springer, 1999. – 3rd ed
- [30] WASSERMAN, L.: *All of Statistics. A Concise Course in Statistical Inference*. Springer, 2004

Index

- a-posteriori-Verteilung, 84
- a-priori-Verteilung, 84
- Ablehnungsbereich, 119
- absolute Häufigkeit, 7
- Abweichung, mittlere quadratische, 16
- Annahmebereich, 119
- arithmetisches Mittel, 12
- asymptotisch erwartungstreu, 51
- asymptotisch normalverteilt, 52
- AusgangsvARIABLE, 34

- Balkendiagramm, 8
- Bandbreite, 26
- Bayes-Schätzer, 84
- Bayesche Formel, 84
- bedingte Wahrscheinlichkeit, 97
- bedingter Erwartungswert, 95
- Bernoulli-Verteilung
 - asymptotisches Konfidenzintervall, 111
- besserer Schätzer, 52
- bester erwartungstreuer Schätzer, 52
- Bestimmtheitsmaß, 38
- Bias, 51
- bimodal, 10
- Binomialverteilung, 138
- Blackwell-Rao, Ungleichung von, 105
- Bootstrap
 - Konfidenzintervall, 89
 - Schätzer, 89
- Bootstrap-Schätzer
 - Monte-Carlo-Methoden, 89
- Box-Plot, 14
 - modifizierter, 14
- Bravais-Pearson-Koeffizient, 31
- Bravais-Pearson-Korrelationskoeffizient, 29
- Brownsche Brücke, 70

- χ^2 -Verteilung, 44
- Cramér-Rao, Ungleichung von, 90

- Daten-Stichproben, 1
- Datenbereinigung, 1
- Datenerhebung, 1
- Dichteschätzung, 25
- Dichtetransformationssatz für Zufallsvektoren, 48
- gleichmäßiger Abstand D_n , 65
- Dvoretzky-Kiefer-Wolfowitz, Ungleichung von, 67

- Eindeutigkeit der besten erwartungstreuen Schätzer, 103
- Einfache lineare Regression, 34
- Einflussfaktor, 34
- einparametrische Exponentialklasse, 137
- empirische(r)
 - Kovarianz, 29
 - Median, 14
 - Standardabweichung, 16
 - Varianz, 16
 - Variationskoeffizient, 16, 17
 - Verteilungsfunktion, 10
- Entscheidungsregel, 118
- Erlangverteilung, 44
- erwartungstreu, 50
- Erwartungstreue, 17
- Erwartungswert, bedingter, 96
- Explorative Datenanalyse, 1

- Faktorisierungssatz von Neyman-Fisher, 100
- Fehler 1. Art, 70
- Fehler 1. und 2. Art, 119
- Fisher
 - Fisher-Information, 80
 - Fisher-Snedecor-Verteilung, F-Verteilung, 48
 - Wölbungsmaß von Fisher, 21
- Fisher-Informationsmatrix, 155

- Gütefunktion, 120

Gammaverteilung
 Faltungsstabilität, 44
 Momenterzeugende und charakteristische Funktion, 43
 geometrisches Mittel, 12, 13
 Gesamtstreuung, 38
 getrimmtes Mittel, 12
 Gini-Koeffizient, 17, 18
 Gini-Koeffizient, Darstellung von, 19
 Gliwenko-Cantelli, Satz von, 65
 Grundgesamtheit, 1

Häufigkeit
 absolute, 7
 relative, 7
 harmonisches Mittel, 13
 Hauptsatz über zweiseitige Tests, 146
 Herfindahl-Index, 17
 Histogramm, 7
 eindimensionales Histogramm, 7
 zweidimensionales Histogramm, 28
 Hoeffding-Ungleichung, 110
 Hypothese, 118
 Alternative, 118
 Haupthypothese, 118

identifizierbar, 41
 Information von Kullback-Leibler, 77
 Informationsmatrix von Fisher, 155
 Invarianzeigenschaften, 32
 Irrtumswahrscheinlichkeit, 106
 Iterationstest, 162

