

Biometrie und Epidemiologie

Max Resch und Johannes Maschler

Sommersemester 2010

Inhaltsverzeichnis

1	Vorlesung am 4. März 2010	6
1.1	Überlebenszeiten und -funktionen	6
1.1.1	Bevölkerungsstatistik und Sterbetafeln	6
1.1.1.1	Sterbetafel	6
1.1.1.2	Lebenserwartung	6
1.1.2	Verteilung von Überlebenszeiten	7
1.1.2.1	Überlebenszeit	7
2	Vorlesung am 11. März 2010	11
2.0.3	Verteilung von Überlebenszeiten	11
2.0.3.1	Exponentialverteilung	11
2.0.3.2	Weibull Verteilung	12
2.0.4	Prognose Studien	13
2.0.4.1	Kaplan-Meier-Methode	13
2.0.4.2	Logrank-Test	14
3	Vorlesung am 18. März 2010	18
3.1	Wahrscheinlichkeiten in der Epidemiologie	18
3.1.1	Bevölkerungsstatistik	19
3.1.2	Diagnostische Tests	19
3.1.2.1	Vorhersagewerte	19
3.2	Diskrete Verteilungen	20
3.2.1	Diskrete Zufallsvariablen	20
3.2.2	Binomialverteilung	22
3.2.2.1	Bernoulli-Experiment	22
3.2.2.2	Bernoulli-Prozess	22
3.2.3	Poissonverteilung	22
4	Vorlesung am 25. März 2010	24
4.0.4	Polynomialverteilung	25
4.0.5	Negative Binomialverteilung	25
4.0.5.1	Geometrische Verteilung	26
4.0.5.2	Definition der negativen Binomialverteilung	26
4.0.6	Hypergeometrische Verteilung	26
4.0.6.1	Annäherung durch die Binomialverteilung	27
4.0.7	Übersicht über die diskreten Verteilungen	27

Inhaltsverzeichnis

4.1	Stetige Verteilungen	27
4.1.1	Normalverteilung	27
4.1.1.1	Standardnormalverteilung	27
4.1.1.2	σ -Bereiche und Referenzbereiche	28
4.1.2	Normalisierende Transformationen	29
4.1.2.1	Lognormalverteilung	29
5	Vorlesung am 15. April 2010	31
5.0.3	Sätze der Wahrscheinlichkeitsrechnung	31
5.0.3.1	Tschebyscheff'sche Ungleichung	31
5.0.3.2	Gesetz der großen Zahlen	31
5.0.3.3	Zentraler Grenzwertsatz	32
5.0.4	Prüfverteilungen	33
5.0.4.1	t-Verteilung	33
5.0.4.2	χ^2 -Verteilung	34
5.0.4.3	F-Verteilung	36
5.1	Schätzverfahren	37
5.1.1	Punktschätzungen	37
5.1.1.1	Kriterien zur Güte einer Schätzung	37
5.1.2	Spezielle Schätzfunktionen	37
5.1.2.1	Erwartungswert	37
5.1.2.2	Median	38
5.1.2.3	Varianz	38
5.1.2.4	Schätzung der Wahrscheinlichkeit	41
6	Vorlesung am 22. April 2010	42
6.0.3	Intervallschätzungen	42
6.0.3.1	Bedeutung eines Konfidenzintervalls	42
6.0.3.2	Konfidenzintervall für einen Erwartungswert	42
6.0.3.3	Konfidenzintervall für eine Wahrscheinlichkeit	44
6.0.4	Bedeutung des Stichprobenumfangs	45
7	Vorlesung am 29. April 2010	47
7.1	Das Prinzip eines statistischen Tests	47
7.1.1	Durchführung eines Statistischen Tests	47
7.1.2	Arten von Fehlern	48
7.1.3	Testentscheidungen und Konsequenzen	49
7.1.3.1	p -Wert und Konfidenzintervall	50
7.1.3.2	Multiples Testen	53
8	Vorlesung am 6. Mai 2010	54
8.1	Lagetests	54
8.1.1	t-Tests	54
8.1.1.1	t-Test für eine Stichprobe	54

Inhaltsverzeichnis

8.1.1.2	t-Test für zwei verbundene Stichproben	54
8.1.1.3	t-Test für zwei unverbundene Stichproben	55
8.1.1.4	Welch-Test	57
8.1.1.5	Vorraussetzungen für t-Tests	57
8.1.1.6	Andere Anwendungen des t-Tests	58
8.1.2	Rangsummentests	58
8.1.2.1	Wilcoxon-Test für eine Stichprobe	58
8.1.2.2	Wilcoxon-Test für zwei verbundene Stichproben	59
9	Vorlesung am 20. Mai 2010	61
9.0.2.3	U-Test (Wilcoxon-Test für unverbundene Stichproben)	61
9.0.2.4	Vergleich von Rangsummentests und t-Tests	62
9.0.3	Vorzeichentest	62
9.0.3.1	Vorzeichentest für eine Stichprobe	62
9.0.3.2	Vorzeichentest für zwei verbundene Stichproben	62
9.1	Tests zum Vergleich von Häufigkeiten	64
9.1.1	Binomialtest für eine Stichprobe	64
9.1.2	χ^2 -Tests	64
9.1.2.1	χ^2 -Vierfelder-Test	64
9.1.2.2	χ^2 -Unabhängigkeits-Test	65
9.1.2.3	χ^2 -Homogenitäts-Test	65
10	Vorlesung am 27. Mai 2010	69
10.0.2.4	χ^2 -Test für $k \cdot l$ Felder	69
10.0.2.5	Assiotiationsmaße für qualitative Merkmale	69
10.0.2.6	McNemar Test	70
10.0.2.7	Mittelwert	71
11	Vorlesung am 11. Juni 2010	72
11.0.3	Korrelationsanalyse	72
11.0.3.1	Punktwolke	72
11.0.3.2	Kovarianz	72
11.0.4	Pearson Koeffizient	73
11.0.4.1	Interpretation des Korrelationskoeffizienten	73
11.0.5	Regressionsanalyse	74
11.0.6	Spearman'scher Korrelationskoeffizient	74
11.0.7	Klassifikation nach formalen Aspekten	74
11.0.7.1	Deskriptiv vs. Analytisch	74
11.0.7.2	Transversial vs. Longitudinal	75
11.0.7.3	Retrospektiv vs. Prospektiv	75
11.0.7.4	Beobachtend vs. Experimentell	75
11.0.7.5	Monozentrisch vs. Multizentrisch	75
11.0.7.6	Typische Studien	76

Inhaltsverzeichnis

11.0.8	Fehlerquellen und -möglichkeiten	76
11.0.8.1	Zufällige Fehler	76
11.0.8.2	Systematische Fehler	76
11.0.9	Fall-Kontroll-Studien	76
11.0.9.1	Matching	76
11.0.9.2	Biasquellen	76
11.0.9.3	κ -Koeffizient nach Cohen	76

1 Vorlesung am 4. März 2010

1.1 Überlebenszeiten und -funktionen

1.1.1 Bevölkerungsstatistik und Sterbetafeln

Abschnitt 6.4
(S.114 – 117)

$$\mathbb{E} X = \sum_i x_i p_i \dots \text{Erwartungswert, Lebenserwartung}$$

1.1.1.1 Sterbetafel¹

Abschnitt 6.4.2

- $l_o \dots$ nominelle Anzahl von Lebendgeborenen
- $l_x \dots$ Anzahl der Personen, die den x -ten Geburtstag erleben
- $d_x := l_x - l_{x+1}$
- $x :=$ Anzahl der vollendeten Lebensjahre

Mit vollendetem Lebensalter ist das ganzzahlige Alter – die Anzahl der vollendeten Lebensjahre – gemeint, dies ist ein diskretes Merkmal. Das kontinuierliche Alter hingegen ist stetig.

$$x \in \mathbb{N}_0 \quad \text{wobei} \quad x := \lfloor \text{kontinuierliches Alter} \rfloor$$

Die Wahrscheinlichkeit, wenn ein Individuum den x -ten Geburtstag erlebt hat, im darauffolgenden Jahr stirbt:

$$q_x = \frac{d_x}{l_x}$$

Für die letzte Zeile, wird ω als Alter x angenommen, meistens ist das maximal berücksichtigte Alter $\omega = 100$, weiters wird angenommen, dass

$$l_{\omega+1} = 0 \quad q_\omega = 1 \quad e_\omega = \frac{1}{2}$$

1.1.1.2 Lebenserwartung

Für die Berechnung der Lebenserwartung e_0 (eines Neugeborenen)

¹Lifetable

Zufallsvariable

X ... Anzahl der Lebensjahre, die beim Tod *vollendet* sind

Wahrscheinlichkeit

$P(X = x) = \frac{d_x}{l_0}$... Wahrscheinlichkeit im ganzzahligen Alter x zu sterben

Herleitung von e_0

Formel 6.16

$$\begin{aligned}
 e_0 &= \frac{1}{2} + \text{Anzahl der erwarteten vollendeten Lebensjahre} \\
 &= \frac{1}{2} + \left[0 \cdot \frac{d_0}{l_0} + 1 \cdot \frac{d_1}{l_0} + 2 \cdot \frac{d_2}{l_0} + \dots + (\omega - 1) \cdot \frac{d_{\omega-1}}{l_0} + \omega \cdot \frac{d_\omega}{l_0} \right] \\
 &= \frac{1}{2} + \frac{1}{l_0} [0 \cdot (l_0 - l_1) + 1 \cdot (l_1 - l_2) + \dots + (\omega - 1) \cdot (l_{\omega-1} - l_\omega) + \omega \cdot (l_\omega - l_{\omega+1})] \\
 &= \frac{1}{2} + \frac{1}{l_0} [0 \cdot l_0 + l_1 + l_2 + \dots + l_\omega + \omega \cdot 0] \\
 &= \frac{1}{2} + \frac{1}{l_0} \cdot \sum_{x=1}^{\omega} l_x
 \end{aligned}$$

e_x ist die fernere Lebenserwartung (die Restliche Lebensdauer) eines x -Jährigen.

$$e_x = \frac{1}{2} + \frac{1}{l_x} \cdot \sum_{y=x+1}^{\omega} l_y \quad \text{für } x \in [0 \dots \omega]$$

Das erwartete Sterbealter ist dann $x + e_x$

Formel 6.18

$$F(X) = 1 - \frac{l_x}{l_0} \quad \text{Verteilungsfunktion}$$

1.1.2 Verteilung von Überlebenszeiten

Abschnitt 8.4

1.1.2.1 Überlebenszeit

(S. 162 – 166)

Abschnitt 8.4.1

Die Überlebenszeit ist die Dauer vom Anfangsereignis (Geburt) bis zum Endereignis (Tod)

X (oder T) ... Überlebenszeit (Überlebensdauer)
 $F(t) = P(X \leq t)$
 $S(t) = P(X > t) = 1 - F(t)$
 $S(t)$... Überlebensfunktion (survival function)
 $f(t) = F'(t)$... Dichtefunktion
 $r(t) = \frac{f(t)}{S(t)}$... Momentane Sterberate, Hazard-Funktion, Hazard-Rate, Ausfallrate

Bedingte Überlebenswahrscheinlichkeit²

$$\begin{aligned}
 P(X > t + \Delta t \mid X > t) &= \frac{P(X > t + \Delta t \wedge X > t)}{P(X > t)} \\
 \frac{P(X > t + \Delta t)}{P(X > t)} &= \frac{S(t + \Delta t)}{S(t)}
 \end{aligned}$$

Bedingte Sterbewahrscheinlichkeit

$$\begin{aligned}
 P(t < X \leq t + \Delta t \mid X > t) &= \frac{P(t < X \leq t + \Delta t \wedge X > t)}{P(X > t)} = \\
 &= \frac{P(t < X \leq t + \Delta t)}{P(X > t)} = \frac{F(t + \Delta t) - F(t)}{S(t)}
 \end{aligned}$$

Momentane Sterberate Die Momentane Sterberate (Ausfallrate) ergibt sich aus der bedingten Sterbewahrscheinlichkeit folgendermaßen:

$$\begin{aligned}
 r &:= \lim_{\Delta t \rightarrow 0} \frac{P(t < X \leq t + \Delta t \mid X > t)}{\Delta t} = \\
 &= \lim_{\Delta t \rightarrow 0} \frac{F(t + \Delta t) - F(t)}{\Delta t} \cdot \frac{1}{S(t)} = \\
 &= F'(t) \cdot \frac{1}{S(t)} = \frac{f(t)}{S(t)}
 \end{aligned}$$

Zur Interpretation von $r(t)$

In linearer Approximation (für hinreichend kleines Δt) gilt bekanntlich folgendes:

$$F(t + \Delta t) \approx F(t) + F'(t) \cdot \Delta t$$

²Nach dem Bayestheorem für bedingte Wahrscheinlichkeiten $P(A \mid B) = \frac{P(A \cap B)}{P(B)}$

Daher gilt für die Wahrscheinlichkeit, dass ein Individuum, nachdem es das Alter t erlebt hat im folgenden Zeitintervall der Δt stirbt:

$$\begin{aligned} P(t < X \leq t + \Delta t \mid X > t) &= \frac{F(t + \Delta t) - F(t)}{S(t)} \approx \\ &\approx \frac{F(t) + F'(t)\Delta t - F(t)}{S(t)} = \frac{F'(t)\Delta t}{S(t)} = \\ &= \frac{f(t)\Delta t}{S(t)} = r(t)\Delta t \end{aligned}$$

Die Wahrscheinlichkeit ist also (in linearer Näherung) das Produkt aus momentaner Sterberate und der Länge des Zeitintervalls, somit proportional zu Δt

Mit $\Delta t := 1$ (Einheitszeit) ergibt sich daraus

$$r(t) \approx P(t < X \leq t + 1 \mid X > t)$$

Das bedeutet, dass die Ausfallrate $r(t)$ folgendermaßen interpretiert werden kann: (näherungsweise in linearer Approximation) $r(t)$ ist die Wahrscheinlichkeit, dass ein Individuum, das das Alter t erlebt hat während der folgenden (kleinen) Zeiteinheit stirbt.

Die momentane Sterberate (Ausfallrate) ist charakterisiert durch

$$r(t) = \frac{f(t)}{S(t)}$$

wobei

$$\begin{aligned} f(t) &= F'(t) \\ S(t) &= 1 - F(t) \end{aligned}$$

Daraus folgt

$$f(t) = -S'(t)$$

und wegen

$$\begin{aligned} F(0) &= 0 & \text{bzw.} & & S(0) &= 1 \\ F(t) &= \int_0^t f(\tau) d\tau & \text{bzw.} & & S(t) &= 1 - \int_0^t f(\tau) d\tau \end{aligned}$$

Fragen

1. Wie kann $r(t)$ nur durch $S(t)$ oder $F(t)$ oder $f(t)$ ausgedrückt werden?

$$\begin{aligned} r(t) &= \frac{-S'(t)}{S(t)} \\ r(t) &= \frac{F'(t)}{-F(t)} \\ r(t) &= \frac{f(t)}{1 - \int_0^t f(\tau) d\tau} \end{aligned}$$

2. Wie können $S(t)$ bzw. $F(t)$ bzw. $f(t)$ nur durch $r(t)$ ausgedrückt werden?
Mit der Notation τ anstelle von t ist

$$r(\tau) = \frac{S'(\tau)}{S(\tau)}$$

Unter der Voraussetzung $S(t) > 0$ gilt auch $S(\tau) > 0$ für $0 \leq \tau \leq t$, daher gilt

$$\begin{aligned} - \int_0^t r(\tau) d\tau &= - \int_0^t \frac{S'(\tau)}{S(\tau)} d\tau \\ \ln S(\tau) \Big|_0^t &= \ln S(t) - \ln S(0) \\ \ln S(t) - 0 &= \ln S(t) \end{aligned}$$

Somit ist

$$\begin{aligned} S(t) &= e^{-\int_0^t r(\tau) d\tau} \\ F(t) &= 1 - e^{-\int_0^t r(\tau) d\tau} \\ f(t) &= r(t) \cdot e^{-\int_0^t r(\tau) d\tau} \end{aligned}$$

Letzteres ergibt sich aus

$$f(t) = F'(t)$$

oder aus

$$r(t) = \frac{f(t)}{S(t)} \iff f(t) = r(t) \cdot S(t)$$

2 Vorlesung am 11. März 2010

2.0.3 Verteilung von Überlebenszeiten

2.0.3.1 Exponentialverteilung

Abschnitt 8.4.2

Im Spezialfall wird die Exponentialverteilung

$$\text{Exp}(\lambda)$$

Die Sterberate ist konstant

$$r(t) = \lambda$$

Also handelt es sich bei der Exponentialverteilung um eine gedächtnislose Verteilung, da die Änderung unabhängig von der Zeit ist.

$$\int_0^t r(\tau) d\tau = \int_0^t \lambda d\tau = \lambda t$$

Damit ergeben sich auch auf diesem Weg die für die Exponentialverteilung bekannten Formeln

$$\begin{aligned} S(t) &= e^{-\lambda t} \\ F(t) &= 1 - e^{-\lambda t} \\ f(t) &= \lambda e^{-\lambda t} \end{aligned}$$

Der Erwartungswert und gefolgt von der Varianz, welche sich mit dem Verschiebesatz¹ errechnen lässt.

$$\begin{aligned} \mu &= \frac{1}{\lambda} \\ \sigma^2 &= \frac{1}{\lambda^2} \end{aligned}$$

¹ $\mathbb{E}(X - \mathbb{E} X)^2 = \mathbb{E} X^2 - (\mathbb{E} X)^2$

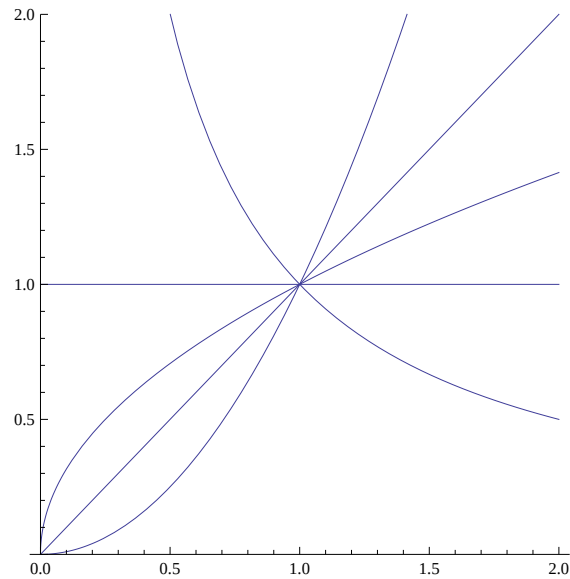


Abbildung 2.0.1: Das Verhalten der Potenzfunktion $f(x) = x^a$ für $a \in \{-1, 0, \frac{1}{2}, 1, 2\}$

Die Schiefe² γ_1 ist immer 2.

