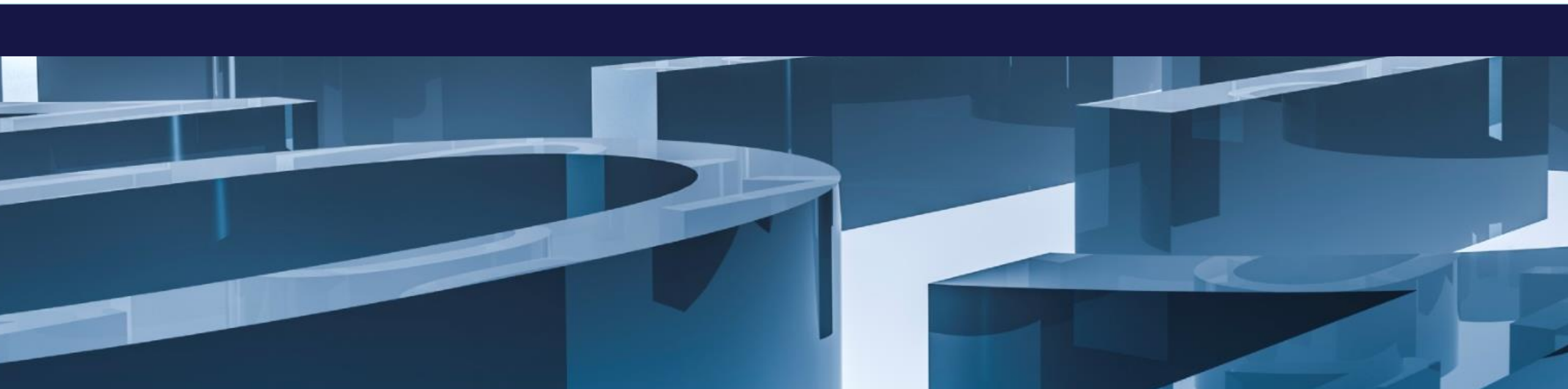




Statistische Grundlagen auf denen wir aufbauen

Allgemeines, Skalenniveau, Deskriptivstatistik, Korrelationen u.v.m.



Ablauf LV-Einheit 1

- Vorstellung und organisatorische Informationen (08:30-08:50)
- Einführung und Grundlagen der Statistik
 - Teil 1 von 08:50-09:50: Zweck der Statistik und Skalenniveaus, Fragebogenkonstruktion (inkl. Übungen; 20 Minuten Vortrag, 40 Minuten Übungen)
 - Pause 09:50-10:00
 - Teil 2 von 10:00-10:50: Graphische Darstellungen, Deskriptivstatistiken „in a rush“, Logik von Signifikanztests, Teststatistiken, Gerichtete und ungerichtete Hypothesen, Alpha-Fehler und Beta-Fehler, Konfidenzintervalle (50 Minuten Vortrag)
 - Pause 10:50-11:00 (falls Teil 2 rechtzeitig abgeschlossen wird)
 - Teil 3 von 11:00-11:30: Beispiele für inferenzstatistische Tests, quantitativer Forschungsprozess (inkl. Übung; 20 Minuten Vortrag, 10 Minuten Übung)

2 3 5 7 11 13 17 19 23 29 31 37 41 43 47 53 59 61 67 71 73
79 83 89 97 101 103 107 109 113 127 131 137 139 149 151
157 163 167 173 179 181 191 193 197 199 2 3 5 7 11 13 17
19 23 29 31 37 41 43 47 53 59 61 67 71 73 79 83 89 97 101
103 107 109 113 127 131 137 139 149 151 157 163 167 173
181 191 193 197 199 2 3 5 7 11 13 17 19 23 29 31 37 41 43
2 3 5 7 11 13 17 19 23 29 31 37 41 43 47 53 59 61 67 71 73
79 83 89 97 101 103 107 109 113 127 131 137 139 149 151
157 163 167 173 179 181 191 193 197 199 2 3 5 7 11 13 17
19 23 29 31 37 41 43 47 53 59 61 67 71 73 79 83 89 97 101
103 107 109 113 127 131 137 139 149 151 157 163 167 173
181 191 193 197 199 2 3 5 7 11 13 17 19 23 29 31 37 41 43
47 53 59 61 67 71 73 79 83 89 97 101 103 107 109 113 127
131 137 139 149 151 157 163 167 173 179 181 191 193 197
199 2 3 5 7 11 13 17 19 23 29 31 37 41 43 47 53 59 61 67 71