Jackknife-Schätzer für die/den
 Erwartungswert, 87
 Varianz, 87
 Verzerrung (Bias), 87

Karl Popper, 119
 Kerndichteschätzer
 eindimensionaler Kerndichteschätzer, 26
 zweidimensionaler Kerndichteschätzer, 29
 Klassenstärke, 148
 Kolmogorow, Satz von, 69
 Kolmogorow-Abstand D_n , 65
 Kolmogorow-Verteilung, 69
 Konfidenzband, 67
 Konfidenzintervall, 58, 106

 asymptotisches, 106, 111
 für die Bernoulli-Verteilung, 111
 für die Poissonverteilung, 112
 Bootstrap, 89
 Lange, 106
 minimales, 106
 Konfidenzniveau, 106
 konsistenter Schätzer, 51
 Konstanzbereich, 68
 Konzentrationsrate, 17, 20
 Korrelationskoeffizient, 29
 Spearman's, 31
 Kovarianz, empirische, 29
 Kreisdiagramm, 7, 8
 kritischer Bereich, *siehe* Ablehnungsbereich
 Kullback-Leibler, Information von, 77
 Kurtosis, 21

Lagemaß, 12
 Lehmann-Scheffé, Satz von, 104
 Lernstichprobe, 3
 Likelihood-Funktion, 74
 linksschief, 10
 linkssteil, 10
 Lorenz-Münzner-Koeffizient, 20
 Lorenzkurve, 17

maximale Streuung, 16
 Maximum-Likelihood-Schätzer, 74, 75
 schwache Konsistenz, 78
 Median, 12, 15
 empirischer, 14
 Mittel
 arithmetisches, 12
 geometrisches, 12, 13
 getrimmtes, 12
 harmonisches, 12, 13
 Mittelwert, 12, 15
 mittlere quadratische Abweichung, 16
 mittlerer quadratischer Fehler, 51
 Modalität, 10
 Modellierung von Daten, 1
 Modellvalidierung, 1
 modifizierter Box-Plot, 14
 Modus, 12, 15
 Momentenmethode, 72
 Momentenschätzer, 72

- multimodal, 10
- Multinomialverteilung, 148
- Neyman-Fisher, Faktorisierungssatz, 100
- Neyman-Pearson
 - Fundamentallemma, 134
 - Optimalitätssatz, 133
- Normalverteilung
 - Konfidenzintervall
 - für eine Stichprobe, 108
 - für zwei Stichproben, 113
 - Signifikanztests, 126
- Ordnungsstatistik, 6, 12, 13
- p -Wert, 123
- Parameterraum, 41
- Parametervektor, 41
- Pearson-Teststatistik, 149
- Plug-in-Methode, 70
- Plug-in-Schätzer, 70, 71
- Poissonverteilung, 115, 128, 130
 - asymptotisches Konfidenzintervall, 112
 - Neyman-Fisher-Test, 158
 - Neyman-Pearson-Test, 136
- Polynomiale Regression, 34
- Punktschätzer, 41
- Quantil, 12, 13
- Quantilplot, 22
- Quartil, 12, 14
- Randomisierungsbereich, 119
- Rangkorrelationskoeffizient, 31
- Realisierung, 5, 7
- rechtsschief, 10
- rechtssteil, 10
- Regressand, 34
- Regression, 34
 - einfache lineare, 34
 - polynomiale, 34
- Regressionsgerade, 34
- Regressionsgerade, Eigenschaften von, 37
- Regressionskoeffizient, 34
- Regressionskonstante, 34
- Regressionsvarianz, 34
- Regressor, 34
- relative Häufigkeit, 7
- Resampling-Methode, 86
- Residualplot, 39
- Residuen, 36
- Säulendiagramm, 8
- Satz
 - χ^2 -Verteilung, Spezialfall, 44
 - Darstellung des Gini-Koeffizient, 19
 - Dichte der t -Verteilung, 46
 - Dichtetransformationssatz für Zufallsvektoren, 48
 - Eigenschaften der empirischen Momente, 52
 - Eigenschaften des bedingten Erwartungswertes, 96
 - Faktorisierungssatz von Neyman-Fisher, 100
 - Gliwenko-Cantelli, 65
 - Invarianzeigenschaften, 32
 - Kolmogorow, 69
 - Lehmann-Scheffé, 104
 - Momenterzeugende und charakteristische Funktion der Gammaverteilung, 43
 - Schwache Konsistenz von ML-Schätzern, 78
 - Ungleichung von Cramér-Rao, 90
 - Ungleichung von Dvoretzky-Kiefer-Wolfowitz, 67
- Schätzer, 50
 - besserer, 52
 - bester erwartungstreuer, 52
 - konsistenter, 51
 - suffizienter, 97
 - Vergleich von, 52
- Schiefe, 12, 20
- Schließende Datenanalyse, 2
- Spannweite, 16
- Spearman's Korrelationskoeffizient, 31
- Stabdiagramm, 7
- Stamm-Blatt-Diagramm, 9
- Standardabweichung, 17
- Statistische Merkmale, 4
- stem-leaf display, 9
- Stichproben, 5
- Stichprobenfunktion, 6
- Stichprobenmittel, 6, 12
- Stichprobenvarianz, 6, 16