2.0.3.2 Weibull Verteilung

Abschnitt 8.4.3

Die Weibullverteilung ist eine Verallgemeinerung der Exponentialverteilung

$$\text{WB}(\lambda, \gamma)$$

Es gelten dabei folgende Formeln

$$\begin{aligned} S(t) &= e^{-\lambda t^\gamma} \\ F(t) &= 1 - e^{-\lambda t^\gamma} \\ f(t) &= \lambda \gamma t^{\gamma-1} e^{-\lambda t^\gamma} \\ r(t) &= \frac{f(t)}{S(t)} = \lambda \gamma t^{\gamma-1} \end{aligned}$$

Die Exponentialverteilung kann als Ausprägung mit dem Parameter $\gamma = 1$ angesehen werden.

$$\text{WB}(\lambda, 1) = \text{Exp}(\lambda)$$

²Der Vollständigkeit: γ_2 gibt die Wölbung an.

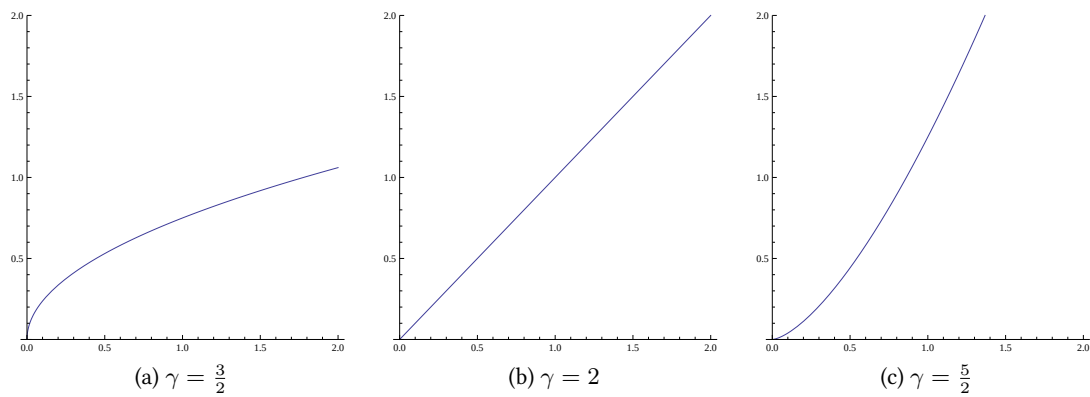


Abbildung 2.0.2: Die Änderungsrate für die Weibull Verteilung $r(t) = \lambda \gamma t^{\gamma-1}$ mit Parameter $\lambda = \frac{1}{2}$ im Bereich $0 < t < 2$ verhält sich wie die Potenzfunktion. Die Graphiken zeigen $\gamma \in \{\frac{3}{2}, 2, \frac{5}{2}\}$

2.0.4 Prognose Studien

2.0.4.1 Kaplan-Meier-Methode

Abschnitt 16.2.3
(S. 307 – 309)

Es wird für die unbekannte Funktion $S(t)$ eine geschätzte Funktion $\hat{S}(t)$ gesucht.

n	...	Anzahl der Personen in der Studie
$t_1 < t_2 < \dots < t_k$	...	Zeitpunkte der kritischen Endereignisse
$t_0 := 0$	...	Beginn des Beobachtungszeitraumes
$t_{k+1} := t_\omega$	...	Ende des Beobachtungszeitraumes
c_i	...	Anzahl der zensierten Beobachtungen in $[t_{i-1}, t_i)$
d_i	...	Anzahl der Endereignisse zum Zeitpunkt t_i
n_i	...	Anzahl der Personen die unmittelbar vor t_i beobachtet werden

Daher ist

$$n_i = n_{i+1} - d_{i-1} - c_i \quad \text{für } 1, \dots, k$$

wobei definiert ist, dass

$$\begin{aligned} n_0 &:= n \\ d_0 &:= 0 \\ n_1 &= n - c_1 \end{aligned}$$

Daher gilt

$$\begin{aligned}\hat{S}(0) &= 1 \\ \hat{S}(t_i) &= \hat{S}(t_{i-1}) \cdot \frac{n_i - d_i}{n_i} \quad \text{für } i = 1, 2, \dots, k\end{aligned}$$

Beispiel 16.1

Example 1. Nach einer Organtransplantation wurden bei 10 Patienten die Überlebenszeiten in Tagen ermittelt. Nach 160 Tagen wurde die Studie beendet. Bei 7 Patienten konnte der Zeitpunkt des Endereignisses ermittelt werden (nach 20, 35, 62, 91, 91, 128 und 148 Tagen). Ein Patient brach nach 98 Tagen die Studie ab; zwei Patienten lebten am Ende der Studie noch.

Aus dem Text schließt man $t_\omega = 160$ und $n = 10$. Desweiteren gibt es folgende Zeitpunkte mit kritischen Ereignissen: $t_1 = 20$, $t_2 = 35$, $t_3 = 62$, $t_4 = 92$, $t_5 = 128$ und $t_6 = 148$. Für alle Zeitpunkte gilt $d_i = 1$ außer für t_4 wo $d_4 = 2$ gilt, da an diesem Tag zwei Patienten verstorben sind.

t_i	c_i	n_i	d_i	$n_i - d_i$	$\hat{S}(t_i)$
0				10	1
20	0	10	1	9	$1 \cdot \frac{9}{10} = 0.9$
35	0	9	1	8	$0.9 \cdot \frac{8}{9} = 0.8$
62	0	8	1	7	$0.8 \cdot \frac{7}{8} = 0.7$
91	0	7	2	5	$0.7 \cdot \frac{5}{7} = 0.5$
128	1	4	1	3	$0.5 \cdot \frac{3}{4} = 0.375$
148	0	3	1	2	$0.375 \cdot \frac{2}{3} = 0.25$
160	0	2			

2.0.4.2 Logrank-Test

Abschnitt 12.2.7
(S. 242 – 243)

- Zum Vergleich (von Verteilungen) von Überlebenszeiten.

$$S_1(t) \quad S_2(t)$$

- bei zwei unverbundenen Stichproben

- mit folgenden Hypothesen

$$\begin{array}{ll} \text{Null Hypothese} & N_0 : S_1(x) \equiv S_2(x) \quad \forall x : S_1(x) = S_2(x) \\ \text{Alternativ Hypothese} & N_1 : S_1(x) \not\equiv S_2(x) \quad \exists x : S_1(x) \neq S_2(x) \end{array}$$

Stichproben sind $j = 1, 2$ (manchmal auch $j = A, B$)

Die Zeitpunkte sind mit t_i bezeichnet, wobei $i = 1, 2, \dots, k$ sowie $t_0 := 0$ gilt.

$$\begin{aligned} d_{j,i} & \dots \text{Anzahl der Endereignisse bei der Stichprobe } j \text{ zum Zeitpunkt } i \\ d_j &:= \sum_{i=1}^k d_{j,i} \quad j = 1, 2 \\ c_{j,i} & \dots \text{Anzahl zensierter Daten in Stichprobe } j \text{ zum Zeitpunkt } [t_{i-1}, t_i) \\ n_{j,0} &:= n_j \quad \dots \text{Anzahl der Beobachtungseinheiten bei Stichprobe } j \\ n_{j,i} &:= n_{j,i-1} - d_{j,i-1} - c_{j,i} \quad \text{bzw.} \quad n_{j,i} - d_{j,i} - c_{j,i+1} =: n_{j,i+1} \\ e_j &:= \sum_{i=1}^k (d_{1,i} + d_{2,i}) \cdot \frac{n_{j,i}}{n_{1,i} + n_{2,i}} \quad \text{für } j = 1, 2 \\ e_1 + e_2 &= \sum_{i=1}^k (d_{1,i} + d_{2,i}) \cdot \frac{n_{1,i}}{n_{1,i} + n_{2,i}} + \sum_{i=1}^k (d_{1,i} + d_{2,i}) \cdot \frac{n_{2,i}}{n_{1,i} + n_{2,i}} = \\ &= \sum_{i=1}^k (d_{1,i} + d_{2,i}) \cdot \left[\frac{n_{1,i} + n_{2,i}}{n_{1,i} + n_{2,i}} \right] = d_1 + d_2 \\ e_2 &= d_1 + d_2 - e_1 \end{aligned}$$

Die zugehörige χ^2 -Verteilung (mit einem Freiheitsgrad) ist somit definiert als

$$\chi^2 := \frac{(d_1 - e_1)^2}{e_1} + \frac{(d_2 - e_2)^2}{e_2}$$

Die Nullhypothese $H_0 : S_1(x) \equiv S_2(x)$ wird abgelehnt, falls $\chi^2 > \chi_{1,1-\alpha}^2$

Example 2. Gegeben sind zwei Studien $j = A, B$ mit $t_\omega = 160$

Therapie A $n_A = 10$

Zeiten der Endereignisse: 20, 35, 62, 91, 91, 128, 148

Zeiten mit Abbruch: 98

Therapie B $n_B = 15$

Zeiten der Endereignisse: 1, 11, 3, 3, 20, 91

Zeiten mit Abbruch: 2, 3, 15, 29, 29, 120

2 Vorlesung am 11. März 2010

t_i	$c_{A,i}$	$n_{A,i}$	$d_{A,i}$	$n_{A,i} \cdot d_{A,i}$	$\hat{S}_A(t_i)$
	$c_{B,i}$	$n_{B,i}$	$d_{B,i}$	$n_{B,i} \cdot d_{B,i}$	$\hat{S}_B(t_i)$
0				10	
				15	
1	0	10	0	10	
	0	15	3	12	
3	0	10	0	10	
	1	11	2	9	
20	0	10	1	9	
	2	7	1	8	
35	0	9	1	8	
	2	4	0	4	
62	0	8	1	7	
	0	4	0	4	
91	0	7	2	5	
	0	4	1	3	
128	1	4	1	3	
	1	2	0	2	
148	0	3	1	2	
	0	2	0	2	
160	0	2	0		
	0	2	0		

Es ergibt sich aus der Tabelle als Summe

$$d_A = 7 \quad \text{und} \quad d_B = 7$$

$$\begin{aligned}e_A &= 3 \cdot \frac{10}{25} + 2 \cdot \frac{10}{21} + 2 \cdot \frac{10}{17} + 1 \cdot \frac{9}{13} + 1 \cdot \frac{8}{12} + 3 \cdot \frac{7}{11} + 1 \cdot \frac{4}{6} + 1 \cdot \frac{3}{5} = \\&= 7.863\,583\,475\end{aligned}$$

$$e_B = 14 - e_A = 6.136\,416\,525$$

$$\chi^2 = 0.216\,372\,145$$

χ^2 Tabelle
S. 322

$\alpha = 10\%$	$\chi^2_{1,0.9} = 2.706 > \chi^2$	$\Rightarrow$ Testentscheidung H_0
$\alpha = 5\%$	$\chi^2_{1,0.95} = 3.841 > \chi^2$	$\Rightarrow$ Testentscheidung H_0
$\alpha = 1\%$	$\chi^2_{1,0.99} = 6.635 > \chi^2$	$\Rightarrow$ Testentscheidung H_0

3 Vorlesung am 18. März 2010

3.1 Wahrscheinlichkeiten in der Epidemiologie

Abschnitt 6.3
(S. 111)

Prävalenz genau gesagt, die Punktprävalenz gibt zu einem bestimmten Zeitpunkt folgendes an:

- Den relativen Krankenbestand
- Den Anteil der Erkrankten (einer bestimmten Krankheit) in der Population
- Die Wahrscheinlichkeit für eine beliebige Person *erkrankt zu sein*

Um die Prävalenz zu ermitteln werden oft Querschnittstudien verwendet.

Periodenprävalenz gibt an am Anfang, während, oder am Ende des Beobachtungszeitraumes erkrankt zu sein.

Für den Rest gilt dasselbe wie für die Prävalenz oben.

Inzidenz

- Die Wahrscheinlichkeit für eine beliebige Person während des Beobachtungszeitraumes neu zu *erkranken*.
- Falls der Beobachtungszeitraum eine Zeiteinheit (ein Jahr, ein Monat, ...) ist spricht man von der Neuerkrankungsrate.

$$\text{Prävalenz} = \text{Inzidenz} \cdot \text{durchschnittliche Krankheitsdauer}$$

Die Prävalenz wird mit $P(K_T)$ bezeichnet und die Inzidenz mit $P(K)$ in der Epidemiologie jedoch herrscht die Bezeichnung $P(K)$ für die Prävalenz vor. Diese Bezeichnung wird auch im Buch und in Folge verwendet.

(krankheitspezifische) Mortalität auch „Todesrate“ genannt

- $P(K \cap T)$
- Die Wahrscheinlichkeit (während des Beobachtungszeitraumes) an der Krankheit zu erkranken *und* zu sterben.

Letalität

- $P(T | K)$
- Die Wahrscheinlichkeit (während des Beobachtungszeitraumes) zu sterben, falls das Individuum an der Krankheit K erkrankt ist.

$$\bullet P(K \cap T) = P(K) \cdot P(T | K)$$

$$\text{Mortalität} = \text{Inzidenz} \cdot \text{Letalität}$$

Kontragonindex

Manifestationsindex

3.1.1 Bevölkerungsstatistik

Abschnitt 6.4
(S. 114)

Natalität (auch: Geburtenziffer)

$$\frac{\| \text{Geborene Kinder pro Jahr} \|}{\text{Populationsgröße}}$$

Fertilitätsrate

$$\frac{\| \text{Lebend Geborene pro Jahr} \|}{\| \text{Frauen im gebärfähigen Alter} \|}$$

Sterbeziffer

$$\frac{\| \text{Gestorbene pro Jahr} \|}{\text{Populationsgröße}}$$

Pearl-Index

Anzahl der Frauen – von 100 Frauen, welche eine bestimmte Verhütungsmethode anwenden – die in einem Jahr ungewollt schwanger werden.

3.1.2 Diagnostische Tests

Abschnitt 6.5
(S. 119)

K	...	erkranktes Individuum
$\overline{K}$	...	gesundes Individuum
T_+	...	positives Testergebnis
T_-	...	negatives Testergebnis
$P(T_+ K)$	...	Sensitivität
$P(T_- \overline{K})$	...	Spezivität

3.1.2.1 Vorhersagewerte

$P(K T_+)$	...	positive Vorhersagewahrscheinlichkeit
$P(\overline{K} T_-)$	...	negative Vorhersagewahrscheinlichkeit

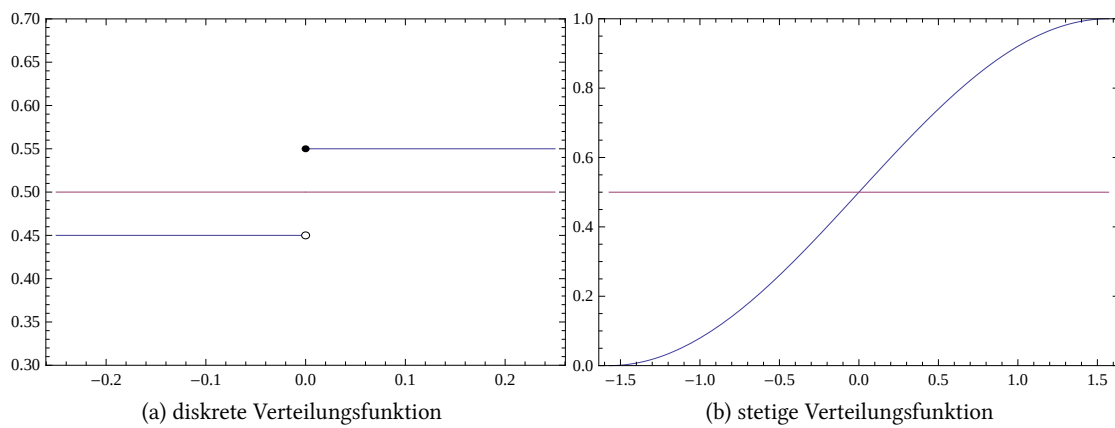


Abbildung 3.2.1: Median bei diskreten und stetigen Verteilungsfunktionen

$$\begin{aligned}
 P(K | T_-) &= \frac{P(K \cap T_-)}{P(T_-)} = \frac{P(K \cap T_-)}{P(K \cap T_-) + P(\bar{K} \cap T_-)} = \\
 &= \frac{P(K) \cdot P(T_- | K)}{P(K) \cdot P(T_- | K) + P(\bar{K}) \cdot P(T_- | \bar{K})} = \\
 &= \frac{P(K) \cdot [1 - P(T_+ | K)]}{P(K) \cdot [1 - P(T_+ | K)] + [1 - P(K)] \cdot P(T_- | \bar{K})}
 \end{aligned}$$

3.2 Diskrete Verteilungen

3.2.1 Diskrete Zufallsvariablen

Abschnitt 7.1
(S. 125)

Bezeichner

X	...	Zufallsvariable
x_i	...	Wert (Ausprägung)
p_i	...	Wahrscheinlichkeit für Wert x_i
$f(x_i) = p_i$	...	Dichtefunktion (Wahrscheinlichkeitsfunktion)

Erwartungswert

$$\begin{aligned}\text{Verteilungsfunktion } F(x) &= P(X \leq x) \\ \text{Erwartungswert } \mathbb{E} X &= \mu \\ \mathbb{E}(a \cdot X + b) &= a \cdot \mathbb{E} X + b \\ \mathbb{E} \left(\sum_i X_i \right) &= \sum_i \mathbb{E} X_i\end{aligned}$$

Median und Quantill

$$\begin{aligned}\text{Median } \tilde{\mu} &= \min \{x \mid F(x) \geq 0.5\} \\ \alpha\text{-Quantill } \tilde{\mu}_\alpha &= \min \{x \mid F(x) \geq \alpha\}\end{aligned}$$

$\tilde{\mu}$ wird immer als 0.5-Quantill bezeichnet: $\tilde{\mu} = \tilde{\mu}_{0.5}$

Varianz und Standardabweichung

$$\begin{aligned}\text{Var } X = \sigma^2 &= \mathbb{E} \left([X - \mu]^2 \right) = \mathbb{E} X^2 - (\mathbb{E} X)^2 \\ &= \sum_i x_i^2 p_i - \left(\sum_i x_i p_i \right)^2 \\ \text{Standardabweichung } \sigma &= \sqrt{\text{Var } X} \\ \text{Variationskoeffizient} &= \frac{\sigma}{\mu}\end{aligned}$$

Für das Rechnen mit Varianzen gilt,

$$\begin{aligned}\text{Var}(a \cdot X + b) &= a^2 \cdot \text{Var } X \\ \text{Var}(X + Y) &= \text{Var } X + \text{Var } Y + 2 \cdot \text{Cov}(X, Y)\end{aligned}$$

Wenn $X_1, \dots, X_n$ unabhängige Zufallsvariablen sind, so gilt,

$$\text{Var} \left(\sum_{i=1}^n X_i \right) = \sum_{i=1}^n \text{Var } X_i$$

Kovarianz Wenn X, Y unabhängig¹ sind, so gilt für die Kovarianz

$$\text{Cov}(X, Y) = 0$$

¹Eigentlich muss es unkorreliert heißen, da es exotische Ausnahmen geben kann, wo dies für unabhängige Zufallsvariablen nicht gilt.