Zweck der Statistik und Skalenniveaus

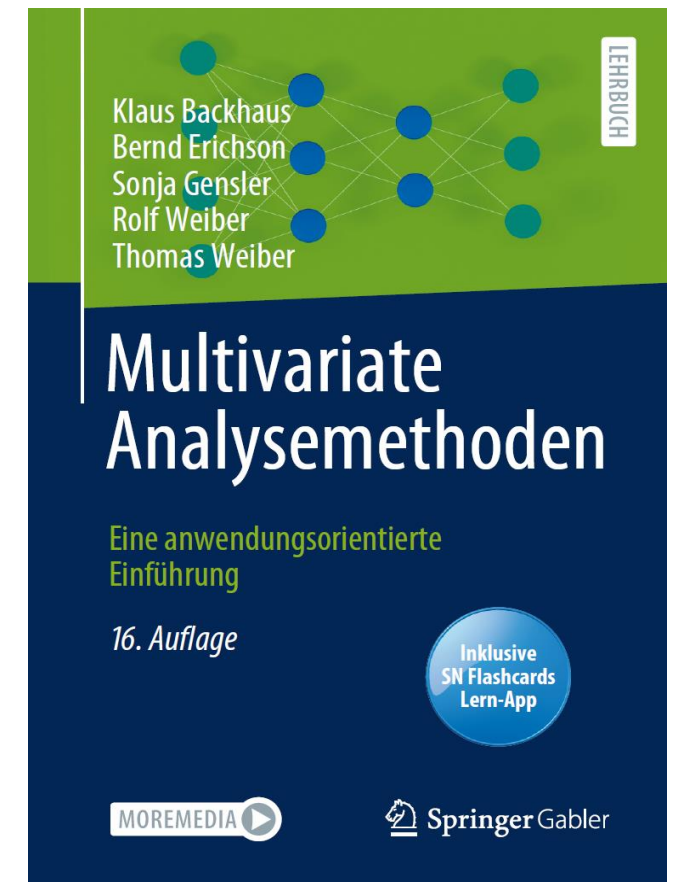
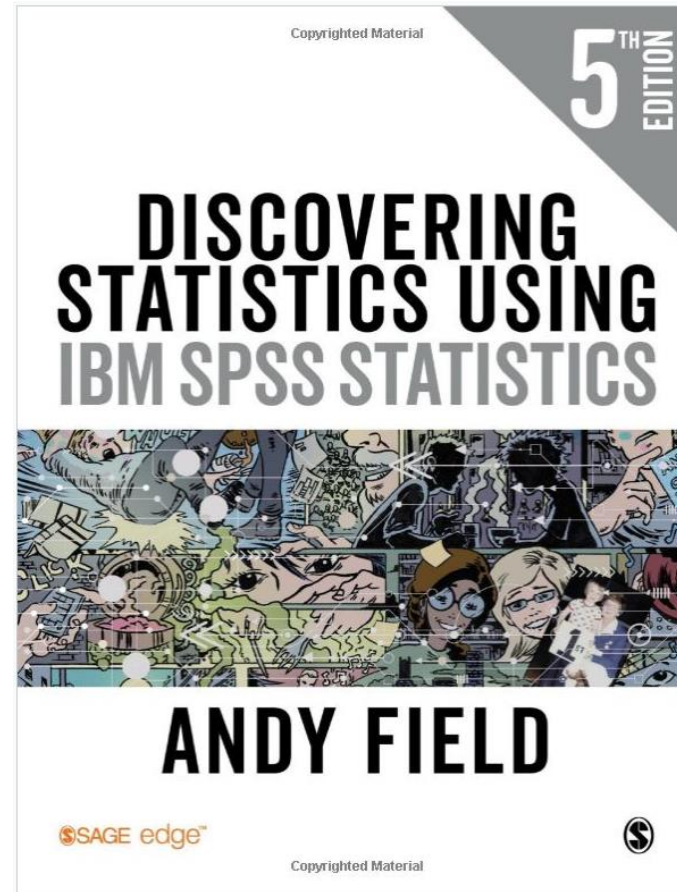
A: Für die Berechnung des Bestimmtheitsmaßes benötigt man jedenfalls metrisch skalierte Variablen.

versus

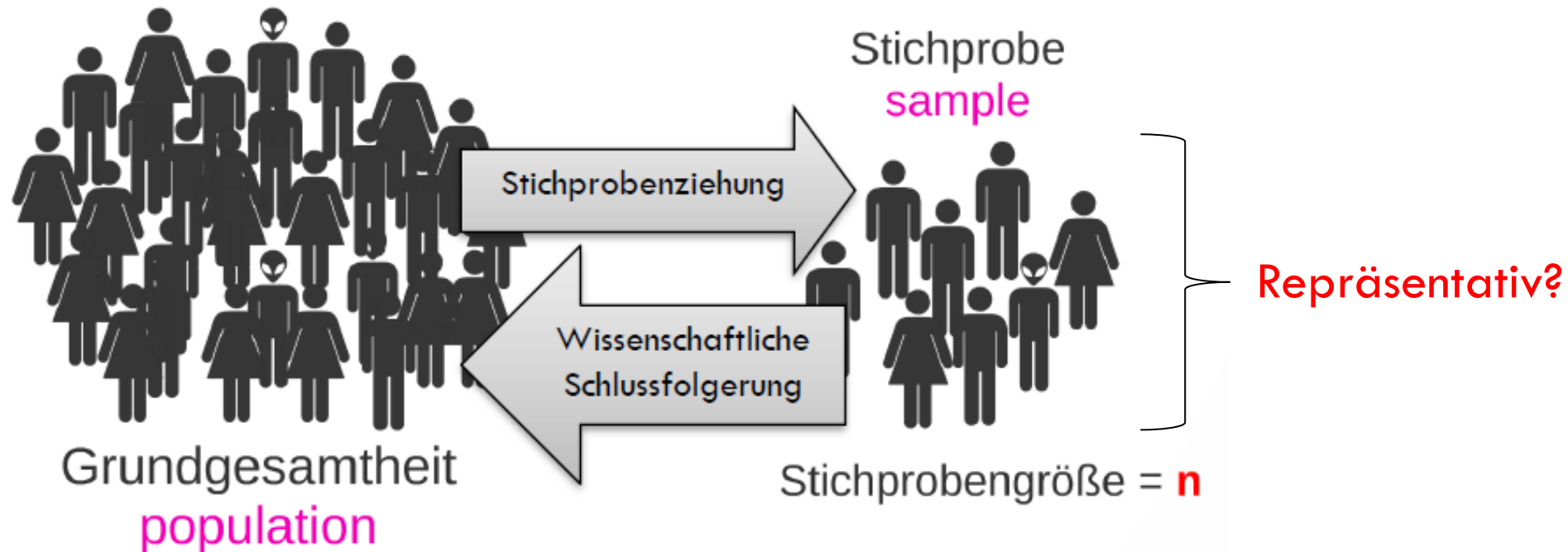
B: Auch mit dichotomen nominal- oder ordinal-skalierten Variablen ist eine Regression und folglich auch das Bestimmtheitsmaß sinnvoll berechenbar.

?