- Streudiagramm, 28, 39
- suffizienter Schätzer, 97
- sum of squared residuals, 38
- sum of squares explained, 38
- sum of squares total, 38
- Symmetriekoeffizient, 20
- symmetrisch, 10
- t -Verteilung, 46
- Test
 - Anpassungstest, 148
 - Anpassungstest von Shapiro, 159
 - asymptotischer, 121, 127
 - besserer, 131
 - bester, 132
 - Binomialtest, 161
 - χ^2 -Anpassungstest, 148
 - χ^2 -Pearson-Fisher-Test, 154
 - Iterationstest, 162
 - Kolmogorov-Smirnov, 148
 - Macht, 120
 - Monte-Carlo-Test, 121
 - Neyman-Pearson-Test, 132
 - Ablehnungsbereich, 132
 - einseitiger, 137
 - modifizierter, 144
 - Parameter der Poissonverteilung, 136
 - Umfang, 132
 - NP-Test, *siehe* Neyman-Pearson-Test
 - Parameter der Normalverteilung, 126
 - parametrischer, 120
 - einseitiger, 120
 - linksseitiger, 120
 - rechtsseitiger, 120
 - zweiseitiger, 120
 - parametrischer Signifikanztest, 126
 - power, *siehe* Macht
 - randomisierter, 119, 131
 - Schärfe, 120
 - von Shapiro-Francia, 160
 - von Shapiro-Wilk, 161
 - Stärke, 120
 - Umfang, 131
 - unverfälschter, 125
 - Wald-Test, 127
 - von Wald-Wolfowitz, 164
- Teststatistik, 107
- Tortendiagramm, 8
- Transformationsregel, 16
- unimodal, 10, 78
- unverzerrt, 50
- Varianz, empirische, 16
- Verlustfunktion, 84
- Verteilung mit monotonem Dichtekoeffizienten, 137
- verteilungsfrei, 67
- Verteilungsfunktion, empirische, 10
- Vertrauensintervall, 58
- Verzerrung, 51
- Visualisierung, 1
- Vollständigkeit, 102
- Wölbung, 12
- Wölbungsmaß von Fisher, 21
- Zielgröße, 34
- Zufallsstichprobe, 5
- Zufallsvektoren
 - Dichtetransformationssatz, 48
- zweidimensionaler Kerndichteschätzer, 29
- zweidimensionales Histogramm, 28