3.2.2 Binomialverteilung

Abschnitt 7.2

3.2.2.1 Bernoulli-Experiment

Abschnitt 7.2.1

Das Bernoulli-Experiment geht davon aus, dass jeder durchgeführte Versuch zwei Ausgänge (Ereignisse) haben kann (Diese sind hier mit A und $\bar{A}$ bezeichnet).

$$\begin{aligned} P(A) &= P(X = 1) = p \\ P(\bar{A}) &= P(X = 0) = q = 1 - p \end{aligned}$$

3.2.2.2 Bernoulli-Prozess

Abschnitt 7.2.2

$$\begin{aligned} X &\sim B(n, p) \\ \mathbb{E} X &= n \cdot p \\ \text{Var } X &= n \cdot p \cdot q \\ P(X = k) &= \binom{n}{k} \cdot p^k \cdot q^{n-k} \quad \text{für } k = 0, \dots, n \\ &= \frac{n!}{k! \cdot (n-k)!} \cdot p^k \cdot q^{n-k} \\ \text{Schiefe } \gamma_1 &= \frac{q-p}{\sigma} \quad \text{wenn } \begin{cases} > 0 & \text{rechtsschief} \\ < 0 & \text{linksschief} \end{cases} \end{aligned}$$

Symmetrische Binomialverteilung

$$\begin{aligned} p &= q = \frac{1}{2} \\ \mathbb{E} X &= \frac{n}{2} \\ \text{Var } X &= \frac{n}{4} \end{aligned}$$

3.2.3 Poissonverteilung

Abschnitt 7.3.1

Die Poissonverteilung ist ein Spezialfall der Binomialverteilung, für die folgendes gilt:

$$\begin{aligned} \mathbb{E} X &= n \cdot p = \lambda \text{ (const)} \\ n &\rightarrow \infty \end{aligned}$$

In der Praxis wird die Poissonverteilung als Annäherung für die Binomialverteilung verwendet:

$$B(n, p) \rightarrow P(\lambda) \quad \text{wenn } n \geq 30 \wedge p \leq 0.1$$

Es gilt also

$$X \sim P(\lambda)$$

unter der Bedingung,

$$\begin{aligned} \mathbb{E} X &= \lim_{n \rightarrow \infty} n \cdot p = \lim_{n \rightarrow \infty} \lambda = \lambda \\ \text{Var } X &= \lim_{n \rightarrow \infty} n \cdot p \cdot (1 - p) = \lim_{n \rightarrow \infty} \lambda \cdot \left(1 - \frac{1}{n}\right) = \lambda \\ \gamma_1 &= \lim_{n \rightarrow 0} \frac{(1 - p)}{\sqrt{n \cdot p \cdot (1 - p)}} = \lim_{n \rightarrow 0} \frac{1 - 2\frac{\lambda}{n}}{\sqrt{n \cdot p \cdot (1 - p)}} = \frac{1}{\sqrt{\lambda}} > 0 \end{aligned}$$

Die Poissonverteilung definiert sich also als

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

4 Vorlesung am 25. März 2010

Beispiel 7.5
(S. 137)

Example 3. In einer Geburtsklinik werden jährlich $n = 2\,000$ Kinder geboren. Die Wahrscheinlichkeit, dass ein Neugeborenes mit einem Down-Syndrom zur Welt kommt, beträgt $p = 0.001$. Unter der Annahme, dass die Ereignisse unabhängig sind, lässt sich die Anzahl der Neugeborenen mit Down-Syndrom durch eine Poisson-verteilte Zufallsvariable X beschreiben. Für die charakteristischen Parameter gilt: $\lambda = n \cdot p = 2\,000 \cdot 0.001 = 2$.

Zur Approximation der Binomialverteilung durch die Poissonverteilung. Wenn X Binomial verteilt ist,

$$X \sim B\left(n, \frac{\lambda}{n}\right) \quad \text{wobei} \quad \frac{\lambda}{n} = p \text{ (const)}$$

Aus der Definition der Binomialverteilung ergibt sich:

$$\begin{aligned} P(X = k) &= \binom{n}{k} \cdot p^k \cdot (1-p)^{n-k} = \frac{n!}{k! \cdot (n-k)!} \cdot \left(\frac{\lambda}{n}\right)^k \cdot \left(1 - \frac{\lambda}{n}\right)^{n-k} = \\ &= \frac{1}{k!} \cdot \underbrace{\frac{n \cdot (n-1) \cdots (n-k+1)}{n^k}}_A \cdot \underbrace{\lambda^k}_{B} \cdot \underbrace{\left(1 - \frac{\lambda}{n}\right)^n \cdot \left(1 - \frac{\lambda}{n}\right)^{-k}}_C = \end{aligned}$$

Diese Gleichung lässt sich vereinfachen, da der Grenzwert einiger Teilterme 1 ist. Der Übersicht genüge wird dies in mehreren Schritten dargestellt.

$$\begin{aligned} A \quad \lim_{n \rightarrow \infty} \frac{n \cdot (n-1) \cdots (n-k+1)}{n^k} &= \\ &= \lim_{n \rightarrow \infty} \underbrace{1 \cdot \overbrace{\left(1 - \frac{1}{n}\right)}^{\rightarrow 1} \cdots \overbrace{\left(1 - \frac{k-1}{n}\right)}^{\rightarrow 1}}_{k\text{-mal}} = 1 \end{aligned}$$

Da es im letzten Term nur k Faktoren von annähern 1 gibt, ergibt sich der Gesamtgrenzwert von 1.

$$\begin{aligned} \text{B} \quad \lim_{n \rightarrow \infty} \left(1 - \frac{\lambda}{n}\right)^n &= \\ &= \lim_{n \rightarrow \infty} \left(1 + \frac{-\lambda}{n}\right)^n = e^{-\lambda} \end{aligned}$$

$$\text{C} \quad \lim_{n \rightarrow \infty} \left(1 - \frac{\lambda}{n}\right)^{-k} = 1^{-k} = 1$$

Folglich ist

$$\lim_{n \rightarrow \infty} P(X = k) = \frac{1}{k!} \cdot 1 \cdot \lambda^k \cdot e^{-\lambda} \cdot 1 = \frac{1}{k!} \lambda^k e^{-\lambda}$$

Die Binomialverteilung kann also durch die Poissonverteilung für hinreichend große n und konstante p approximiert werden.

$$B(n, p) \xrightarrow[n \rightarrow \infty]{\substack{\lambda \\ = p \text{ const}}} P(\lambda)$$

4.0.4 Polynomialverteilung

Abschnitt 7.3.2

Bei der Polynomialverteilung (oder auch Multinomialverteilung) werden n Zufallsexperimente mit k möglichen Ereignissen $A_1, A_2, \dots, A_k$ mit den zugehörigen Wahrscheinlichkeiten $p_1, p_2, \dots, p_k$ und Häufigkeiten $n_1, n_2, \dots, n_k$.

$$P(n_1, n_2, \dots, n_k \mid p_1, p_2, \dots, p_k) = \frac{n!}{n_1! \dots n_k!} \cdot p_1^{n_1} \cdot p_2^{n_2} \dots p_k^{n_k}$$

Wobei gilt, dass

$$\sum_{i=1}^n n_i = n \quad \wedge \quad \sum_{i=1}^n p_i = 1 \quad \wedge \quad \forall p_i \geq 0 \quad \wedge \quad \forall n_i \in \{0, 1, 2, \dots, n\}$$

4.0.5 Negative Binomialverteilung

Abschnitt 7.3.3

Die Negative Binomialverteilung beschreibt die Wahrscheinlichkeit, dass bei einem Zufallsexperiment mit zwei Ausgängen ein Ereignis A – mit der Wahrscheinlichkeit p – bei der j -ten Beobachtung zum r -ten Mal eintritt.

$$X \sim \text{NB}(r, p)$$

4.0.5.1 Geometrische Verteilung

Die geometrische Verteilung ist ein Sonderfall der negativen Binomialverteilung mit $r = 1$.

$$X \sim \text{NB}(1, p)$$

$P(X = j)$ ist die Wahrscheinlichkeit, dass das Ereignis bei der j -ten Beobachtung erstmals Eintritt.

$$P(X = j) = q^{j-1}p \quad \text{wobei } q = 1 - p \quad \wedge \quad j = 1, 2, 3, \dots$$

4.0.5.2 Definition der negativen Binomialverteilung

$$P(X = j) = \binom{j-1}{r-1} p^{r-1} q^{j-r} p \quad \text{für } j = r, r+1, \dots$$

Example 4. Eine Blutbank...

Beispiel 7.8
(S. 139)

4.0.6 Hypergeometrische Verteilung

Abschnitt 7.3.4

Bei der Hypergeometrischen Verteilung nimmt man an, dass es bei einer Menge von N Objekten insgesamt M Objekte gibt, die eine gesuchte Eigenschaft A besitzen. Der Rest hat diese Eigenschaft nicht ($\bar{A}$). Der Anteilswert $p = \frac{M}{N}$ gibt an, wie groß der Anteil der gesuchten Objekte an der Gesamtmenge ist. Es werden nun zufällig n Objekte aus der Gesamtmenge zufällig ausgewählt und bezüglich ihrer Eigenschaft untersucht.

$$X \sim \text{HG}(n; N, M)$$

Die Zufallsvariable X gibt nun an, wie häufig das gesuchte Merkmal in den gewählten n Beobachtungen auftritt.

$$\begin{aligned} P(X = k) &= \frac{\binom{M}{k} \cdot \binom{N-M}{n-k}}{\binom{N}{n}} \\ \mathbb{E} X &= n \cdot p = n \cdot \frac{M}{N} \\ \text{Var } X &= \frac{N-n}{N-1} \cdot n \cdot p \cdot (p-1) \end{aligned}$$

Bei der Varianz wird der Bruch $\frac{N-n}{N-1}$ als Endlichkeitskorrektur bezeichnet.

4.0.6.1 Annäherung durch die Binomialverteilung

Allgemein gesagt, ist die Approximation möglich, wenn n im Vergleich zu N sehr groß ist.

$$\begin{aligned} n = 1 &\Rightarrow \frac{N - n}{N - 1} = 1 \\ N \rightarrow \infty &\Rightarrow \frac{N - n}{N - 1} = \frac{1 - \frac{n}{N}}{1 - \frac{1}{N}} \rightarrow 1 \end{aligned}$$

4.0.7 Übersicht über die diskreten Verteilungen

Übersicht 6
(S. 142)

4.1 Stetige Verteilungen

Abschnitt 8

Die Schiefe der Verteilungsfunktion

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

Die Wölbung, Krümmung der Verteilungsfunktion

$$\gamma_2 = \frac{\mathbb{E} (X - \mathbb{E} X)^4}{\sigma^4} - 3$$

Wobei für die Standardnormalverteilung gilt,

$$\gamma_1 = 0, \quad \gamma_2 = 0$$

4.1.1 Normalverteilung

Abschnitt 8.2.1

$$\begin{aligned} X &\sim N(\mu, \sigma^2) \\ f(x) &= \frac{e^{-\frac{(x-\mu)^2}{2\sigma^2}}}{\sqrt{2\pi} \cdot \sigma} \end{aligned}$$

Wobei μ den Mittelwert beschreibt, und σ^2 die Varianz, das Quadrat der Standardabweichung σ .

4.1.1.1 Standardnormalverteilung

Abschnitt 8.2.2

Die Standardnormalverteilung ist eine spezielle Form der Normalverteilung mit den Parametern

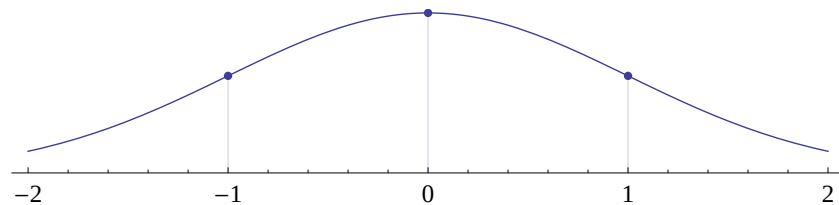


Abbildung 4.1.1: Dichtefunktion der Normalverteilung $X \sim N(0, 1)$, auch Standardnormalverteilung genannt. Das Maximum liegt bei $0, \frac{1}{\sqrt{2\pi \cdot \sigma}}$. Die Wendepunkte sind bei $\mu - \sigma = -1$ und $\mu + \sigma = 1$.

$$\mu = 0 \quad \text{und} \quad \sigma^2 = 1$$

Die Zufallsvariable wird zur besseren Unterscheidung, mit Z bezeichnet.

$$Z \sim N(0, 1)$$

Es können alle Normalverteilungen auf die Standardnormalverteilung transformiert werden, so wird nur eine Nachschlagetabelle (Tabelle A auf Seite 317) benötigt.

$$\frac{X - \mu}{\sigma} := Z \sim N(0, 1)$$

Die Dichtefunktion wird ebenso nicht mit F sondern mit Φ bezeichnet.

$$\Phi(z) = P(Z \leq z)$$

Example 5. Die Körpergröße einer Population...

Beispiel 8.1
(S. 150)

$$\begin{aligned} P(x_1 \leq X \leq x_2) &\sim N(\mu, \sigma^2) \\ P\left(\frac{x_1 - \mu}{\sigma} \leq \frac{X - \mu}{\sigma} \leq \frac{x_2 - \mu}{\sigma}\right) &\sim N(0, 1) \end{aligned}$$

4.1.1.2 σ -Bereiche und Referenzbereiche

Abschnitt 8.2.3

Tabellen für Werte der Φ -Funktion konzentrieren sich auf den Bereich $\pm 3\sigma$ um μ herum, da innerhalb dieses Bereichs 99.7% aller Werte sind. Für medizinische Anwendungen sind vorallem Referenzbereiche interessant, das sind bereiche, in denne genau 95% und 99% aller Werte liegen.

Tabelle 8.1
(S. 152)

Bezeichnung	Intervall	P
1 σ -Bereich	$\mu - \sigma \leq X \leq \mu + \sigma$	0.6827
2 σ -Bereich	$\mu - 2\sigma \leq X \leq \mu + 2\sigma$	0.9545
3 σ -Bereich	$\mu - 3\sigma \leq X \leq \mu + 3\sigma$	0.9973
95%-Referenzbereich	$\mu - 1.96\sigma \leq X \leq \mu + 1.96\sigma$	0.95
99%-Referenzbereich	$\mu - 2.58\sigma \leq X \leq \mu + 2.58\sigma$	0.99

Tabelle 4.1: σ -Bereiche und Referenzbereiche

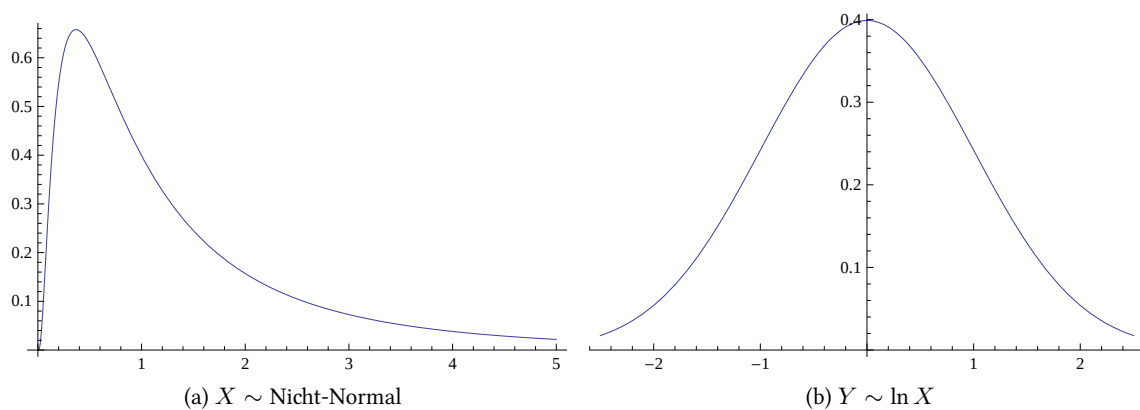


Abbildung 4.1.2: Normalisierende Transformation

4.1.2 Normalisierende Transformationen

Abschnitt 8.2.4

Nicht alle Zufallsdaten für, die für eine Verteilung zugrunde liegen sind Normalverteilt. Daher muss die Verteilung transformiert werden, damit die Merkmale für die Normalverteilung geeignet sind.

In medizinischen Anwendungen besitzen viele Daten von Natur aus eine Rechtsschiefeverteilung, sie müssen mithilfe einer logarithmischen Transformation behandelt werden.

4.1.2.1 Lognormalverteilung

Die Logarithmische Transformation kann nur für $X > 0$ verwendet werden, da sonst das Ergebnis komplex wird, bzw. für $X = 0$ nicht definiert ist, da es asymptotisch gegen negativ unendlich geht.

$$\begin{aligned} X \sim \text{LN}(\mu, \sigma^2) &\Leftrightarrow Y \sim \text{N}(\mu, \sigma^2) \\ X = e^Y &\Leftrightarrow Y = \ln X \end{aligned}$$

Daraus folgt:

$$\begin{aligned} P(X \leq d) &= P(\ln X \leq \ln d) = P(Y \leq \ln d) = \Phi\left(\frac{\ln d - \mu}{\sigma}\right) \\ P(Y \leq c) &= P(e^Y \leq e^c) = P(X \leq e^c) = \Phi\left(\frac{c - \mu}{\sigma}\right) \end{aligned}$$

Example 6. Sei, $\tilde{\mu}_X$ der Median von X und $\tilde{\mu}_Y$ der Median von Y .

$$\begin{aligned} \tilde{\mu}_Y = \mu &\Rightarrow \\ \frac{1}{2} = P(Y \leq \tilde{\mu}_Y) &= P(Y \leq \mu) = \\ &= P(X \leq e^\mu) = P(e^Y \leq e^\mu) \\ &\Rightarrow \tilde{\mu}_X = e^\mu \end{aligned}$$

Example 7. Die Konzentration von Serum IgM bei Kindern...