Persönliche Literaturempfehlungen



Stichprobe und Grundgesamtheit



Modelle: Brückenmetapher

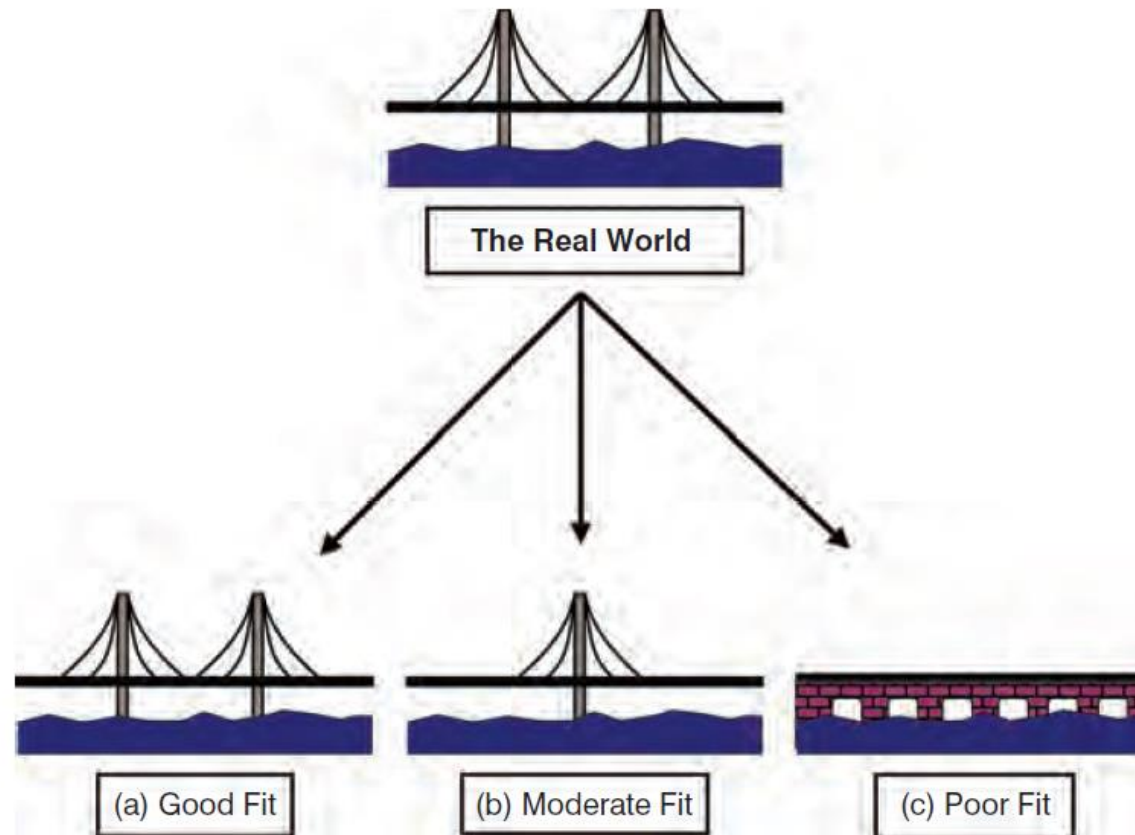


FIGURE 2.2
Fitting models
to real-world
data (see text
for details)

Modelle: Brückenmetapher

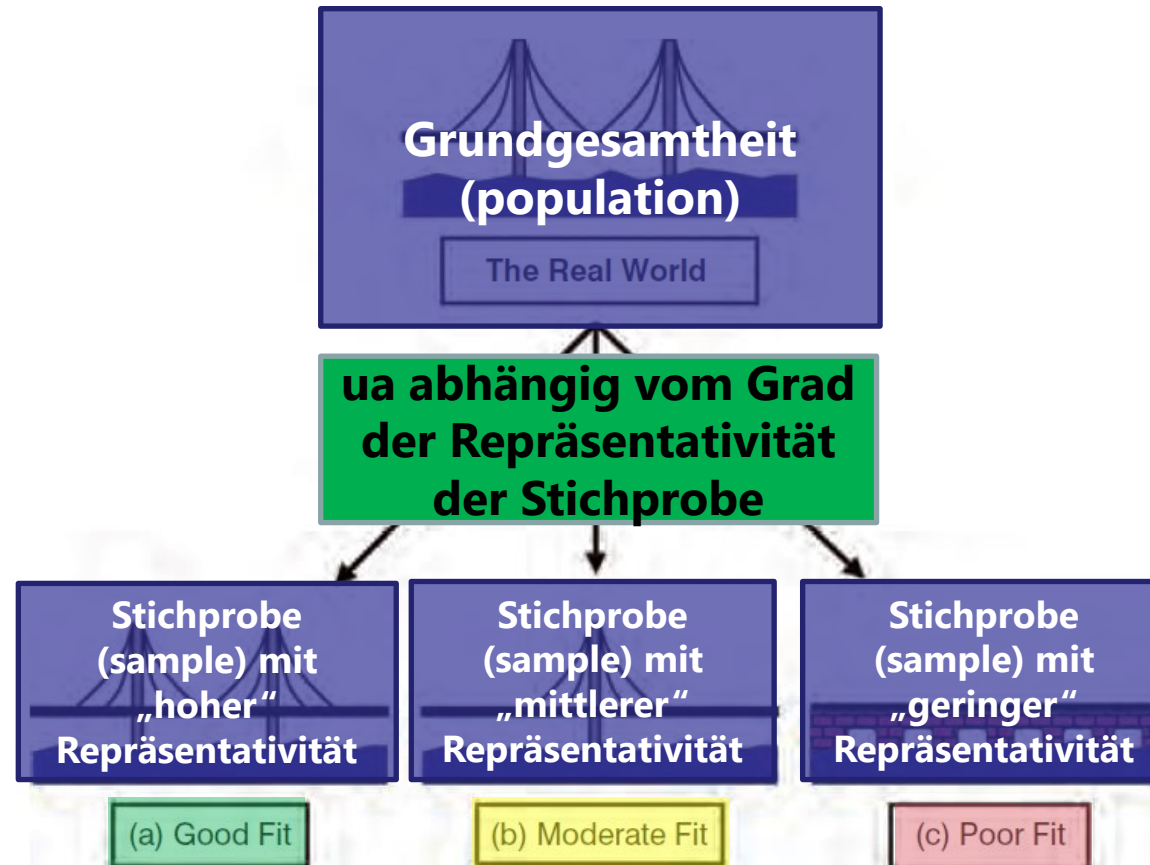


FIGURE 2.2
Fitting models
to real-world
data (see text
for details)

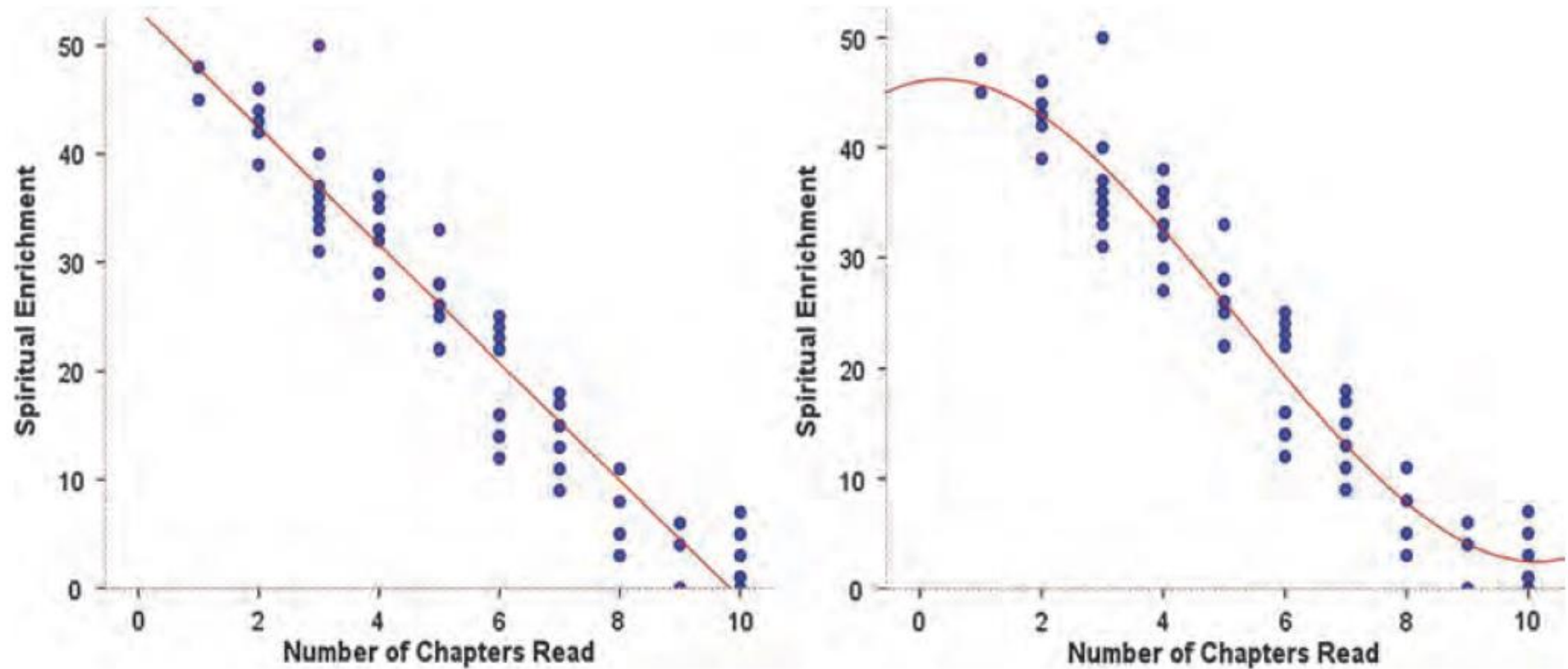
Modelle

- Ob man ein gutes Modell hat, hängt neben der Repräsentativität der Stichprobe aber auch ab von der Wahl des Modelltyps...

Modelle: linear oder doch nicht?

FIGURE 2.3

A scatterplot of the same data with a linear model fitted (left), and with a non-linear model fitted (right)



Notation Grundgesamtheit/Stichprobe

	Populationsparameter (griechische Buchstaben)	Stichprobenkennwerte (lateinische Buchstaben)
Arithmetischer Mittelwert	μ (<i>my</i>) z.B. durchschnittliche TV- Nutzungszeit in der Population	$M = \bar{x}$ z.B. durchschnittliche TV- Nutzungszeit in der Stichprobe
Relative Häufigkeit	π (<i>pi</i>) z.B. Anteil der Personen mit Depression in der Population	p z.B. Anteil der Personen mit Depression in der Stichprobe
Bivariates lineares Zusammenhangsmaß	ρ (<i>rho</i>) z.B. linearer Zusammenhang (bivariater Korrelationskoeffizient) zwischen TV-Nutzungszeit und Internet-Nutzungszeit in der Population	r z.B. linearer Zusammenhang (bivariater Korrelationskoeffizient) zwischen TV-Nutzungszeit und Internet-Nutzungszeit in der Stichprobe

Variablen – der essentielle Baustein für statistische Auswertungen

- Variablen kann man sich vorstellen wie Schubladen, die etwas enthalten, nämlich Daten
- Wichtiges Merkmal von Variablen: die Daten in der Schublade sind nicht gleich, sondern (zumindest teilweise) unterschiedlich
- Damit sind Variablen durch Veränderlichkeit gekennzeichnet

Variablen – der essentielle Baustein für statistische Auswertungen

- Beispiel: Hypothese: „Der Intelligenzquotient hängt mit der Einwohnerzahl eines Wohnortes zusammen“ und es wurden von insgesamt 100 Testpersonen in Niederösterreich sowohl der IQ als auch die Einwohnerzahl des Wohnortes mit dem klassischen Messinstrument = einem Fragebogen erhoben
- Es gibt hier 2 Variablen, nämlich den Intelligenzquotienten sowie die Größe des Wohnortes; in jeder Variable sind Werte enthalten, die zum Teil nicht gleich sind (theoretisch wäre das möglich, aber praktisch äußerst unwahrscheinlich)
- Die Variable Intelligenzquotient enthält die IQ-Werte je Testperson, die Variable Einwohnerzahl die Anzahl der Einwohner des Wohnortes pro Testperson
- ➔ damit könnte man jetzt bereits eine statistische Auswertung vornehmen – dazu aber später mehr!

Variablen – der essentielle Baustein für statistische Auswertungen

- Denkt man jetzt an klassische Fragebögen, dann kommen da unterschiedliche Fragen zum Einsatz
- Bei vielen dieser Fragen ist allerdings auf den ersten Blick nicht leicht erkennbar, wie man hier von der Beantwortung zu einer statistischen Auswertung kommt
- Beispiel:
- Der Fragebogen enthält die folgende Frage: Geschlecht: ☐ männlich ☐ weiblich ☐ divers
- Eine Testperson kreuzt jetzt z.B. „weiblich“ an – wie komme ich da zu einer statistischen Auswertung?
- In diesen und vielen ähnlichen Fällen, wo es textbasierte Antwortmöglichkeiten gibt, muss man für die einzelnen Antworten Zahlen hinterlegen (das macht die Person, die den Fragebogen erstellt, mit wenigen Ausnahmen – dazu noch später - bereits im Vorfeld der Erhebung!) – erst dann kann man damit etwas rechnen

Variablen – der essentielle Baustein für statistische Auswertungen

- Beispiel:
- Der Fragebogen enthält die folgende Frage: Geschlecht: ☐ männlich ☐ weiblich ☐ divers
- Eine Testperson kreuzt jetzt z.B. „weiblich“ an – wie komme ich da zu einer statistischen Auswertung?
- Nehmen wir einmal an, die Person, die den Fragebogen erstellt hat, hat beschlossen, für die Antwortoption „männlich“ eine 1 zu hinterlegen, für „weiblich“ die Zahl 2 und für „divers“ die Zahl 3
- Ok, Zahlen haben wir jetzt. Aber es stellt sich weiters die Frage, was ich mit den Zahlen rechnen darf bzw. was sinnvoll ist.
- Für das Beispiel: Angenommen, Testperson 1 gibt den Wert 1 an, und Testperson 2 den Wert 2. Macht es hier Sinn, eine Multiplikation vorzunehmen? Was würde dabei herauskommen? „männlich“ x „weiblich“ wird dann zu „meiblich“?
- Das macht hier keinen Sinn, d. h. es hängt auch von der Qualität der Zuordnung von den Zahlen zu den Antwortmöglichkeiten einer Variablen ab, was man damit rechnen darf oder anders ausgedrückt: es hängt davon ab, wie genau die Zahlen das repräsentieren, was man messen möchte – damit kommen wir zum sogenannten Skalenniveau