Beispiel 8.3
(S. 154)

Anmerkungen zur Lognormalverteilung Bei extrem rechtsschiefen Verteilungen reicht oft das einfache Logarithmieren nicht aus. Solche Verteilungen müssen Doppellogarithmiert werden.

Grundsätzlich gilt also:

$$Y = \ln(X + a)$$

Wobei a der Versatz ist, der verwendet werden muss, damit die Werte im Bereich > 0 liegen.

Bei linksschiefen Verteilungsdichtefunktionen wird stattdessen potenziert.

$$Y = X^a \quad \text{mit } a > 1$$

5 Vorlesung am 15. April 2010

5.0.3 Sätze der Wahrscheinlichkeitsrechnung

5.0.3.1 Tschebyscheff'sche Ungleichung

Abschnitt 8.3.1

Die Tschebyscheff'sche Ungleichung¹ dient zur Abschätzung der Wahrscheinlichkeit, dass die Zufallsvariable X außerhalb des Bereichs $k \cdot \sigma \pm \mu$. Wobei $k > 0$.

$$P(|X - \mu| > k \cdot \sigma) \leq \frac{1}{k^2}$$

5.0.3.2 Gesetz der großen Zahlen

Abschnitt 8.3.2

Grundaussage des Gesetzes ist, dass bei einem größeren Stichprobenumfang die Schätzung des Erwartungswerts durch den Mittelwert genauer wird.

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \quad \forall X_i : \mathbb{E} X_i = \mu \wedge \text{Var } X_i = \sigma^2$$

Nun gilt für den Erwartungswert und für die Varianz:

¹Auch: Tschebescheff'sche Ungleichung, Tschebyschow-Ungleichung, eng.: Chebyshev, rus.: Чебышёв

$P(X - \mu > k\sigma)$	$> \sigma$	$> 2\sigma$	$> 3\sigma$
allgemeine Verteilung		$\leq \frac{1}{4} = 0.25$	$\leq \frac{1}{9} \approx 0.11$
eingipfelige, symmetrische Verteilung		$\leq \frac{1}{9} \approx 0.11$	$\leq \frac{4}{81} \approx 0.049$
Gauß-Verteilung	$= 0.32 \approx \frac{1}{3}$	$= 0.046$	$= 0.003$

Tabelle 5.1: Anwendung der Tschebyscheff'schen Ungleichung

$$\begin{aligned}\mathbb{E} \bar{X} &= \mu \\ \text{Var} \bar{X} &= \frac{\sigma^2}{n} \\ \sigma_{\bar{X}} &= \frac{\sigma}{\sqrt{n}}\end{aligned}$$

5.0.3.3 Zentraler Grenzwertsatz

Abschnitt 8.3.3

Der zentrale Grenzwertsatz sagt aus, dass – unter recht allgemeinen Bedingungen – die Summe einer großen Anzahl von Zufallsvariablen normalverteilt ist.

Sind alle X_i unabhängig identisch verteilt

$$X_i \quad \text{für } i = 1, \dots, n$$

dann haben Erwartungswert und Varianz die gleichen Werte für alle X_i

$$\mathbb{E} X_i = \mu \quad \wedge \quad \text{Var} X_i = \sigma^2$$

und die Summe aller Zufallsvariablen ist – asymptotisch – normalverteilt.

$$\sum_{i=1}^n X_i \underset{\text{asym}}{\sim} N(n\mu, n\sigma^2)$$

Durch Transformation auf die Standardnormalverteilung ergibt sich

$$Z_n = \frac{\sum X_i - n\mu}{\sqrt{n} \cdot \sigma} \underset{\text{asym}}{\sim} N(0, 1)$$

Verteilung von Mittelwerten Es geht aus dem Gesetz hervor, dass wenn der Stichprobenumfang n hinreichend groß ist – bei realistischen Fällen $n \geq 25$ – der Mittelwert normalverteilt ist.

Binomialverteilung Die Binomialverteilung – für hinreichend großes n – kann ebenso durch eine Normalverteilung approximiert werden.

$$X \sim B(n, p)$$

Hinreichend groß kann mit der Erfüllung folgender Ungleichung angenommen werden.

$$n \cdot p \cdot q \geq 9$$

Nun gilt – approximativ – als Ersatz für die Binomialverteilung

$$X \underset{\text{approx}}{\sim} N(np, npq)$$

5.0.4 Prüfverteilungen

Abschnitt 8.5

5.0.4.1 t-Verteilung

Abschnitt 8.5.1

Voraussetzungen

1. $X_1, \dots, X_n$ sind unabhängig
2. $\mathbb{E} X_i = \mu \quad \wedge \quad \text{Var } X_i = \sigma^2$ für $i = 1, \dots, n$
3. *alle X_i sind normalverteilt*

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

$$Z = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim N(0, 1)$$

$$T = \frac{\bar{X} - \mu}{S/\sqrt{n}} \sim t_{n-1}$$

Der Index bei t gibt die Anzahl der Freiheitsgrade f der Verteilung an. Als Faustregel gilt hier, dass es immer $f = n - \|\text{Geschätzten Parameter}\|$ sind, da die geschätzten Parameter vorherbestimmt sind und sozusagen *Freiheit* wegnehmen.

Für die Schätzung der Standardabweichung S gilt

$$S^2 := \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n-1} \left[\sum_{i=1}^n X_i^2 - n\bar{X}^2 \right]$$

Für den Erwartungswert und die Varianz der t-Verteilung gilt nun:

$$\begin{aligned} T &\sim t_f \\ \mathbb{E} T &= 0 \\ \text{Var } T &= \frac{f}{f-2} \end{aligned}$$

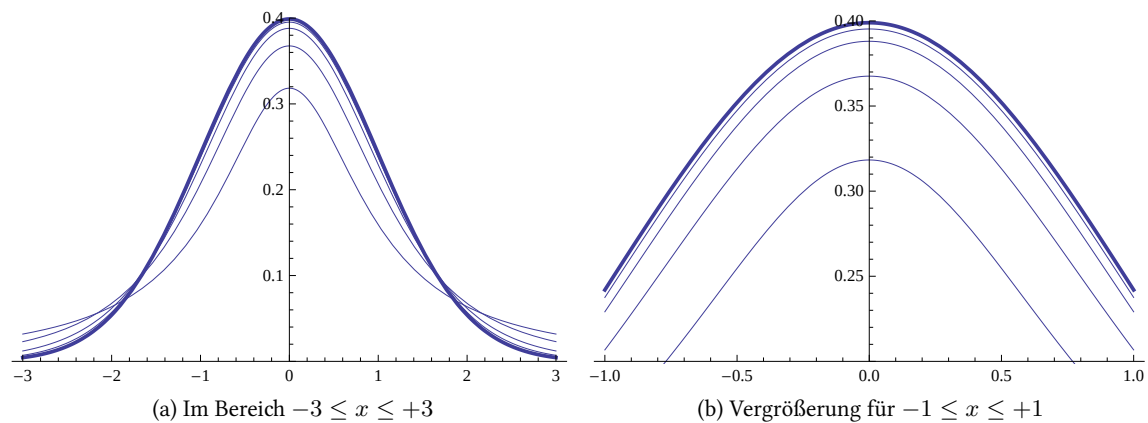


Abbildung 5.0.1: Approximation der Standardnormalverteilung durch die t-Verteilung mit Freiheitsgraden $f \in \{1, 3, 9, 27\}$. Die Dichtefunktion der Normalverteilung ist dick eingezeichnet.

Die Varianz gilt insbesondere nur für $f \geq 3$ und kann daher immer als > 1 angenommen werden. Desweiteren kann die t-Verteilung an die Standardnormalverteilung angenähert werden.

$$t_f \xrightarrow{f \rightarrow \infty} N(0, 1)$$

Ausgewählte γ -Quantile für die t-Verteilung können in der Tabelle B auf Seite 318 nachgeschlagen werden.

$$T \sim t_f \Rightarrow P(T \leq t_{f,\gamma}) = \gamma$$

Example 8. Nachschlagen des Wertes ab dem 95% der Werte links von ihm liegen. (Suchen des 0.95-Quantils) bei einer t-Verteilung mit $f = 100$ und der Standardnormalverteilung.

$$0.95 = P(T \leq t_{100,0.95}) \Rightarrow t_{100,0.95} = 1.660$$

$$0.95 = P(Z \leq z_{0.95}) \Rightarrow z_{0.95} = 1.645$$

5.0.4.2 χ^2 -Verteilung

Abschnitt 8.5.2

In der einfachsten Form – mit Freiheitsgrad $f = 1$ – beschreibt die χ^2 -Verteilung die Verteilung des Quadrats einer standardnormalverteilten Zufallsvariablen $Z \sim N(0, 1)$. Nun gilt für Varianz und Erwartungswerte von $Z^2 \sim \chi_1^2$

$$\begin{aligned} \text{Var } Z &= \mathbb{E} Z^2 - (\mathbb{E} Z)^2 \\ \mathbb{E} Z^2 &= \text{Var } Z + (\mathbb{E} Z)^2 = 1 + 0 = 1 \end{aligned}$$

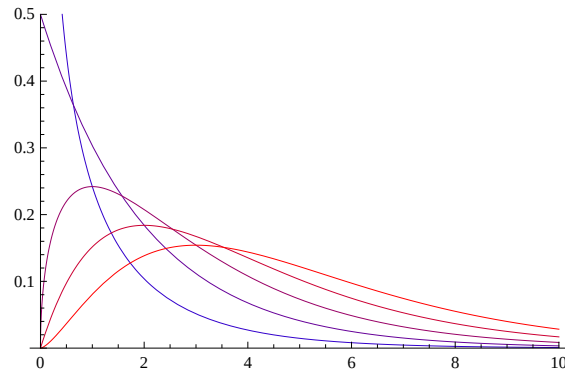


Abbildung 5.0.2: Dichtefunktion der χ^2 -Verteilung mit Freiheitsgraden $1 \leq f \leq 5$

Unter den Voraussetzungen

1. $Z_1, \dots, Z_n$ sind unabhängig verteilt
2. $Z_i \sim N(0, 1) \quad \forall i = 1, \dots, n$

gilt

$$\sum_{i=1}^n Z_i^2 \sim \chi_n^2$$

mit

$$\begin{aligned} \mathbb{E} \left(\sum_{i=1}^n Z_i^2 \right) &= \sum \mathbb{E} Z_i^2 = \sum 1 = n \\ \text{Var} \left(\sum_{i=1}^n Z_i^2 \right) &= 2n \\ \gamma_1 &= \sqrt{\frac{8}{n}} \end{aligned}$$

Da die Schiefe γ_1 immer nur positiv sein kann – mit Annäherung an 0 bei $n \rightarrow \infty$ – ist die χ^2 -Verteilung immer rechtsschief und wird mit steigendem n der Standardnormalverteilung immer ähnlicher.

Es gilt also, für Standardnormalverteilungen, für allgemeine Normalverteilungen $N(\mu, \sigma^2)$ und für geschätzte Normalverteilungen mit Mittelwertschätzer $\bar{X}$.

$$\begin{aligned}\sum_{i=1}^n Z_i^2 &\sim \chi_n^2 \\ \sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma} \right)^2 &\sim \chi_n^2 \\ \sum_{i=1}^n \left(\frac{X_i - \bar{X}}{\sigma} \right)^2 &\sim \chi_{n-1}^2\end{aligned}$$

Da der Mittelwert $\bar{X}$ die Bedingungen der χ^2 -Verteilung einschränkt hat die Verteilung nur $n - 1$ Freiheitsgrade.

Für die Schätzung der Standardabweichung S gilt folgender Hinweis:

$$\begin{aligned}\sum_{i=1}^n \left(\frac{X_i - \bar{X}}{\sigma} \right)^2 &= \frac{n-1}{n-1} \cdot \frac{1}{\sigma^2} \cdot \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{n-1}{\sigma^2} \cdot S^2 \\ \Rightarrow \quad \frac{n-1}{\sigma^2} \cdot S^2 &\sim \chi_{n-1}^2\end{aligned}$$

Spezielle Werte für γ -Quantilen von $\chi_{f,\gamma}^2$ sind in Tabelle E auf Seite 322 zu finden.

5.0.4.3 F-Verteilung

Abschnitt 8.5.3

Unter den Voraussetzungen,

1. S_1^2, S_2^2 sind (empirische) Varianzen von zwei unabhängigen Stichproben, aus zwei normalverteilten Grundgesamtheiten mit derselben Varianz σ^2
2. n_1, n_2 sind die Stichprobenumfänge der zwei Stichproben

ist die F-Verteilung definiert als

$$F_{n_1-1, n_2-1} := \frac{S_1^2}{S_2^2} = F$$

Die F-Verteilung hat somit zwei Freiheitsgrade f_1, f_2 (manchmal auch als m, n bezeichnet).

$$f_1 = n_1 - 1 \quad f_2 = n_2 - 1$$

Mit der vereinfachten Bezeichnung χ_f^2 für eine *Zufallsvariable*, die χ_f^2 -verteilt ist, gilt auch

$$F_{n_1-1, n_2-1} = \frac{S_1^2}{S_2^2} = \frac{\frac{\sigma^2}{n_1-1} \cdot \chi_{n_1-1}^2}{\frac{\sigma^2}{n_2-1} \cdot \chi_{n_2-1}^2} = \frac{\chi_{n_1-1}^2}{\chi_{n_2-1}^2} \cdot \frac{n_2-1}{n_1-1}$$

Beziehungsweise, mit alternativer Bezeichnung:

$$F_{m,n} = \frac{\chi_m^2}{\chi_n^2} \cdot \frac{n}{m}$$

Im Speziellen gilt somit auch, dass die t-Verteilung ein Spezialfall der F-Verteilung ist

$$F_{1,n} = t_n^2$$

5.1 Schätzverfahren

Abschnitt 9

5.1.1 Punktschätzungen

Abschnitt 9.2

$$\bar{X} = \frac{1}{n} \cdot \sum_{i=1}^n X_i$$

5.1.1.1 Kriterien zur Güte einer Schätzung

Abschnitt 9.2.2

Erwartungstreue Schätzer sollten möglichst mit dem unbekannten Parameter übereinstimmen, so sollten Schätzvorschriften nicht systematisch einen zu hohen oder zu niedrigen Wert liefern. Eine erwartungstreue Schätzung nennt man unverzerrt.

Konsistenz Eine Schätzung sollte konsistent sein, dass heißt mit wachsender Stichprobengröße soll auch die Genauigkeit steigen. Für eine konsistente Schätzung gilt:

$$\text{Var } \bar{X} \xrightarrow{n \rightarrow \infty} 0$$

Effizienz Für eine Schätzung ist eine möglichst kleine Varianz von großer Bedeutung. Eine Schätzung mit hoher Effizienz liefert schon mit Stichproben geringeren Umfangs schon brauchbare Werte. Die Effizienz ist insbesondere dann wichtig, wenn zwei Schätzverfahren für einen Parameter miteinander verglichen werden sollen.

5.1.2 Spezielle Schätzfunktionen

Abschnitt 9.2.3

5.1.2.1 Erwartungswert

Der Mittelwert $\bar{X}$ schätzt den Erwartungswert μ der Grundgesamtheit. Es gilt daher

$$\begin{aligned} \mathbb{E} \bar{X} &= \mu \\ \text{Var } \bar{X} &= \frac{\sigma^2}{n} \xrightarrow{n \rightarrow \infty} 0 \end{aligned}$$

Diese Schätzung ist sowohl erwartungstreu und konsistent.

5.1.2.2 Median

Der Median $\tilde{X}$ ist ein erwartungstreuer Schätzer, falls die Verteilung stetig und symmetrisch ist, in diesem Fall gilt auch $\tilde{\mu} = \mu$. Daher ist bei der Normalverteilung der Median ein erwartungstreuer Schätzer für den Erwartungswert. Es gilt,

$$\text{Var } \tilde{X} = \frac{\pi}{2} \cdot \frac{\sigma^2}{n} \xrightarrow{n \rightarrow \infty} 0$$

Daher ist $\tilde{X}$ auch eine konsistente Schätzung, jedoch ist $\bar{X}$ ein effizienterer Schätzer für μ als $\tilde{X}$, da die Varianz größer ist.

5.1.2.3 Varianz

Herleitung von $\mathbb{E} S^2$

Nebenrechnung Zuerst wird gezeigt, dass

$$\sum_{i=1}^n (X_i - \bar{X})^2 = \sum (X_i - \mu)^2 - n(\bar{X} - \mu)^2 \quad (5.1.1)$$

$$\begin{aligned} \sum_{i=1}^n (X_i - \bar{X})^2 &= \sum (X_i^2 - 2X_i\bar{X} + \bar{X}^2) \\ &= \sum X_i^2 - 2\left(\sum X_i\right)\bar{X} + n\bar{X}^2 \\ &= \sum X_i^2 - 2(n\bar{X})\bar{X} + n\bar{X}^2 \\ &= \sum X_i^2 - n\bar{X}^2 \end{aligned} \quad (5.1.2)$$

$$\begin{aligned} \sum_{i=1}^n (X_i - \mu)^2 &= \sum (X_i^2 - 2X_i\mu + \mu^2) \\ &= \sum X_i^2 - 2\left(\sum X_i\right)\mu + n\mu^2 \\ &= \sum X_i^2 - 2n\bar{X}\mu + n\mu^2 \end{aligned} \quad (5.1.3)$$

$$-n(\bar{X} - \mu)^2 = -n\bar{X}^2 + 2n\bar{X}\mu - n\mu^2 \quad (5.1.4)$$

Aus (5.1.2) sowie aus (5.1.3) und (5.1.4) ist ersichtlich, dass (5.1.1) stets stimmt.