Variablen und Skalenniveaus

- **Kategorial**
 - › Nominalskala
 - › Ordinalskala
 - ... dichotom: nur 2 Ausprägungen

- **Numerisch**
 - › Intervallskala
 - › Rationalskala
 - › diskret vs. stetig

■ **Tab. 2.3** Überblick über die vier üblichen Skalenniveaus mit Angabe der jeweils relevanten (bedeutsamen) Relationen und Beispielen

Relation	Skala			
	Nominalskala	Ordinalskala	Intervallskala	Verhältnisskala
Verschiedenheit	ja	ja	ja	ja
Rangordnung	nein	ja	ja	ja
Differenzen	nein	nein	ja	ja
Verhältnisse	nein	nein	nein	ja
Beispiel	Geschlecht, Haarfarbe	Schulnote, Tabellenplatz	Intelligenzquotient	Reaktionszeit, Körpergröße

■ **Tab. 2.3** Überblick über die vier üblichen Skalenniveaus mit Angabe der jeweils relevanten (bedeutsamen) Relationen und Beispielen

Relation	Skala			
	Nominalskala	Ordinalskala	Intervallskala	Verhältnisskala
Verschiedenheit	ja	ja	ja	ja
Rangordnung	nein	ja	ja	ja
Differenzen	nein	nein	ja	ja
Verhältnisse	nein	nein	nein	ja
Beispiel	Geschlecht, Haarfarbe	Schulnote, Tabellenplatz	Intelligenzquotient	Reaktionszeit, Körpergröße

Zuordnung von Zahlen zu Ausprägungen der Variablen sind beliebig und kennzeichnen lediglich eine Verschiedenheit

Beispiel

Variable Geschlecht: 1 = weiblich, 2 = männlich, 3 = divers; möglich wäre aber auch 7 = weiblich, 5 = männlich, 9 = divers

Skalenniveaus

■ **Tab. 2.3** Überblick über die vier üblichen Skalenniveaus mit Angabe der jeweils relevanten (bedeutsamen) Relationen und Beispielen

Relation	Skala			
	Nominalskala	Ordinalskala	Intervallskala	Verhältnisskala
Verschiedenheit	ja	ja	ja	ja
Rangordnung	nein	ja	ja	ja
Differenzen	nein	nein	ja	ja
Verhältnisse	nein	nein	nein	ja
Beispiel	Geschlecht, Haarfarbe	Schulnote, Tabellenplatz	Intelligenzquotient	Reaktionszeit, Körpergröße

Zuordnung von Zahlen zu Ausprägungen der Variablen sind **größenmäßig nicht mehr beliebig** und kennzeichnen eine Rangordnung

Beispiel

Variable Schulnote: 1 = sehr gut, 2 = gut, 3 = befriedigend, 4 = genügend, 5 = nicht genügend

Nicht mehr sinnvoll ist hier eine beliebige Zuordnung, etwa:

1 = sehr gut, 5 = gut, 2 = befriedigend, 4 = genügend, 3 = nicht genügend

Skalenniveaus

■ **Tab. 2.3** Überblick über die vier üblichen Skalenniveaus mit Angabe der jeweils relevanten (bedeutsamen) Relationen und Beispielen

Relation	Skala			
	Nominalskala	Ordinalskala	Intervallskala	Verhältnisskala
Verschiedenheit	ja	ja	ja	ja
Rangordnung	nein	ja	ja	ja
Differenzen	nein	nein	ja	ja
Verhältnisse	nein	nein	nein	ja
Beispiel	Geschlecht, Haarfarbe	Schulnote, Tabellenplatz	Intelligenzquotient	Reaktionszeit, Körpergröße

Zuordnung von Zahlen zu Ausprägungen der Variablen sind nicht mehr beliebig und kennzeichnen nicht nur eine Rangordnung, es lassen sich auch die Differenzen sinnvoll interpretieren

Beispiel

Variable Intelligenzquotient: Person 1 = IQ von 110, Person 2 = IQ von 120, Person 3 = IQ von 90, Person 4 = IQ von 100

→ Die Differenz zwischen Person 2 und Person 1 = 10; ebenso die Differenz zwischen Person 4 und Person 3 = 10

Skalenniveaus

■ **Tab. 2.3** Überblick über die vier üblichen Skalenniveaus mit Angabe der jeweils relevanten (bedeutsamen) Relationen und Beispielen

Relation	Skala			
	Nominalskala	Ordinalskala	Intervallskala	Verhältnisskala
Verschiedenheit	ja	ja	ja	ja
Rangordnung	nein	ja	ja	ja
Differenzen	nein	nein	ja	ja
Verhältnisse	nein	nein	nein	ja
Beispiel	Geschlecht, Haarfarbe	Schulnote, Tabellenplatz	Intelligenzquotient	Reaktionszeit, Körpergröße

Zuordnung von Zahlen zu Ausprägungen der Variablen sind nicht mehr beliebig und kennzeichnen nicht nur eine Rangordnung, es lassen sich auch die Differenzen und die Verhältnisse sinnvoll interpretieren; es gibt auch einen natürlichen Nullpunkt (z.B. kann die Variable Körpergröße oder Reaktionszeit keine Werte kleiner 0 aufweisen)

Beispiel

Variable Reaktionszeit: Person 1 = 20 ms, Person 2 = 40 ms, Person 3 = 10 ms, Person 4 = 30 ms

→ Die Differenz zwischen Person 2 und Person 1 = 20; ebenso die Differenz zwischen Person 4 und Person 3 = 20

→ Die Reaktionszeit etwa von Person 1 ist halb so groß wie von Person 2 und doppelt so groß wie von Person 3

Skalenniveaus

Was man hier wieder sehen kann: oftmals müssen Variablen Zahlenwerte für Antworten in Form von Text zugeordnet werden, z. B. der Variable Geschlecht: hier weist man den textbasierten Antwortmöglichkeiten („weiblich“, „männlich“, „divers“) Zahlen zu (also 1 für weiblich, 2 für männlich und 3 für divers)

Oder auch bei den Schulnoten: den ursprünglich textbasierten Beurteilungen („sehr gut“, „gut“, „befriedigend“, „genügend“, „nicht genügend“) werden auch Zahlen zugewiesen (also 1, 2, 3, 4, 5)

Dieses Vorgehen der Zuordnung von textbasierten Antwortmöglichkeiten zu Zahlenwerten nennt man Kodierung!

Beispiel Fragebogen: hier kreuzt die befragte Person zwar zum Beispiel am Fragebogen die Antwortoption „weiblich“ an, jedoch ist für die Personen, die den Fragebogen nachher auswerten, insbesondere auch die Zahl wichtig, die der Antwortoption zugeordnet ist

Die Kodierung ist wichtig, um nachher statistische Berechnungen durchführen zu können!

Skalenniveaus

Aber nicht alle Variablen benötigen eine solche Kodierung:

Beispiel: Es könnte im Fragebogen auch die folgende Frage vorkommen:

„Alter (in ganzen Jahren): ____“

Hier tragen die befragten Personen gleich Zahlenwerte ein, mit denen man nachher statistische Auswertungen vornehmen kann → diese Variable könnte man den Variablennamen ALTER geben; es handelt sich um eine verhältnisskalierte Variable

Oder aber eine Kodierung ist erst nach der Umfrageerhebung überhaupt erst möglich:

Beispiel: Oftmals findet sich auch am Fragebogen ganz am Ende die folgende letzte Frage:

„Falls Sie Anmerkungen zur Umfrage haben, dann können Sie das jetzt im Textfeld unterhalb tun:“

In diesem Beispiel tragen die befragten Personen Textzeichen ein (das nennt man auch string); es handelt sich damit um nominalskalierte Variablen, für die man z.B. für jeden Aspekt, der hier erwähnt wird, einen beliebigen Zahlenwert vergibt (wie auch bei der Variable Geschlecht); im Unterschied aber etwa zur Variable Geschlecht macht man das erst im Nachhinein, also wenn die Personen die Umfrage schon ausgefüllt haben, geht die Person an die Auswertung und bestimmt im Nachhinein, wie diese Antworten kodiert werden sollen

SKALENNIVEAU

Nominalskala.....

Ordinalskala.....

Intervallskala.....

Rationalskala.....

BEISPIEL

Geschlecht, Religion

Gefallen am Produkt

Temperatur in Grad Celsius

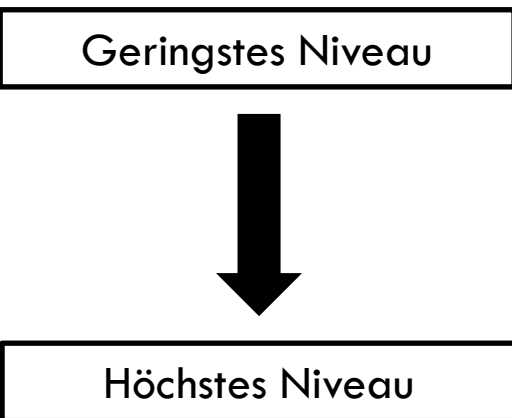
Hautleitfähigkeitswert
(kontinuierlich), Alter (diskret)

Geringstes Niveau



Höchstes Niveau

Skalenniveaus (in SPSS)



SKALENNIVEAU

Nominalskala.....

Ordinalskala.....

Metrisch.....

BEISPIEL

Geschlecht, Religion

Gefallen am Produkt

{

 Temperatur in Grad Celsius

 Hautleitfähigkeitswert

 (kontinuierlich), Alter (diskret)

Übung 1 (15 Minuten)

Kodierplan

Auf Moodle finden Sie einen Übungsfragebogen: **LV_Einheit_1_ Übung_1_Fragebogen.pdf**.

Legen Sie zu diesem Fragebogen einen *Datenkodierungsplan an, indem Sie sich für jede Frage bzw. jedes Item folgende Punkte überlegen:

- Einen möglichst kurzen aber aussagekräftigen Variablennamen
- Variablentyp: numerisch (enthält die Variable Zahlenwerte?) oder string (enthält die Variable Textzeichen?)
- Falls Variablen numerisch sind und von vornherein keine numerischen Werte einzugeben sind: welche Zahlen erhalten die textbasierten Antwortoptionen? (z.B. stimme voll zu – stimme eher zu – weder/noch - stimme eher nicht zu – stimme gar nicht zu → welche Zahlen werden den Antwortoptionen zugeordnet?)
- Skalenniveau (nominal/ ordinal/ intervall/ rational)?

*Eine Vorlage für den Datenkodierungsplan finden Sie auch auf Moodle (Excel-File):