Hauptrechnung Vorrausgesetzt wird

1. $\mathbb{E} X_i = \mu$
2. $\text{Var } X_i = \sigma^2$
3. $X_1, \dots, X_n$ sind unabhängig
4. Anmerkung: Normalverteilung ist nicht erforderlich

Daher ist,

$$\mathbb{E} \bar{X} = \mu \quad \wedge \quad \text{Var } \bar{X} = \frac{\sigma^2}{n}$$

Allgemein gilt:

$$\sum (X_i - \bar{X})^2 = \sum (X_i - \mu)^2 - n (\bar{X} - \mu)^2$$

somit ist

$$\begin{aligned} \mathbb{E} S^2 &= \mathbb{E} \left(\frac{1}{n-1} \cdot \sum (X_i - \bar{X})^2 \right) \\ &= \frac{1}{n-1} \cdot \mathbb{E} \left(\sum (X_i - \bar{X})^2 \right) \\ &= \frac{1}{n-1} \cdot \left[\mathbb{E} \left(\sum (X_i - \mu)^2 \right) - n \mathbb{E} (\bar{X} - \mu)^2 \right] \end{aligned} \quad (5.1.5)$$

Nun ist einerseits

$$\mathbb{E} \left(\sum (X_i - \mu)^2 \right) = \sum \mathbb{E} (X_i - \mu)^2 = \sum \text{Var } X_i = n\sigma^2 \quad (5.1.6)$$

Andererseits

$$\mathbb{E} (\bar{X} - \mu)^2 = \text{Var } \bar{X} = \frac{\sigma^2}{n} \quad (5.1.7)$$

Aus (5.1.6) und (5.1.7) folgt, dass

$$\mathbb{E} S^2 = \frac{1}{n-1} \left[n\sigma^2 - \frac{n \cdot \sigma^2}{n} \right] = \sigma^2$$

daraus folgt

$$\mathbb{E} S^2 = \sigma^2$$

Dies ist ein erwartungstreuer Schätzer für σ^2 .

Anmerkung: Aus (5.1.6) ist zu ersehen, dass $\frac{1}{n} \sum (X_i - \mu)^2$ eine erwartungstreue Schätzfunktion für σ^2 ist, sofern μ bekannt ist, denn

$$\mathbb{E} \left[\frac{1}{n} \cdot (X_i - \mu)^2 \right] = \frac{\mathbb{E} (X_i - \mu)^2}{n} = \frac{n \cdot \sigma^2}{n} = \sigma^2$$

Alternative Herleitung von $\mathbb{E} S^2$ für unbekannte μ und bekannte $\bar{x}$, unter den Voraussetzungen,

1. $X_i \sim N(\mu, \sigma^2)$
2. $X_1, \dots, X_n$ sind unabhängig

$$\frac{n-1}{\sigma^2} \cdot S^2 \sim \chi_{n-1}^2$$

Daraus folgt,

$$\mathbb{E} \left(\frac{n-1}{\sigma^2} \cdot S^2 \right) = n-1$$

Daher gilt für den Erwartungswert:

$$\frac{n-1}{\sigma^2} \cdot \mathbb{E} S^2 = n-1 \quad \Leftrightarrow \quad \mathbb{E} S^2 = \sigma^2$$

Herleitung von $\text{Var } S^2$ unter den Voraussetzungen

1. $X_i \sim N(\mu, \sigma^2)$
2. $X_1, \dots, X_n$ sind unabhängig

$$\frac{n-1}{\sigma^2} \cdot S^2 \sim \chi_{n-1}^2$$

$$\begin{aligned} \text{Var} \left(\frac{n-1}{\sigma^2} \cdot S^2 \right) &= 2(n-1) \\ \frac{(n-1)^2}{\sigma^4} \cdot \text{Var } S^2 &= 2(n-1) \\ \text{Var } S^2 &= \frac{2\sigma^4}{n-1} \xrightarrow{n \rightarrow \infty} 0 \end{aligned}$$

Herleitung von $\mathbb{E} S < \sigma$ unter der Voraussetzung, dass gilt

$$0 < \text{Var } S = \mathbb{E} S^2 - (\mathbb{E} S)^2 = \sigma^2 - (\mathbb{E} S)^2$$

Diese Bedingung ist normalerweise bei allen realistischen Fällen erfüllt. Daraus folgt,

$$\begin{aligned} (\mathbb{E} S)^2 &= \sigma^2 - \text{Var } S < \sigma^2 \\ \mathbb{E} S &< \sigma \end{aligned}$$

5.1.2.4 Schätzung der Wahrscheinlichkeit

Unter den Voraussetzungen

1. $X_i \sim B(1, p)$
2. $X_1, \dots, X_n$ sind unabhängig

gilt,

$$\begin{aligned} X &:= \sum X_i \sim B(n, p) \\ \hat{p} &:= \frac{X}{n} \end{aligned}$$

Wobei $\hat{p}$ in diesem Zusammenhang sowohl eine Zufallsvariable (so wie hier), als auch eine konkrete Zahl beschreiben kann.

$$\mathbb{E} \hat{p} = \frac{\mathbb{E} X}{n} = \frac{n \cdot p}{n} = p$$

Daher ist die Schätzung erwartungstreu, für die Varianz gilt,

$$\text{Var } \hat{p} = \text{Var} \left(\frac{X}{n} \right) = \frac{\text{Var } X}{n^2} = \frac{n \cdot p \cdot (1 - p)}{n^2} = \frac{p \cdot q}{n} \xrightarrow{n \rightarrow \infty} 0$$

Daher ist die Schätzung auch konsistent.

6 Vorlesung am 22. April 2010

6.0.3 Intervallschätzungen

Abschnitt 9.3

6.0.3.1 Bedeutung eines Konfidenzintervalls

Abschnitt 9.3.1

Punktschätzungen sind oft ungenügend, da ein einzelner Wert keine Informationen enthält, wie sehr er vom tatsächlichen Parameter abweichen kann. Es können bei einer Schätzung natürlich keine genauen Angaben getroffen werden, jedoch ist es wünschenswert, dass sich der tatsächliche – unbekannte – Wert in der *näheren Umgebung* der Schätzung liegt.

Es wird aus den Daten der Stichprobe ein *Konfidenzintervall* konstruiert, in der Hoffnung, dass bei dem passend gewählten Verfahren, der gesuchte Parameter innerhalb des Intervalls liegt.

Irrtumswahrscheinlichkeit α ist die Wahrscheinlichkeit, dass der Parameter *nicht* im Intervall liegt.

Konfidenzwahrscheinlichkeit γ oder Konfidenzniveau bezeichnet die Wahrscheinlichkeit, dass der Parameter im Intervall liegt

$$\gamma = 1 - \alpha$$

6.0.3.2 Konfidenzintervall für einen Erwartungswert

Abschnitt 9.3.2

Herleitung für beliebige Konfidenzintervalle für μ Voraussetzung

$$\begin{aligned}\bar{X} &\sim N\left(\mu, \frac{\sigma^2}{n}\right) \Rightarrow \frac{\bar{X} - \mu}{\sigma^2/\sqrt{n}} \sim N(0, 1) \\ P\left(z_{\alpha_1} \leq \frac{\bar{X} - \mu}{\sigma^2/\sqrt{n}} \leq z_{1-\alpha_2}\right) &= 1 - (\alpha_1 + \alpha_2)\end{aligned}$$

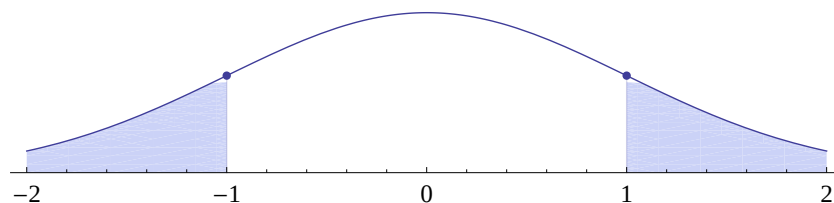


Abbildung 6.0.1: Intervallgrenzen, Links und Rechts eingezeichnet α_1 und α_2 , Punkte zeigen jeweils z_{α_1} und $z_{1-\alpha_2}$

Wobei gilt:

$$z_0 = -\infty \quad \wedge \quad z_1 = +\infty$$

$$\alpha_1 + \alpha_2 = \alpha$$

Nun gilt ja:

$$\begin{aligned} z_{\alpha_1} &\leq \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \leq z_{1-\alpha_2} \\ \frac{z_{\alpha_1} \cdot \sigma}{\sqrt{n}} &\leq \bar{X} - \mu \leq \frac{z_{1-\alpha_2} \cdot \sigma}{\sqrt{n}} \\ \bar{X} - \frac{z_{\alpha_1} \cdot \sigma}{\sqrt{n}} &\geq \mu \geq \bar{X} - \frac{z_{1-\alpha_2} \cdot \sigma}{\sqrt{n}} \end{aligned}$$

Daher

$$P \left(\bar{X} - \frac{z_{1-\alpha_2} \cdot \sigma}{\sqrt{n}} \leq \mu \leq \bar{X} - \frac{z_{\alpha_1} \cdot \sigma}{\sqrt{n}} \right) = 1 - \alpha$$

wobei $\alpha = \alpha_1 + \alpha_2$.

Konfidenzintervalle zu Konfidenzniveau $1 - (\alpha_1 + \alpha_2)$ sind allgemein alle Intervalle der Form

$$\left[\bar{x} - \frac{z_{1-\alpha_2} \cdot \sigma}{\sqrt{n}}; \bar{x} - \frac{z_{\alpha_1} \cdot \sigma}{\sqrt{n}} \right]$$

beziehungsweise, durch die Eigenschaft der Standardnormalverteilung: $z_{\alpha_1} = -z_{1-\alpha_1}$ kann das Intervall auch angescheiben werden als:

$$\left[\bar{x} - \frac{z_{1-\alpha_2} \cdot \sigma}{\sqrt{n}}; \bar{x} + \frac{z_{1-\alpha_1} \cdot \sigma}{\sqrt{n}} \right]$$

Für die Spezialfälle $\alpha_1 = \alpha_2 = \frac{\alpha}{2}$ ergibt sich

$$\left[\bar{x} - \frac{z_{1-\frac{\alpha}{2}} \cdot \sigma}{\sqrt{n}}; \bar{x} + \frac{z_{1-\frac{\alpha}{2}} \cdot \sigma}{\sqrt{n}} \right]$$

für $\alpha_1 = 0$ und $\alpha_2 = \alpha$ gilt

$$\left[\bar{x} - \frac{z_{1-\alpha} \cdot \sigma}{\sqrt{n}}; +\infty \right)$$

für $\alpha_1 = \alpha$ und $\alpha_2 = 0$ gilt

$$\left(-\infty; \bar{x} + \frac{z_{1-\alpha} \cdot \sigma}{\sqrt{n}} \right]$$

Grenzen zweiseitiger Konfidenzintervalle für μ unter der Voraussetzung $\bar{X} \sim N(\mu, \sigma^2)$
bzw. $\bar{X} \sim N(\mu, \frac{\sigma^2}{n})$

Falls σ bekannt ist:

$$\bar{x} \pm \frac{z_{1-\frac{\alpha}{2}} \cdot \sigma}{\sqrt{n}}$$

Falls σ unbekannt ist:

$$\bar{x} \pm \frac{t_{n-1, 1-\frac{\alpha}{2}} \cdot s}{\sqrt{n}} \cdot \sqrt{\frac{N-n}{N-1}}$$

Wobei die Endlichkeitskorrektur $\sqrt{\frac{N-n}{N-1}}$ nur für Werte von $\frac{n}{N} > 0.01$ benötigt wird.

6.0.3.3 Konfidenzintervall für eine Wahrscheinlichkeit

Abschnitt 9.3.3

Unter der Voraussetzung, dass X binomialverteilt ist,

$$X \sim B(n, p)$$

ist die geschätzte Wahrscheinlichkeit definiert als,

$$\hat{p} := \frac{X}{n} \quad \text{bzw.} \quad \hat{p} := \frac{x}{n}$$

Weiters ist x approximativ normalverteilt

$$X \sim N(np, npq) \quad \text{für } npq \geq 9$$

daher

$$\frac{X}{n} \sim N\left(p, \frac{p \cdot q}{n}\right)$$

Approximativ heißt dass,

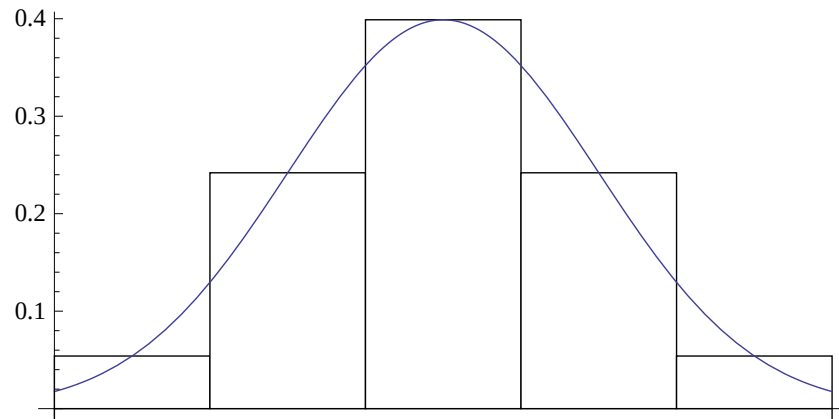


Abbildung 6.0.2: Annäherung der Binomialverteilung

$$\frac{X}{n} \sim N \left(\hat{p}, \frac{\hat{p} \cdot \hat{q}}{n} \right)$$

wenn gilt,

$$n\hat{p} > 5 \quad \wedge \quad n\hat{q} > 5$$

Grenzen zweiseitiger Konfidenzintervalle für p

$$\hat{p} \pm \left[z_{1-\frac{\alpha}{2}} \cdot \sqrt{\frac{\hat{p} \cdot \hat{q}}{n}} \cdot \sqrt{\frac{N-n}{N-1}} + \frac{1}{2n} \right]$$

Bei Bedarf wird die Endlichkeitskorrektur hinzugefügt, falls $\frac{n}{N} > 0.1$ und daher nicht vernachlässigbar ist.

Die Stetigkeitskorrektur $\frac{1}{2n}$ resultiert aus der Annäherung der Dichtefunktion mit über die Fläche. Siehe Abbildung 6.0.2.

6.0.4 Bedeutung des Stichprobenumfangs

Abschnitt 9.4.1

Die Breite des Konfidenzintervalls (BK) ist (etwa) indirektproportional zur Wurzel der Stichprobengröße n .

$$BK = \frac{2 \cdot s \cdot t_{n-1, 1-\frac{\alpha}{2}}}{\sqrt{n}}$$

$$n \geq 60 : \quad t_{n-1, 0.975} \leq 2$$

Wenn die Breite des Konfidenzintervalls vorgegeben ist, gilt:

$$n = \frac{4 \cdot t_{n-1, 1-\frac{\alpha}{2}}^2 \cdot s^2}{BK^2} \approx \frac{16 \cdot s^2}{BK^2}$$

Wenn $\alpha = 5\%$ für $n \gtrsim 60$

$$t_{n-1, 1-\frac{\alpha}{2}} \leq 2$$

7 Vorlesung am 29. April 2010

7.1 Das Prinzip eines statistischen Tests

Abschnitt 10

7.1.1 Durchführung eines Statistischen Tests

Abschnitt 10.1

H_0 Nullhypothese

H_1 Alternativhypothese

Sei nun μ_0 eine vorgegebene Zahl, so kann man nun unterscheiden

1. Einseitige Fragestellung

$$H_0 : \mu = \mu_0 \quad \text{gegen} \quad H_1 : \mu \neq \mu_0$$

2. Einseitige Fragestellung

a)

$$H_0 : \mu = \mu_0 \Leftrightarrow H_0 : \mu \geq \mu_0$$

$$H_1 : \mu < \mu_0 \Leftrightarrow H_1 : \mu < \mu_0$$

b)

$$H_0 : \mu = \mu_0 \Leftrightarrow H_0 : \mu \leq \mu_0$$

$$H_1 : \mu > \mu_0 \Leftrightarrow H_1 : \mu > \mu_0$$

In Kurzform $H_0 : \mu = \mu_0$

1. Zweiseitige Fragestellung

$$H_1 : \mu \neq \mu_0$$

2. Einseitige Fragestellung

a)

$$H_1 : \mu < \mu_0$$

b)

$$H_1 : \mu > \mu_0$$

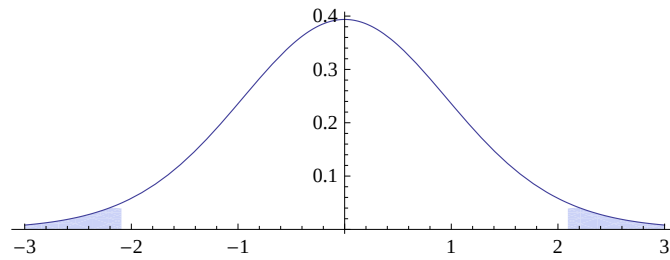


Abbildung 7.1.1: Verteilung von T beim Zutreffen von H_0 , die beiden markierten Teile sind jeweils $\frac{\alpha}{2}$ in der Mitte befindet sich daher $(1 - \alpha)$ -Annahmebereich.

7.1.2 Arten von Fehlern

Fehler 1. Art α -Fehler

Fehler 2. Art β -Fehler

T ... Teststatistik

t ... Testgröße, Prüfgröße, Wert der Teststatistik

Angenommen die Teststatistik

$$T \sim t_{n-1}$$

$$T = \frac{\bar{X} - \mu_0}{S/\sqrt{n}}$$

Fragestellung für Nullhypothese $H_0 : \mu = \mu_0$

1. Zweiseitige Fragestellung

$$H_1 : \mu \neq \mu_0$$

2. Einseitige Fragestellung

a)

$$H_1 : \mu < \mu_0$$

b)

$$H_1 : \mu > \mu_0$$

Kritischer Bereich

1. Zweiseitige Fragestellung

$$\left(-\infty; t_{n-1, \frac{\alpha}{2}}\right) \cup \left(t_{n-1, 1-\frac{\alpha}{2}}; +\infty\right)$$

2. Einseitige Fragestellung

a)

$$(-\infty; t_{n-1, \alpha})$$

b)

$$(t_{n-1, 1-\alpha}; +\infty)$$

Annahmehereich

1. Zweiseitige Fragestellung

$$\left[t_{n-1, \frac{\alpha}{2}}; t_{n-1, 1-\frac{\alpha}{2}}\right]$$

2. Einseitige Fragestellung

a)

$$[t_{n-1, \alpha}; +\infty)$$

b)

$$(-\infty; t_{n-1, 1-\alpha}]$$

$$|t| > t_{n-1, 1-\frac{\alpha}{2}} \Leftrightarrow t < -t_{n-1, 1-\frac{\alpha}{2}} \vee t > t_{n-1, 1-\frac{\alpha}{2}}$$

Falls H_0 zutrifft, aber – aufgrund der Stichprobenwerte – H_0 abgelehnt – und daher H_1 angenommen – wird, handelt es sich um einen Fehler 1. Art / α -Fehler.

Falls H_1 zutrifft, aber – aufgrund der Stichprobenwerte – H_1 nicht angenommen – und daher H_0 beibehalten – wird, handelt es sich um einen Fehler 2. Art / β -Fehler.

7.1.3 Testentscheidungen und Konsequenzen

Abschnitt 10.2

α	Ergebnis
10%	schwach signifikant
5%	signifikant
1%	hoch signifikant
1‰	höchst signifikant

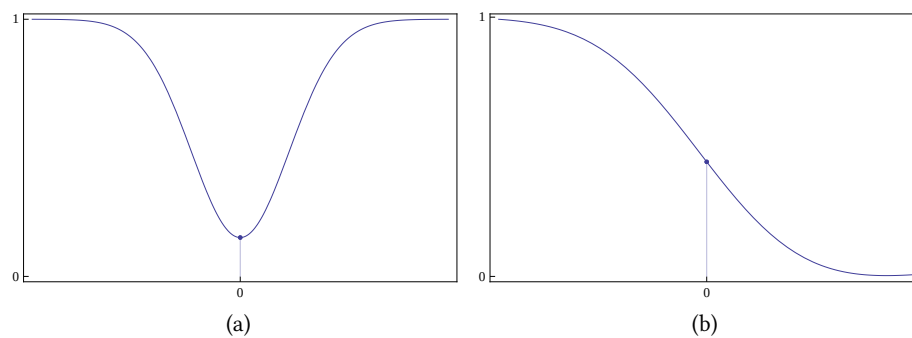


Abbildung 7.1.2: Gütefunktion $g(\mu)$ für Einseitige Fragestellungen (7.1.2a) und Zweiseitige Fragestellungen (7.1.2b) (Typ a).

Testentscheidung	Wirklichkeit	
	H_0 gilt	H_1 gilt
für H_0	richtige Entscheidung $1 - \alpha$	Fehler 2. Art β
für H_1	Fehler 1. Art α	richtige Entscheidung $1 - \beta$
Summe	1	1

Tabelle 7.1: Testentscheidungen und Fehler

$1 - \beta \dots$ Güte, Teststärke, Trennschärfe, Macht, Power

Allgemein ist die Gütefunktion $g(\mu)$ die Wahrscheinlichkeit, H_0 abzulehnen, also H_1 anzunehmen, falls μ der tatsächlich (jedoch unbekannte) Parameterwert ist.

7.1.3.1 p -Wert und Konfidenzintervall

Abschnitt 10.2.2

$$p < \alpha \Rightarrow \text{Testentscheidung } H_1$$

$$p \geq \alpha \Rightarrow \text{Testentscheidung } H_0$$

Interpretation des p -Werts Der p -Wert – zu einer aktuellen Stichprobe – ist jenes Signifikanzniveau, bei dem H_0 gerade noch beibehalten werden würde (maximaler α -Wert). Oder, in einer Alternativen Definition, es ist jener Wert, bei dem H_1 gerade schon angenommen werden könnte (minimaler α -Wert).

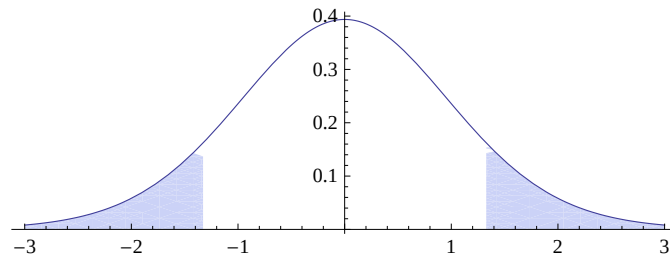


Abbildung 7.1.3: p_1 ist nur der linke Teil, p_2 beide Teile zusammen

Zusammenhang des p -Werts zu ein- und zweiseitigen Fragestellungen Bezeichnen wir mit

p_1 p -Wert eines einseitigen Tests – bei einseitiger Fragestellung

p_2 p -Wert des entsprechenden zweiseitigen Tests – bei entsprechender zweiseitiger Fragestellung

Dann gilt:

$$p_2 = 2 \cdot p_1 \quad \text{bzw.} \quad p_1 = \frac{p_2}{2}$$

Zusammenhang zwischen Testergebniss und Konfidenzintervall des t-Tests am Beispiel für eine Stichprobe

$$T = \frac{\bar{X} - \mu_0}{S/\sqrt{n}} \quad \text{im Falle von } \mu = \mu_0$$

Die Wahrscheinlichkeit, dass $H_0 : \mu = \mu_0$ beibehalten wird, falls H_0 zutrifft – also tatsächlich $\mu = \mu_0$ ist, beträgt:

$$P(t_{n-1, \alpha_1} \leq T \leq t_{n-1, 1-\alpha_2}) = 1 - (\alpha_1 + \alpha_2)$$

für den Fall dass $\mu = \mu_0$.

Wenn wir definieren $\alpha := \alpha_1 + \alpha_2$, dann gilt für

1. Zweiseitige Fragestellungen $\alpha_1, \alpha_2 = \frac{\alpha}{2}$

2. Einseitige Fragestellungen

$$H_1 : \mu < \mu_0 \quad \alpha_1 = \alpha \quad \wedge \quad \alpha_2 = 0$$

$$H_1 : \mu > \mu_0 \quad \alpha_1 = 0 \quad \wedge \quad \alpha_2 = \alpha$$

Es gilt

$$\begin{aligned}
 t_{n-1, \alpha_1} &\leq T \leq t_{n-1, 1-\alpha_2} \\
 t_{n-1, \alpha_1} &\leq \frac{\bar{X} - \mu_0}{S/\sqrt{n}} \leq t_{n-1, 1-\alpha_2} \\
 \frac{t_{n-1, \alpha_1} \cdot S}{\sqrt{n}} &\leq \bar{X} - \mu_0 \leq \frac{t_{n-1, 1-\alpha_2} \cdot S}{\sqrt{n}} \\
 \bar{X} - \frac{t_{n-1, \alpha_1} \cdot S}{\sqrt{n}} &\geq \mu_0 \geq \bar{X} - \frac{t_{n-1, 1-\alpha_2} \cdot S}{\sqrt{n}}
 \end{aligned}$$

Daraus ergibt sich folgendes Intervall – mit Schreibweise $t_{n-1, \alpha_1} = t_{n-1, 1-\alpha_1}$, da nur diese Werte sich in einer Tabelle finden.

$$\mu_0 \in [\bar{X} - t_{n-1, 1-\alpha_2} \cdot S/\sqrt{n}; \bar{X} + t_{n-1, 1-\alpha_1} \cdot S/\sqrt{n}]$$

$(1-\alpha)$ -Konfidenzintervall μ_0 Somit gilt also, sofern die Nullhypothese $H_0 : \mu = \mu_0$ zutrifft, wird H_0 auf dem Signifikanzniveau gehalten, genau dann, wenn μ_0 von entsprechenden $(1-\alpha)$ -Konfidenzintervall überdeckt wird.

Im Speziellen bedeutet dies

1. Bei der zweiseitigen Fragestellung, mit $\alpha_1 = \alpha_2 = \frac{\alpha}{2}$

$$H_1 : \mu \neq \mu_0$$

H_0 wird beibehalten – falls¹ $|t| \leq t_{n-1, 1-\frac{\alpha}{2}}$ – genau dann, wenn

$$\mu_0 \in \left[\bar{x} - t_{n-1, 1-\frac{\alpha}{2}} \cdot s/\sqrt{n}; \bar{x} + t_{n-1, 1-\frac{\alpha}{2}} \cdot s/\sqrt{n} \right]$$

2. Bei einseitiger Fragestellung

- a) mit $\alpha_1 = \alpha$ und $\alpha_2 = 0$

$$H_1 : \mu < \mu_0$$

H_0 wird beibehalten – falls $t \geq -t_{n-1, 1-\alpha}$ – genau dann, wenn

$$\mu_0 \in (-\infty; \bar{x} + t_{n-1, 1-\alpha} \cdot s/\sqrt{n}]$$

- b) mit $\alpha_1 = 0$ und $\alpha_2 = \alpha$

$$H_1 : \mu > \mu_0$$

H_0 wird beibehalten – falls $t \leq t_{n-1, 1-\alpha}$ – genau dann, wenn

$$\mu_0 \in [\bar{x} - t_{n-1, 1-\alpha} \cdot s/\sqrt{n}; +\infty)$$

¹Wobei hier t – dies gilt auch für $\bar{x}$ und s im folgenden – anstelle von T steht, da es sich um eine konkrete Stichprobe handelt.

Example 9. Geburtsgewicht von Raucherkindern

Beispiel 10.1
(S. 198)

7.1.3.2 Multiples Testen

Bei k *unabhängig* durchgeführten Tests ist die Wahrscheinlichkeit H_0 beizubehalten, sofern H_0 auch tatsächlich zutrifft nur $(1 - \alpha)^k$ hoch.

$$(1 - \alpha)^k = 1 - \binom{k}{1}\alpha + \binom{k}{2}\alpha^2 - \binom{k}{3}\alpha^3 + \dots + (-1)^k \alpha^k \approx 1 - k\alpha$$

Die Wahrscheinlichkeit H_0 zu verwerfen, obwohl H_0 tatsächlich zutrifft (ein Fehler 1. Art) ist

$$1 - (1 - \alpha)^k \approx k\alpha$$

Der Fehler multipliziert sich also mit der Anzahl der durchgeführten Test

Wenn also $\frac{\alpha}{k}$ statt α als Signifikanzniveau für jeden einzelnen dieser k unabhängigen Tests gewählt wird (*Bonferroni-Korrektur*) beträgt die Wahrscheinlichkeit für einen Fehler 1. Art insgesamt nur noch

$$\left(1 - \left(1 - \frac{\alpha}{k}\right)^k\right) \approx \alpha$$

8 Vorlesung am 6. Mai 2010

8.1 Lagetests

Abschnitt 11

8.1.1 t-Tests

Abschnitt 11.1

8.1.1.1 t-Test für eine Stichprobe

Abschnitt 11.1.1

Unter der Voraussetzung, dass alle Werte normalverteilt sind,

$$X \sim N(\mu, \sigma^2)$$

gilt für die Teststatistik:

$$t = \frac{\bar{x} - \mu_0}{s/\sqrt{n}} = \frac{\bar{x} - \mu_0}{s \cdot \sqrt{\frac{1}{n}}}$$

$$\begin{array}{ll} H_0 : \mu = \mu_0 & H_0 \text{ wird abgelehnt, falls:} \\ H_1 : \mu \neq \mu_0 & |t| > t_{n-1, 1-\frac{\alpha}{2}} \\ H_1 : \mu > \mu_0 & t > t_{n-1, 1-\alpha} \\ H_1 : \mu < \mu_0 & t < t_{n-1, \alpha} = -t_{n-1, 1-\alpha} \end{array}$$

8.1.1.2 t-Test für zwei verbundene Stichproben

Abschnitt 11.1.2

Voraussetzungen für die Zufallsvariablen X, Y :

$$\begin{array}{ll} D & := X - Y \\ \delta & := \mathbb{E} D \\ D & \sim N(\delta, \sigma_d^2) \end{array}$$

$$t = \frac{\bar{d} - 0}{s_d/\sqrt{n}}$$

wobei hier gilt:

$$d_i := x_i - y_i \quad s_d = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (d_i - \bar{d})^2}$$

$H_0 : \delta = 0$ H_0 wird abgelehnt, falls:

$H_1 : \delta \neq 0$ $|t| > t_{n-1, 1-\frac{\alpha}{2}}$

$H_1 : \delta > 0$ $t > t_{n-1, 1-\alpha}$

$H_1 : \delta < 0$ $t < t_{n-1, \alpha} = -t_{n-1, 1-\alpha}$

Konfidenzintervall, dass einem zweiseitigen Test entspricht:

$$\left[\bar{d} - t_{n-1, 1-\frac{\alpha}{2}} \cdot \frac{s_d}{\sqrt{n}}; \bar{d} + t_{n-1, 1-\frac{\alpha}{2}} \cdot \frac{s_d}{\sqrt{n}} \right]$$

8.1.1.3 t-Test für zwei unverbundene Stichproben

Abschnitt 11.1.3

Herleitung für die Prüfgröße Es ist

$$\bar{X} := \frac{1}{n_1} \cdot \sum_{i=1}^n X_i \quad \wedge \quad \bar{Y} := \frac{1}{n_2} \cdot \sum_{i=1}^n Y_i$$

wobei die Zufallsvariablen normalverteilt sind

$$X_i \sim N(\mu_1, \sigma^2) \quad \wedge \quad Y_i \sim N(\mu_2, \sigma^2)$$

und alle X_i, Y_i sind unabhängig, daher gilt

$$\bar{X} \sim N\left(\mu_1, \frac{\sigma^2}{n_1}\right) \quad \wedge \quad \bar{Y} \sim N\left(\mu_2, \frac{\sigma^2}{n_2}\right)$$

So gilt nun für den Erwartungswert

$$\begin{aligned} \mathbb{E}(\bar{X} - \bar{Y}) &= \mathbb{E}\bar{X} + (-1)\mathbb{E}\bar{Y} \\ &= \mu_1 - \mu_2 \end{aligned}$$

und für die Varianz beider Verteilungen

$$\begin{aligned}\text{Var}(\bar{X} - \bar{Y}) &= \text{Var} \bar{X} + \text{Var}(-\bar{Y}) \\ \text{Var} \bar{X} + (-1)^2 \text{Var} \bar{Y} &= \sigma^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)\end{aligned}$$

Daher ist die Differenz normalverteilt mit:

$$\bar{X} - \bar{Y} \sim N \left(\mu_1 - \mu_2, \sigma^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right) \right)$$

Falls $H_0 : \mu_1 = \mu_2$ zutrifft, gilt daher:

$$\bar{X} - \bar{Y} \sim N \left(0, \sigma^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right) \right)$$

Eigenschaften Voraussetzungen

$$X \sim N(\mu_1, \sigma^2) \quad \wedge \quad Y \sim N(\mu_2, \sigma^2)$$

wobei X, Y unabhängig sind. Für die Teststatistik gilt nun (siehe Herleitung):

$$t = \frac{(\bar{x} - \bar{y})}{s \cdot \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

wobei für s^2 gilt:

$$s^2 = \frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{n_1 + n_2 - 2}$$

$$H_0 : \mu_1 = \mu_2 \Leftrightarrow H_0 : \mu_1 - \mu_2 = 0 \quad H_0 \text{ wird abgelehnt, falls:}$$

$$H_1 : \mu_1 \neq \mu_2 \Leftrightarrow H_1 : \mu_1 - \mu_2 \neq 0 \quad |t| > t_{f, 1-\frac{\alpha}{2}}$$

$$H_1 : \mu_1 > \mu_2 \Leftrightarrow H_1 : \mu_1 - \mu_2 > 0 \quad t > t_{f, 1-\alpha}$$

$$H_1 : \mu_1 < \mu_2 \Leftrightarrow H_1 : \mu_1 - \mu_2 < 0 \quad t < -t_{f, 1-\alpha}$$

Wobei für den Freiheitsgrad $f := n_1 + n_2 - 2$

Speziell für $n := n_1 = n_2$

$$t = \frac{\bar{x} - \bar{y}}{s \cdot \sqrt{\frac{2}{n}}} \quad s^2 := \frac{s_1^2 + s_2^2}{2}$$

8.1.1.4 Welch-Test

Abschnitt 11.1.4

$$\begin{aligned} X_i &\sim N(\mu_1, \sigma_1^2) & Y_i &\sim N(\mu_2, \sigma_2^2) \\ \bar{X} &\sim N(\mu_1, \frac{\sigma_1^2}{n_1}) & \bar{Y} &\sim N(\mu_2, \frac{\sigma_2^2}{n_2}) \end{aligned}$$

$$\text{Var}(\bar{X} - \bar{Y}) = \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}$$

$$(\bar{X} - \bar{Y}) \sim N\left(\mu_1 - \mu_2, \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}\right)$$

$$t = \frac{(\bar{x} - \bar{y}) - 0}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

Für die Freiheitsgrade der t-Verteilung gilt:

$$f := \left\lfloor \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{\left(\frac{s_1^2}{n_1}\right)^2}{n_1-1} + \frac{\left(\frac{s_2^2}{n_2}\right)^2}{n_2-1}} \right\rfloor$$

8.1.1.5 Voraussetzungen für t-Tests

Abschnitt 11.1.5

Der t-Test ist robust gegenüber Abweichungen von einer Normalverteilung, jedoch ist folgendes zusätzlich zu beachten.

t-Test für eine Stichprobe

$$\begin{array}{ll} n \geq 10 & \wedge \quad \text{Daten annähernd symmetrisch} \\ n \geq 25 & \text{sonst} \end{array}$$

t-Test für zwei verbundene Stichproben

$$n \geq 10 \quad \wedge \quad d_i \text{ annähernd symmetrisch}$$

insbesondere, wenn X, Y eine gleiche Verteilungsform haben.

t-Test für zwei unverbundene Stichproben Vorrausgesetzt ist eine gleiche Verteilung der Zufallsvariablen

$$\begin{array}{ll} n_1, n_2 \geq 10 & \wedge \quad \text{symetrische Verteilung} \\ n_1, n_2 \geq 20 & \text{sonst} \end{array}$$

und X, Y sind vom gleichen Verteilungstyp.

8.1.1.6 Andere Anwendungen des t-Tests

Abschnitt 11.1.6

Der t-Test kann auch nicht als Lagetest eingesetzt werden, um zu testen, ob sich ein empirische Korrelationskoeffizient nach Pearson ρ signifikant von 0 unterscheidet, kann folgende Prüfgröße verwendet werden. r steht hier für den empirisch ermittelten Wert.

$$t = \frac{r}{\sqrt{\frac{1-r^2}{n-2}}}$$

$$H_0 : \rho = 0$$

$$H_1 : \rho > 0 \quad t > t_{n-2, 1-\alpha}$$

$$H_1 : \rho < 0 \quad t < -t_{n-2, 1-\alpha}$$

8.1.2 Rangsummentests

Abschnitt 11.2

Rangsummentests sind *verteilungsfreie*, also *nicht-prametrische*, Tests, da sie keine bestimmte Verteilungsform der Stichproben voraussetzen. Die Prüfgrößen werden nicht aus den Messwerten, sondern aus deren Rangzahlen berechnet.

8.1.2.1 Wilcoxon-Test für eine Stichprobe

Abschnitt 11.2.1

Dieser Test überprüft, ob und in welchen Maß die Werte einer Stichprobe von einem vorgegebenen Sollwert $\tilde{\mu}_0$ abweichen. Vorraussetzung für diese Verfahren ist, eine symetrische Verteilung der Stichprobenwerte, also $\tilde{\mu} = \mu$.