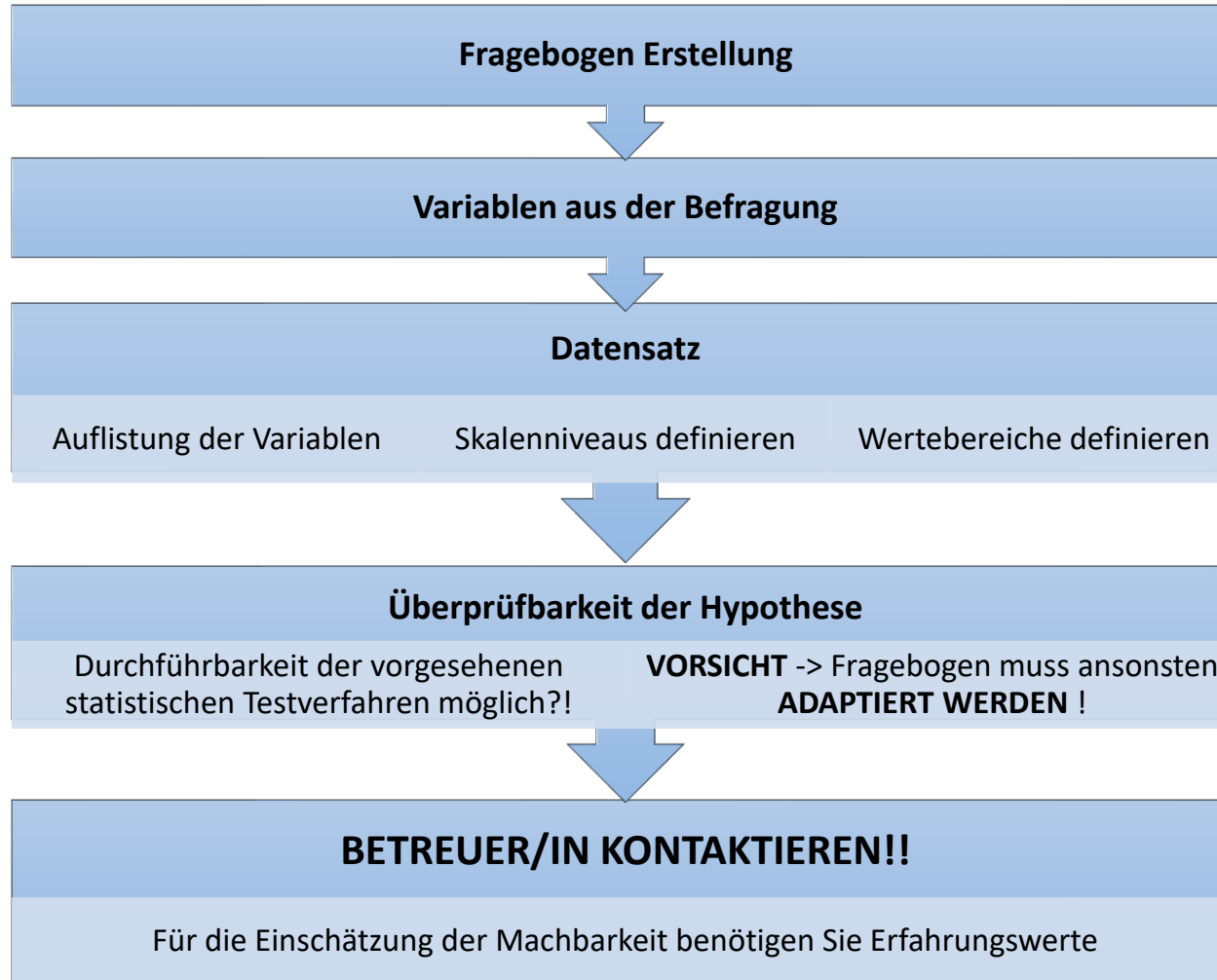
LV_Einheit_1_ Übung_1_Fragebogen.xlsx

Übung 1

Kodierplan

Fragennummer	Fragenwortlaut	Variablenname	Variablentyp	Kodierung (falls textbasiertes Antwortformat)	Skalenniveau
1	Wohnen Sie in Österreich?	WOHNORT	NUMERISCH	JA = 1; NEIN = 2	NOMINAL (DICHOTOM)
2	Was denken Sie über das Thema Armut in Österreich:	MEINUNG_ARMUT	STRING		NOMINAL (KODIERUNG ERFOLGT IM NACHHINEIN)
3	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis Ihrer eigenen Entscheidungen.	ARMUT_SKALA_1	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
4	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis Ihrer eigenen Einstellungen.	ARMUT_SKALA_2	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
5	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis Ihrer eigenen Fähigkeiten und Fertigkeiten.	ARMUT_SKALA_3	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
6	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis Ihrer Prioritäten.	ARMUT_SKALA_4	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
7	Wenn Sie hart genug arbeiten, können Sie aus der Armut herauskommen.	ARMUT_SKALA_5	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
8	Wenn Sie motiviert genug sind, können Sie aus der Armut herauskommen.	ARMUT_SKALA_6	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
9	Wenn Sie bereit sind, Unterstützung von Freunden und Familie anzunehmen, können Sie aus der Armut herauskommen.	ARMUT_SKALA_7	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
10	Wenn Sie bereit sind, Unterstützung von Sozialdiensten anzunehmen, können Sie aus der Armut herauskommen.	ARMUT_SKALA_8	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
11	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis von Problemen in unserem Bildungssystem.	ARMUT_SKALA_9	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
12	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis von Problemen in unserem Regierungssystem.	ARMUT_SKALA_10	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
13	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis von Problemen in unserem Wirtschaftssystem.	ARMUT_SKALA_11	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
14	Wenn Sie in Österreich von Armut betroffen sind, ist dies das Ergebnis des Verhaltens aller in der Gesellschaft.	ARMUT_SKALA_12	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
15	Wohlhabende Menschen sind für das Ausmaß der Armut in Österreich verantwortlich.	ARMUT_SKALA_13	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
16	Maßnahmen der Regierung sind für das Ausmaß der Armut in Österreich verantwortlich.	ARMUT_SKALA_14	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
17	Die Regierungspolitik ist für das Ausmaß der Armut in Österreich verantwortlich.	ARMUT_SKALA_15	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
18	Diskriminierung führt in Österreich zu Armut.	ARMUT_SKALA_16	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
19	Richtlinien in Bezug auf den Arbeitsplatz (z. B. Mindestlohn, Kinderbetreuung, Gesundheitsfürsorge) verursachen in Österreich Armut.	ARMUT_SKALA_17	NUMERISCH	1 = STIMME GAR NICHT ZU; 2 = STIMME EHER NICHT ZU; 3= WEDER/NOCH; 4 = STIMME EHER ZU; 5 = STIMME VOLL ZU	ORDINAL (ODER PSEUDOINTERVALLSKALIERT)
20	Bitte geben Sie Ihr Geschlecht an:	GESCHLECHT	NUMERISCH	1 = MÄNNLICH; 2 = WEIBLICH; 3 = DIVERS	NOMINAL
21	Alter (Angabe in Jahren):	ALTER	NUMERISCH		RATIONAL
22	Wie hoch ist Ihr monatliches Nettoeinkommen (in €)?	EINKOMMEN	NUMERISCH	< 500 = 1; 500-1000 = 2; 1001-1500 = 3; 1501-2000 = 4; 2001-2500 = 5; 2501-3000 = 6; >-3000 = 7	ORDINAL

Datenstruktur



Konstruktion und Implementierung



Einzelfrage vs. Fragebatterien

Filterfragen ?!

- wo werden sie platziert ?

Demografische Daten am Ende des Fragebogens

Antwortmöglichkeiten und Skalen

- deutlich, angemessen und erschöpfend

Einheitlich definierte Skalen

- Hohe Werte = hohe Zustimmung. (Schulnotenskala = Ausnahme)

Datenkodierung mitbedenken!