Die Nullhypothese lautet:

$$H_0 : \tilde{\mu} = \tilde{\mu}_0$$

Für die Durchführung wird folgender Ablauf eingehalten:

1. Die Differenz zum Sollwert wird bestimmt: $x_i - \tilde{\mu}_0$
2. Alle Werte mit $x_i = \tilde{\mu}_0$ werden eliminiert. Der Umfang der Stichprobe muss angepasst werden.
3. Die Differenzen werden nun nach der Größe ihres Betrags sortiert, die kleinste erhält die Rangzahl 1.
4. Bei gleichen Differenzen, wird eine mittlere Rangzahl zugeordnet: $\text{Rang} + \frac{1}{\|\text{Gleiche Differenzen}\|}$
Hierbei spricht man von *verbundenen Rängen*.
5. Die Rangzahlen der positiven und negativen Differenzen werden addiert und mit R^- , R^+ bezeichnet.
6. Prüfgröße $R := \min \{R^-, R^+\}$
7. Nun muss in einer Tabelle¹ der kritische Wert R_{krit} in Abhängigkeit von Stichprobenumfang n und Irrtumswahrscheinlichkeit α nachgeschlagen werden.
8. H_0 wird angenommen, wenn $R \leq R_{krit} \Leftrightarrow R^- \leq R_{krit} \vee R^+ \leq R_{krit}$

Die Nullhypothese, und die Alternativhypothesen für zweiseitigen und einseitige Tests lauten wie folgt:

$$\begin{array}{ll}
 H_0 : \tilde{\mu} = \tilde{\mu}_0 & H_0 \text{ wird abgelehnt, falls:} \\
 H_1 : \tilde{\mu} \neq \tilde{\mu}_0 & R^- \leq R_{krit} \vee R^+ \leq R_{krit} \\
 H_1 : \tilde{\mu} > \tilde{\mu}_0 & R^- \leq R_{krit} \\
 H_1 : \tilde{\mu} < \tilde{\mu}_0 & R^+ \leq R_{krit}
 \end{array}$$

8.1.2.2 Wilcoxon-Test für zwei verbundene Stichproben

Abschnitt 11.2.2

Dieser Test verläuft wie sein Pendant zum t-Test mit verbundenen zwei Stichproben. Es werden wiederum die Mediane der beiden Stichproben $\tilde{\mu}_1, \tilde{\mu}_2$ verglichen.

Der Ablauf ist ähnlich dem des Tests für eine Stichprobe:

1. Es wird für jedes Wertepaar die Differenz gebildet: $d_i := x_i - y_i$
2. Alle $d_i = 0$ werden eliminiert.
3. Der Rest verläuft gleich.

Ist nun

$$D := X - Y$$

die Verteilung der Differenzen, so ist der Median der Differenzen $\tilde{\delta} = \tilde{\mu}_D$

¹Tabelle C Anhang Seite 319

Dies gilt natürlich nur unter der Voraussetzung, dass X, Y annähernd die gleiche Verteilungsform haben. D ist daher eher symmetrisch verteilt.

Mit dieser Notation definiert sich folgende Nullhypothese und Alternativhypothesen für zweiseitigen und einseitige Tests. Die Bedingungen zur Annahme, oder Ablehnung von H_0 bleiben dieselben wie für den Test mit einer Stichprobe.

$$H_0 : \tilde{\delta} = 0 \Leftrightarrow H_0 : \tilde{\mu}_1 = \tilde{\mu}_2$$

$$H_1 : \tilde{\delta} \neq 0$$

$$H_1 : \tilde{\delta} > 0$$

$$H_1 : \tilde{\delta} < 0$$

Example 10. Zehn Personen nehmen sechs Monate lang eine Diät ...

Beispiel 11.3
(S. 217)

Mit Hilfe der kritischen Werte aus Tabelle C lässt sich dieses Beispiel sowohl in einseitiger als auch zweiseitiger Fragestellung lösen.

$H_1 : \tilde{\delta} \neq 0$	$R = 7$	Testentscheidung	Begründung
$\alpha = 10\%$	$R_{krit} = 10$	$H_1 : \tilde{\delta} \neq 0$	$R \leq R_{krit}$
$\alpha = 5\%$	$R_{krit} = 8$	$H_1 : \tilde{\delta} \neq 0$	$R \leq R_{krit}$
$\alpha = 2\%$	$R_{krit} = 5$	$H_0 : \tilde{\delta} = 0$	$R > R_{krit}$
$\alpha = 1\%$	$R_{krit} = 3$	$H_0 : \tilde{\delta} = 0$	$R > R_{krit}$

$H_1 : \tilde{\delta} > 0$	$R^- = 7$	Testentscheidung	Begründung
$\alpha = 5\%$	$R_{krit} = 10$	$H_1 : \tilde{\delta} > 0$	$R^- \leq R_{krit}$
$\alpha = 2.5\%$	$R_{krit} = 8$	$H_1 : \tilde{\delta} > 0$	$R^- \leq R_{krit}$
$\alpha = 1\%$	$R_{krit} = 5$	$H_0 : \tilde{\delta} = 0$	$R^- > R_{krit}$
$\alpha = 0.5\%$	$R_{krit} = 3$	$H_0 : \tilde{\delta} = 0$	$R^- > R_{krit}$

9 Vorlesung am 20. Mai 2010

9.0.2.3 U-Test (Wilcoxon-Test für unverbundene Stichproben)

Abschnitt 11.2.3

$$H_0 : \tilde{\mu}_1 = \tilde{\mu}_2$$

Die Nullhypothese nimmt an, dass die Mediane beider Stichproben gleich sind. Die Umfänge der Stichproben n_1, n_2 muss nicht gleich sein, jedoch müssen als Voraussetzung beide Stichproben die gleiche Verteilungsform haben. Um die Tabelle D¹ aus dem Anhang zu verwenden muss zusätzlich gelten: $n_1 \geq n_2$.

Vorgehensweise:

1. Alle Werte beider Stichproben werden aufsteigend sortiert und mit Rangzahlen versehen. (Siehe Wilcoxon-Test)
2. Die Summe der Rangzahlen wird dann als R_1 und R_2 bezeichnet.
3. Danach wird für die Testgröße berechnet:

$$\begin{aligned}U_1 &= n_1 n_2 + \frac{n_1 \cdot (n_1 + 1)}{2} - R_1 \\U_2 &= n_1 n_2 + \frac{n_2 \cdot (n_2 + 1)}{2} - R_2\end{aligned}$$

Dabei gilt:

$$\begin{aligned}U_1 + U_2 &= n_1 \cdot n_2 \\R_1 + R_2 &= \frac{(n_1 + n_2) \cdot (n_1 + n_2 + 1)}{2}\end{aligned}$$

Die Testgröße:

$$U = \min \{U_1, U_2\}$$

4. Die Nullhypothese wird abgelehnt,

$$U \leq U_{krit}$$

Falls nun gilt, $n_1 \geq 10 \wedge n_2 \geq 10$, so ist $U \sim N\left(\frac{n_1 \cdot (n_1 + n_2 + 1)}{2}, \frac{n_1 n_2}{2}\right)$.

¹Die Tabelle ist so konstruiert, dass n_1 vertikal ($\downarrow$) und n_2 horizontal ($\rightarrow$) angeordnet sind.

9.0.2.4 Vergleich von Rangsummentests und t-Tests

Abschnitt 11.2.4

9.0.3 Vorzeichentest

Abschnitt 11.3

9.0.3.1 Vorzeichentest für eine Stichprobe

Abschnitt 11.3.1

$$H_0 : \tilde{\mu}_0 = \tilde{\mu}_1$$

$$\forall x_i \begin{cases} x_i > \tilde{\mu}_0 & \text{gewertet als } + \\ x_i < \tilde{\mu}_0 & \text{gewertet als } - \end{cases}$$

Verallgemeinert lässt sich jetzt sagen, die Nullhypothese wird abgelehnt, wenn es nun *entscheidend* mehr Plus als Minus gibt, oder vice versa. Im Detail heißt dass, zuerst wird die Prüfgröße errechnet

$$V_+ := ||+||$$

Anmerkung: Es wird nur eine Größe benötigt, und nicht wie bei den Rangsummentests zwei, von denen dann die kleinere genommen wird, da in Tabelle F für die kritischen Werte ein Bereich angegeben ist. Wäre nur ein Wert angegeben, so wäre dies der kleinere, und man müsste als Testgröße das Minimum aus Plus und Minus Zahl nehmen, beziehungsweise kann die andere Zahl auch errechnet werden.

Falls nun die Nullhypothese zutrifft, gilt:

$$V_+ \sim B \left(n, \frac{1}{2} \right)$$

Da nun die Binomialverteilung durch die Normalverteilung approximiert werden kann:

$$B(n, p) \approx N(np, npq) \quad \text{wenn } npq \geq 9$$

so gilt nun

$$B \left(n, \frac{1}{2} \right) \approx N \left(\frac{n}{2}, \frac{n}{4} \right) \quad \text{wenn } npq \geq 9 \Leftrightarrow n \geq 36$$

9.0.3.2 Vorzeichentest für zwei verbundene Stichproben

Abschnitt 11.3.2

Die Stichproben werden als X, Y und die Differenz als $D := X - Y$ bezeichnet. Der Median von D wird als $\tilde{\delta} := \tilde{\mu}_D$ bezeichnet.

$$\begin{aligned}
 H_0 : P(X < Y) &= P(X > Y) \Leftrightarrow P(X - Y < 0) = P(X - Y > 0) \Leftrightarrow \tilde{\delta} = 0 \\
 H_1 : P(X < Y) &\neq P(X > Y) \Leftrightarrow P(X - Y < 0) \neq P(X - Y > 0) \Leftrightarrow \tilde{\delta} \neq 0 \\
 H_1 : P(X < Y) &> P(X > Y) \Leftrightarrow P(X - Y < 0) > P(X - Y > 0) \Leftrightarrow \tilde{\delta} < 0 \\
 H_1 : P(X < Y) &< P(X > Y) \Leftrightarrow P(X - Y < 0) < P(X - Y > 0) \Leftrightarrow \tilde{\delta} > 0
 \end{aligned}$$

$$\begin{aligned}
 H_0 : \tilde{\delta} = 0 & \quad H_0 \text{ wird abgelehnt, falls:} \\
 H_1 : \tilde{\delta} \neq 0 & \quad V_+ < V_u \vee V_+ > V_o \\
 H_1 : \tilde{\delta} < 0 & \quad V_+ < V_u \\
 H_1 : \tilde{\delta} > 0 & \quad V_+ > V_o
 \end{aligned}$$

Ähnlich wie beim Vorzeichentest für eine Stichprobe wird hier jeder Differenz $d_i := x_i - y_i$ ein Vorzeichen zugeordnet. Die Zahl der positiven wird als Prüfgröße V_+ bezeichnet. Für die Entscheidung wird in Tabelle F die obere bzw. untere Grenze des Annahmebereichs V_o, V_u nachgeschlagen.

Falls nun $n \geq 36$, so gilt bei Zutreffen von H_0 , approximativ

$$V_+ \sim N\left(\frac{n}{2}, \frac{n}{4}\right)$$

Anmerkung: Für $\mu = n \cdot p = n \cdot \frac{1}{2}$ und $\sigma^2 = n \cdot p \cdot q = n \cdot \frac{1}{2} \cdot \frac{1}{2}$

Schranken für den Annahmebereich

Für eine zweiseitige Fragestellung

$$V_u, V_o := \frac{n}{2} \pm \left(z_{1-\frac{\alpha}{2}} \cdot \frac{\sqrt{n}}{2} + \frac{1}{2} \right)$$

Für eine einseitige Fragestellung

$$\begin{aligned}
 V_u &:= \frac{n}{2} - \left(z_{1-\alpha} \cdot \frac{\sqrt{n}}{2} + \frac{1}{2} \right) \\
 V_o &:= \frac{n}{2} + \left(z_{1-\alpha} \cdot \frac{\sqrt{n}}{2} + \frac{1}{2} \right)
 \end{aligned}$$

9.1 Tests zum Vergleich von Häufigkeiten

Abschnitt 12

9.1.1 Binomialtest für eine Stichprobe

Abschnitt 12.1

Dieser Test überprüft die relative Häufigkeit der Ausprägung A mit der vorgegebenen Wahrscheinlichkeit p_0 vereinbar ist.

$$H_0 : p = p_0$$

$$H_1 : p \neq p_0$$

Falls H_0 zutrifft, ist die Stichprobe binomialverteilt:

$$X \sim B(n, p_0)$$

$$H_1 : p < p_0$$

$$H_1 : p > p_0$$

Wenn n hinreichend groß ist mit $np_0q_0 \geq 9$, lässt sich die Binomialverteilung durch eine Normalverteilung approximieren

$$B(n, p_0) \approx N(np_0, np_0q_0)$$

Wobei hier gilt $q_0 := 1 - p_0$

9.1.2 χ^2 -Tests

Abschnitt 12.2

9.1.2.1 χ^2 -Vierfelder-Test

Abschnitt 12.2.1

	A	$\bar{A}$	Randsummen
B	a	b	$n_1 = a + b$
$\bar{B}$	c	d	$n_2 = c + d$
Randsummen	$a + c$	$b + d$	$n = a + b + c + d$

Tabelle 9.1: Vierfeldertabelle für den χ^2 -Vierfelder-Test

9.1.2.2 χ^2 -Unabhängigkeits-Test

Diesem Test liegt eine Stichprobe mit zwei alternativen Merkmalen und dem Umfang n . Die Häufigkeiten sind in der Vierfeldertafel (Tabelle 9.1) dargestellt.

$$H_0 : P(A | B) = P(A)$$

Wenn die Nullhypothese erfüllt ist, dann sind die Ereignisse unabhängig voneinander und es gilt:

$$\begin{aligned} P(A | B) = P(A) &\Leftrightarrow P(A \cap B) = P(A) \cdot P(B) \\ &\Leftrightarrow P(A | B) = P(A | \bar{B}) = P(A) \\ &\Leftrightarrow P(\bar{A} | B) = P(\bar{A} | \bar{B}) = P(\bar{A}) \\ &\Leftrightarrow P(B | A) = P(B | \bar{A}) = P(B) \\ &\Leftrightarrow P(\bar{B} | A) = P(\bar{B} | \bar{A}) = P(\bar{B}) \end{aligned}$$

Die Idee des χ^2 -Tests ist es nun, die beobachteten Häufigkeiten a, b, c, d mit den erwarteten Häufigkeiten – welche die Ereignisse unter der Nullhypothese annehmen würden – zu vergleichen.

$$\sum_{i \in \{a, b, c, d\}} \frac{(\text{beobachtete Häufigkeit} - \text{erwartete Häufigkeit})^2}{\text{erwartete Häufigkeit}}$$

Die Prüfgröße ist annähernd χ^2 -verteilt mit $f = 1$. Für den Vierfeldertest berechnet sie sich als:

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

H_0 wird verworfen, falls $\chi^2 > \chi_{1, 1-\alpha}$ (beim zweiseitigen Test mit $\neq$) (Siehe Tabelle F im Anhang)

Example 11.

$$\begin{aligned} A &= \sigma \\ B &= \varphi \\ A, B &= \text{Rhesusfaktor positiv} \\ \bar{A}, \bar{B} &= \text{Rhesusfaktor negativ} \end{aligned}$$

9.1.2.3 χ^2 -Homogenitäts-Test

Eine andere Interpretation des Vierfelder-Tests ist eine Überprüfung der relativen Häufigkeit zweier unabhängiger Stichproben; es wird also untersucht, ob ein bestimmtes Merkmal identverteilt ist).

$$H_0 : P(A | B) = P(A | \bar{B})$$

$$H_1 : P(A | B) \neq P(A | \bar{B})$$

$$H_1 : P(A | B) > P(A | \bar{B})$$

$$H_1 : P(A | B) < P(A | \bar{B})$$

Bei der einseitigen Fragestellung wird H_1 angenommen, falls $\chi^2 > \chi_{1,1-2\alpha}^2 \wedge \frac{a}{n_1} > \frac{c}{n_2}$

Example 12. Medikamententest

B = Medikament

$\bar{B}$ = Placebo

A = Besserung

$\bar{A}$ = keine Besserung

Daher gibt es n_1 Personen, die ein Medikament bekommen und n_2 Personen, die ein Placebo bekommen haben.

Überlegungen zum einseitigen χ^2 -Homogenitätstest

$$H_0 : P(A | B) = P(A | \bar{B})$$

Falls H_0 zutrifft, gilt

$$\underbrace{P(\chi^2 > \chi_{1,1-2\alpha}^2)}_{2\alpha} \cdot \underbrace{P\left(\frac{a}{n_1} > \frac{c}{n_2} \mid \chi^2 > \chi_{1,1-2\alpha}^2\right)}_{\frac{1}{2}} = \underbrace{P\left(\frac{a}{n_1} > \frac{c}{n_2} \wedge \chi^2 > \chi_{1,1-2\alpha}^2\right)}_{\alpha}$$

$H_0 : P(A | B) = P(A | \bar{B})$ H_0 wird abgelehnt, falls:

$H_1 : P(A | B) \neq P(A | \bar{B})$ $\chi^2 > \chi_{1,1-\alpha}^2$

$H_1 : P(A | B) > P(A | \bar{B})$ $\chi^2 > \chi_{1,1-2\alpha}^2 \wedge \frac{a}{n_1} > \frac{c}{n_2}$

$H_1 : P(A | B) < P(A | \bar{B})$ $\chi^2 > \chi_{1,1-2\alpha}^2 \wedge \frac{a}{n_1} < \frac{c}{n_2}$

Example 13. zum zwei- bzw. einseitigen χ^2 -Test

M ... Medikament

Pl ... Placebo

B ... Besserung

$\bar{B}$... keine Besserung

Hierbei handelt es sich um einen einseitigen Test, da es die Placebowirkung sowohl beim Medikament als logischerweise auch beim Placebo existiert.

	B	$\bar{B}$	Randsummen
M	151	572	$n_1 = 723$
Pl	51	452	$n_2 = 502$
Randsummen	202	1024	$n = 1226$

$$H_0 : P(B | M) = P(B | Pl)$$

$$H_1 : P(B | M) \neq P(B | Pl)$$

$$H_1 : P(B | M) > P(B | Pl)$$

$$\chi^2 = \frac{1226 \cdot (151 \cdot 452 - 51 \cdot 572)^2}{202 \cdot 1024 \cdot 503 \cdot 723} = 24.890\,978\,62$$

$$\frac{a}{n_1} = \frac{151}{723} = 0.208\,852\,006$$

$$\frac{c}{n_2} = \frac{51}{503} = 0.101\,391\,650$$

$H_1 : P(B M) \neq P(B Pl)$	$\chi^2_{1,1-\alpha}$	Testentscheidung	Begründung
$\alpha = 10\%$	2.706	H_1	$\chi^2 > \chi^2_{1,0.9}$
$\alpha = 5\%$	3.841		$\chi^2 > \chi^2_{1,0.95}$
$\alpha = 2.5\%$	5.024		$\chi^2 > \chi^2_{1,0.975}$
$\alpha = 1\%$	6.635		$\chi^2 > \chi^2_{1,0.99}$
$\alpha = 0.5\%$	7.879		$\chi^2 > \chi^2_{1,0.995}$

$H_1 : P(B M) > P(B Pl)$	$\chi^2_{1,1-2\alpha}$	Testentscheidung	Begründung
$\alpha = 5 \%$	2.706	H_1	$\chi^2 > \chi^2_{1,0.95} \wedge \frac{a}{n_1} > \frac{c}{n_2}$
$\alpha = 2.5 \%$	3.841		$\chi^2 > \chi^2_{1,0.975} \wedge \frac{a}{n_1} > \frac{c}{n_2}$
$\alpha = 1.25\%$	5.024		$\chi^2 > \chi^2_{1,0.9875} \wedge \frac{a}{n_1} > \frac{c}{n_2}$
$\alpha = 0.5 \%$	6.635		$\chi^2 > \chi^2_{1,0.995} \wedge \frac{a}{n_1} > \frac{c}{n_2}$
$\alpha = 0.25\%$	7.879		$\chi^2 > \chi^2_{1,0.9975} \wedge \frac{a}{n_1} > \frac{c}{n_2}$

Auf jeden dieser Signifikanzniveaus wirkt das Medikament besser als das Placebo.

10 Vorlesung am 27. Mai 2010

10.0.2.4 χ^2 -Test für $k \cdot l$ Felder

Abschnitt 12.2.3

Wie Vierfelder-Test nur mit Freiheitsgrad

$$f = (k - 1) \cdot (l - 1)$$

für

erwartete Häufigkeiten > 5

10.0.2.5 Assotiationsmaße für qualitative Merkmale

Abschnitt 12.2.4

ϕ -Koeffizient

$$\phi = \sqrt{\frac{\chi^2}{n}}$$

Zusätzlich gilt folgendes:

$$0 \leq \phi \leq 1 \quad \wedge \quad b = c = 0 \Rightarrow \phi = 1$$

Der ϕ -Koeffizient ist 0 bei vollkommener Unabhängigkeit der Merkmale. Im Falle von $\phi = 1$ lässt sich aufgrund eines Merkmals das andere präzise vorhersagen. Der ϕ -Koeffizient ist signifikant größer als 0, falls das Ergebnis des Vierfelder-Tests signifikant ist.

Abschnitt 3.4.2
Seite 52

Assoziationskoeffizient nach Yule

$$Q = \frac{ad - bc}{ad + bc}$$

$$ad = bc \Leftrightarrow a : b = c : d \Leftrightarrow Q = 0$$

$$-1 \leq Q \leq 1$$

10.0.2.6 McNemar Test

Abschnitt 12.2.5

Häufigkeitstest für 2 verbundene Stichproben hinsichtlich eines Alternativmerkmals

		Stichprobe 1	
		A	$\bar{A}$
Stichprobe 2	A	a	b
	$\bar{A}$	c	d

 $X \dots$ Stichprobe 1 $\hat{=}$ Therapie 1 $Y \dots$ Stichprobe 2 $\hat{=}$ Therapie 2

$$H_0 : P(X = A) = P(Y = A)$$

$$H_1 : P(X = A) \neq P(Y = A)$$

Für die Testgröße gilt

$$\chi^2 = \frac{(b - c)^2}{b + c} \quad f = 1$$

Unter der Bedingung

$$b + c \leq 30$$

gilt

$$\chi^2 = \frac{(|b - c| - 1)^2}{b + c}$$

Vorraussetzung

$$\frac{b + c}{2} \geq 5$$

Einseitiger McNemar Test

$$H_0 : P(X = A) = P(Y = A)$$

$$H_1 : P(X = A) > P(Y = A)$$

 H_1 wird angenommen, falls $\chi^2 > \chi^2_{1,1-2\alpha} \wedge c > b$ **Example 14.** ZweiseitigBeispiel 12.5
(S. 239)

10.0.2.7 Mittelwert

Minimum Eigenschaft des Mittelwerts

$$\sum (x_i - \bar{x})^2 < \sum (x_i - c)^2 \quad \forall c \in \mathbb{R} \setminus \{\bar{x}\}$$

Minimum Eigenschaft des Medians

$$\sum |x_i - \tilde{x}| \leq \sum |x_i - c| \quad \forall c \in \mathbb{R} \setminus \{\tilde{x}\}$$

Abschnitt 4.2.6/7

Geometrisches Mittel und Harmonisches Mittel

Das geometrische Mittel x_i sind relative Änderungen (Änderungsfaktoren)

$$\bar{x}_G = \sqrt[n]{\prod x_i}$$

Das harmonische Mittel x_i Quotienten, die sich nur im Nenner unterscheiden

$$\bar{x}_H = \frac{n}{\sum \frac{1}{x_i}}$$

Die Mittel ergeben sich aus dem arithmetischen Mittel, allerdings von entsprechend transformierenden Stichprobenwerten

$$\begin{aligned} \ln \bar{x}_G &= \frac{1}{n} \sum_{i=1}^n \ln x_i \Leftrightarrow \bar{x}_G = e^{\frac{1}{n} \sum \ln x_i} \\ \sqrt[n]{\prod e^{\ln x_i}} &= \sqrt[n]{\prod x_i} \end{aligned}$$

Analog gilt dies auch für $^{10}\log$ und jede andere Basis.

Für das harmonische Mittel gilt:

$$\begin{aligned} \frac{1}{\bar{x}_H} &= \frac{1}{n} \sum_{i=1}^n \frac{1}{x_i} \\ \bar{x}_H &= \frac{1}{\frac{1}{n} \cdot \sum \frac{1}{x_i}} \\ \bar{x}_H &= n \cdot \left(\sum x_i^{-1} \right)^{-1} \end{aligned}$$

11 Vorlesung am 11. Juni 2010

11.0.3 Korrelationsanalyse

Abschnitt 5.2

11.0.3.1 Punktwolke

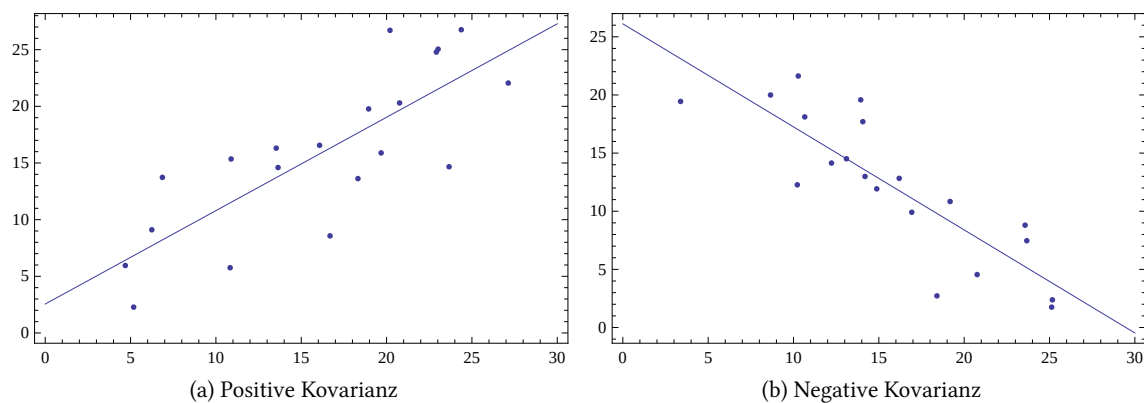


Abbildung 11.0.1: Punktwolke mit Regressionsgerade

Die Punktwolke ist eine quantitative graphische Darstellung der Wertepaare (x_i, y_i) zweier Stichproben X, Y . Hierbei wird für jedes Wertepaar ein Punkt (x_i, y_i) im kartesischen Koordinatensystem gezeichnet.

11.0.3.2 Kovarianz

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

Eine positive Kovarianz heißt die Regressionsgerade steigt und es herrscht ein gleichsinniger Zusammenhang.

Eine negative Kovarianz bedeutet die Regressionsgerade fällt, es herrscht daher ein gegensinniger Zusammenhang

Im Falle, dass die Kovarianz 0 ergibt, so ergibt sich kein linearer Zusammenhang zwischen den beiden Stichproben X und Y .

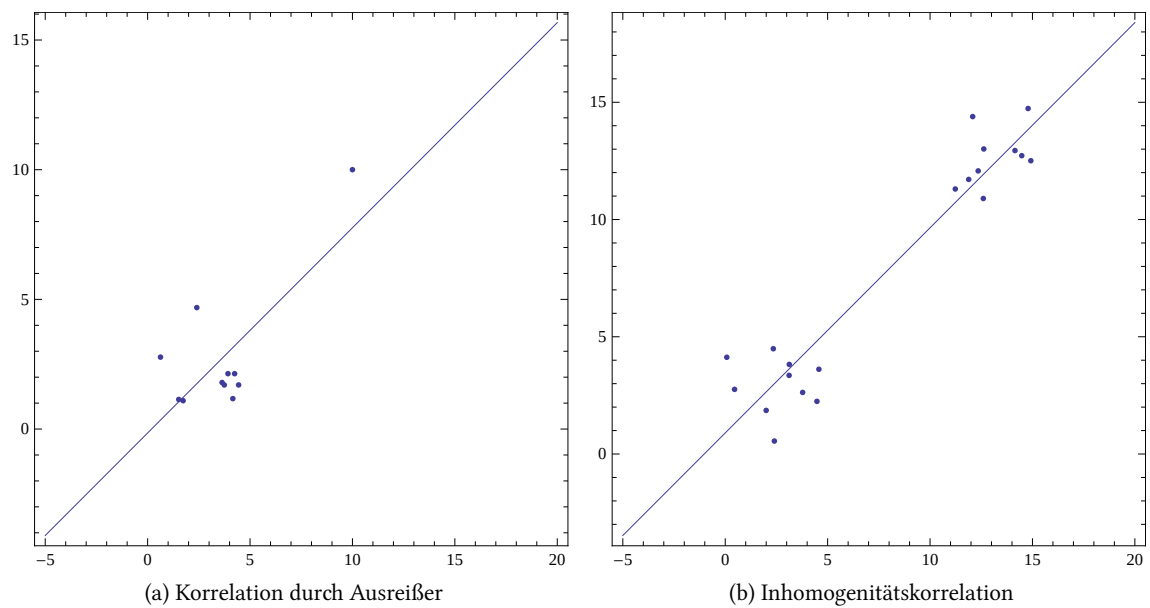


Abbildung 11.0.2: Arten der Scheinkorrelation

11.0.4 Pearson Koeffizient

Abschnitt 5.2.4

$$r = \frac{s_{xy}}{s_x \cdot s_y}$$

Für den Pearsonkoeffizient r gilt,

$$\begin{aligned} -1 &\leq r \leq 1 \\ 0 &\leq r^2 \leq 1 \end{aligned}$$

11.0.4.1 Interpretation des Korrelationskoeffizienten

Wird der Korrelationskoeffizient falsch, das heißt von dem Kontext der Werte getrennt, interpretiert, so kann es zu Scheinrelationen (Nonsensrelationen) kommen.

Formale Korrelation können entstehen, wenn sich zwei Merkmale immer zu einem Ganzen (100%) addieren lassen. Dieser Zusammenhang ist Mathematisch gegeben und es ergibt sich immer ein Korrelationskoeffizient von $r = -1$.

Selektions Korrelation In einer Stichprobe müssen alle Variationen von Merkmalen mit ähnlicher Wahrscheinlichkeit wie in der Grundgesamtheit vorkommen, damit die Stichprobe repräsentativ ist. Wird mit nicht repräsentativen Stichproben eine Korrelation errechnet, so muss diese nicht auch für die Grundgesamtheit gelten.

Korrelation durch Ausreißer Eine Ausreißer in der Stichprobe kann einen stark abweichenden Korrelationskoeffizienten verursachen. Ausreißer lassen sich in Punktwolken meist leicht erkennen, siehe Abbildung 11.0.2a.

Inhomogenitätskorrelation ergibt sich, wenn für zwei inhomogene Gruppen ein gemeinsamer Korrelationskoeffizient ermittelt wird. Auch hier lässt sich das bei einer Punktwolke leicht erkennen, siehe Abbildung 11.0.2b.

Gemeinsamkeitskorrelation liegt vor, wenn zwei Merkmale durch ein drittes beeinflusst werden. Wird – zum Beispiel – der Storchbestand eines Gebiets mit der Geburtenrate verglichen, so ergibt sich ein positiver Zusammenhang, obwohl jedem bekannt sein sollte, dass der Storch nicht die Kinder bringt. Der Grund ist eine dritte Größe, nämlich die allgemeine zeitliche Tendenz.

11.0.5 Regressionsanalyse

Abschnitt 5.3

Die Regressionsgerade in der Form

$$y = a + bx$$

Mit den Werten

$$b = \frac{s_{xy}}{s_x^2} \quad \wedge \quad a = \bar{y} - b\bar{x}$$

$$r^2 = \frac{s_{\hat{y}}^2}{s_y^2} = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} = \frac{\text{erklärte Varianz}}{\text{Gesamtvarianz}}$$

11.0.6 Spearman'scher Korrelationskoeffizient

Abschnitt 5.4.1

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

Für die Berechnung des Spearman Koeffizienten r_S werden alle x -Werte sortiert und mit Rangzahlen versehen. Ebenso wird so mit den y -Werten verfahren. Danach wird die Differenz d_i der Rangzahlen für jedes Wertepaar (x_i, y_i) ermittelt.

11.0.7 Klassifikation nach formalen Aspekten

Abschnitt 13.3

11.0.7.1 Deskriptiv vs. Analytisch

Deskriptive Studien sind rein beschreibend, sie werden zum Beispiel für Register (Krebsreg., Geburtenreg., Sterbereg.) sowie Fallberichte, Fallserien und Querschnittsstudien verwendet.

Analytische Studien werden durchgeführt, um den Zusammenhang mehrerer Einflussgrößen auf eine Zielgröße zu ermitteln. Bei sogenannten ökologischen¹ Studien werden verschiedene Register miteinander verknüpft, mit dem Ziel eine deskriptive Studie in eine analytische zu überführen.

11.0.7.2 Transversal vs. Longitudinal

Transversale Studien sind Querschnittsstudien, eine Momentaufnahme einer Population bei der eine oder Mehrere Eigenschaften der Studienteilnehmer erfasst werden. Eine einfache Form ist die Fallserie. Die Prävalenzstudie ist eine transversale Studie bei der die Prävalenz einer Krankheit zu einem bestimmten Zeitpunkt festgestellt wird. Transversale Studien sind meist deskriptiv.

Longitudinale Studien sind Längsschnittstudien, sie haben das Ziel, einen zeitlichen Verlauf zu beschreiben oder einen zeitlichen Zusammenhang zu ermitteln.

11.0.7.3 Retrospektiv vs. Prospektiv

Retrospektive Studien betrachten zunächst die Zielgröße und dann die Einflussgrößen. Ein Beispiel sind die Fall-Kontroll-Studien. Der Vorteil ist, dass diese Studien mit vorhandenen Daten schnell durchgeführt werden kann, jedoch der Nachteil liegt auf der Hand: oft schlechte Datenqualität bzw. Informationbias.

Prospektive Studien betrachten zunächst die Einflussgrößen und dann die Zielgrößen, hier gibt es eine Kontrollmöglichkeit bezüglich der Stichprobe. Der Vorteil ist die bessere Datenqualität, der Nachteil, der erhöhte Zeit- und Kostenaufwand, da die Daten hier durch Experimente (randomisierte klinische Studien) oder Kohorten Studien ermittelt werden.

11.0.7.4 Beobachtend vs. Experimentell

Beobachtend

Experimentell

11.0.7.5 Monozentrisch vs. Multizentrisch

Monozentrisch

Multizentrisch

¹vom Griechischen οἶκος (oikos) *Haus, Haushalt* und λόγος (logos) *Lehre*

11.0.7.6 Typische Studien

Risikostudien

Diagnosestudien

Therapiestudien

Präventionsstudien

11.0.8 Fehlerquellen und -möglichkeiten

11.0.8.1 Zufällige Fehler

11.0.8.2 Systematische Fehler

Bias als Systematischer Erfassungsfehler

Selektionsbias

Bias durch Confounder

11.0.9 Fall-Kontroll-Studien

Abschnitt 14.3

11.0.9.1 Matching

Abschnitt 14.3.3

11.0.9.2 Biasquellen

Abschnitt 14.3.4

11.0.9.3 κ -Koeffizient nach Cohen

Abschnitt 15.1.4

Der κ -Koeffizient wird verwendet um die interindividuelle² Variabilität und die intraindividuelle³ Variabilität zu messen.

$$\kappa = \frac{p_o - p_e}{1 - p_e}$$

p_o ... Beobachtete Wahrscheinlichkeit

p_e ... Erwarteter Anteil übereinstimmender Urteile

²Grad der Übereinstimmung zwischen zwei Beobachtern

³Grad der Übereinstimmung der z.B. Beurteilung desselben Beobachter zu zwei verschiedenen Zeitpunkten

		Beobachter A		
		1	2	3
Beobachter B	1	n_{11}	n_{12}	n_{13}
	2	n_{21}	n_{22}	n_{23}
	3	n_{31}	n_{32}	n_{33}
		$n_{\cdot 1}$	$n_{\cdot 2}$	$n_{\cdot 3}$

$$\begin{aligned}
 p_o &= \sum_{i=1}^k \frac{n_{ii}}{n} = \frac{1}{n} \sum_{i=1}^k n_{ii} \\
 e_{ii} &= \frac{n_{i\cdot}}{n} \cdot \frac{n_{\cdot i}}{n} \cdot n = \frac{n_{i\cdot} \cdot n_{\cdot i}}{n} \\
 p_e &= \sum_{i=1}^n \frac{e_{ii}}{n} = \frac{1}{n^2} \sum_{i=1}^n n_{i\cdot} \cdot n_{\cdot i}
 \end{aligned}$$

Interpretation des κ -Wertes Wenn $\kappa \leq 1$ gilt, so gibt es einen Zusammenhang:

$\kappa \leq 0.40$	schwache Übereinstimmung
$0.40 \leq \kappa \leq 0.60$	erkennbare Übereinstimmung
$0.60 \leq \kappa \leq 0.80$	gute Übereinstimmung
$0.80 \leq \kappa \leq 0.90$	exzellente Übereinstimmung
$0.90 \leq \kappa \leq 1.00$	perfekte Übereinstimmung

Bei $\kappa = 0$ ist kein Zusammenhang vorhanden. Es sind aber auch Werte mit $\kappa < 0$ möglich.