Frageformate

multiple/single Choice



- Am häufigste verwendet
- Eine oder mehrere Antwortoptionen

Textfeld



- Offene Fragen mittels Texteingabe
- Kann für Formularfelder genutzt werden
- Konkrete Formatbeschränkung

Rangordnung



- Liste nach Präferenzen /Priorität
- Liste in randomisierter Reihenfolge

Matrix Tabelle



- Liste an Aussagen/ Fragen/Produkte
- Spezifizierung der Skala/ Skalenwerte
- Verwendung bei Likert-skalierten Items

Schieberegler



- Verwendet für Einschätzung/ Zustimmung

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

1. einfache, eindeutige Begriffe, die von allen Befragten in gleicher Weise verstanden werden

- *Ich bin in Gesprächen angriffslustig.*
(trifft nicht zu 1-2-3-4-5 trifft zu)
 - a) *Ich vertrete meine Meinung offensiv*
 - b) *Ich bin angriffslustig*
- *Ich fühle mich depressiv.*
(trifft nicht zu 1-2-3-4-5 trifft zu)
- *vs. Ich fühle mich oft traurig.*

2. Vermeide Frage, die **unwahrscheinlich beantwortbar** sind

- *Sind in ihrer Gemeinde Maßnahmen zur Umsetzung der lokalen Agenda 21 getroffen worden?*
- Um welche Maßnahmen geht es?

Relevante Themen gezielt abfragen!



3. Definiere **unklare Begriffe**

- *Wie viele Kinder haben Sie?**
(*leibliche Kinder, auch solche die nicht im selben Haushalt leben, adoptierte Kinder, Kinder in Obsorge...)
- *Wie viel verdienen Sie im Monat? **vs.** Wie hoch ist Ihr monatliches Nettoeinkommen?*

Eindeutigkeit vor Einfachheit!

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

4. Vermeide **lange, komplexe oder verschachtelte** Fragen

- *„Wie Sie wissen, sind manche Leute politisch ziemlich aktiv, andere Leute finden dagegen oft keine Zeit oder haben kein Interesse, sich an politischen Dingen aktiv zu beteiligen. Bitte sagen Sie mir wie oft Sie persönlich so etwas tun bzw. wie häufig das bei Ihnen vorkommt.“*

vs. *Wie häufig nehmen Sie an einer Diskussion zu politischen Themen teil?*

- **nicht mehr als 20 Wörter**
- **keine redundanten/überflüssigen Wörter**
- **Verständlichkeit**

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

5. Vermeide **hypothetische Fragen**

- „Stellen Sie sich einmal vor, Sie wären verheiratet und hätten einen 16 jährigen Sohn, der seine Lehre abbrechen möchte, um Fußballprofi zu werden. Würden Sie ihn in diesem Wunsch unterstützen oder würden Sie ihm raten, zuerst seine Ausbildung zu Ende zu bringen?“

Hineinversetzten gelingt bei:

- **gedankliche Auseinandersetzung mit Situation**
- **Realitätsnähe/-entfernung der hypothetischen Situation**

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

6. Vermeide **doppelte Verneinungen**

- *Ich bin nicht oft traurig.*
(trifft nicht zu 1-2-3-4-5 trifft zu)
- *Es ist **nicht** gut, wenn Kinder ihren Eltern **nicht** gehorchen.*
(trifft nicht zu 1-2-3-4-5 trifft zu)

Verwechslungsgefahr bei Antwortoptionen:
Was bedeutet Zustimmung?

7. Lege jedem Item nur einen Gedanken zugrunde

- *Ich höre gerne Musik von Chopin und Calvin Harris.*
(trifft nicht zu 1-2-3-4-5 trifft zu)
- *Ich fahre sehr gerne und sehr schnell Auto.*
(trifft nicht zu 1-2-3-4-5 trifft zu)

Zerlegen in Einzelitems!



8. Vermeide Verallgemeinerungen

- *Alle ManagerInnen sind skrupellos.*
(trifft nicht zu 1-2-3-4-5 trifft zu)

9. Verwende Fragen mit eindeutigem zeitlichen Bezug

- *In den letzten Wochen war ich häufig niedergeschlagen.*
(trifft nicht zu 1-2-3-4-5 trifft zu)

Was bedeutet häufig?

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

10. Antwortkategorien müssen disjunkt sein

- *Wie viele Vorträge zum Thema „Gesundes Leben“ haben Sie im Jahre 2000 bisher gehört?*
– „keinen“ – „einen Vortrag“ – „zwei bis fünf Vorträge“ – „fünf Vorträge oder mehr“

versus

- „keinen“ – „einen Vortrag“ – „**zwei bis einschließlich vier** Vorträge“ – „**fünf Vorträge oder mehr**“

**Antwortmöglichkeiten schließen sich aus.
Befragte können sich zweifelsfrei einer
Antwortmöglichkeit zuordnen.**

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

11. Antwortkategorien müssen erschöpfend sein

- *Wie viele Stunden beschäftigen Sie sich in einer normalen Arbeitswoche mit der Entwicklung von Fragebögen?*
 - „überhaupt nicht“ – „bis unter 3 Stunden“ – „3 bis unter 5 Stunden“ – „5 Stunden bis unter 10 Stunden“
 - „überhaupt nicht“ – „bis unter 3 Stunden“ – „3 bis unter 5 Stunden“ – „5 Stunden bis unter 10 Stunden“ – „**10 Stunden oder mehr**“

Vollständigkeit der Antwortmöglichkeiten: alle möglichen Antworten der Befragten müssen abgedeckt sein.

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

12. Bedenke mögliche Einflüsse von Vorfragen auf Folgefragen

- *Haben Sie in den Medien von den Problemen erfahren, die die Einführung der Zentralmatura erschwert haben?*
(trifft nicht zu 1-2-3-4-5 trifft zu)
- *Wie stehen Sie der Zentralmatura gegenüber?*
(1-10, 1= ablehnend, 10=befürwortend)

Die Vorfrage -> ruft Probleme in Erinnerung -> schlechtere Zustimmungswerte für Folgefrage!

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

13. Vermeide Suggestivfragen

- *Führende Wissenschaftler sind der Ansicht, dass Autoabgase das Wachstum von Kindern hemmen können. Stimmen Sie dieser Ansicht zu?*

Alternative

- *Autoabgase hemmen das Wachstum von Kindern. (stimme zu 1-2-3-4-5 stimme nicht zu) “*

- **Gefahr konformer Beantwortung**
- **Bewusstes vertreten der gegenteiligen Auffassung!**

Fragebogenkonstruktion-Regeln (Bühner 2011, Kapitel 3)

14. Vermeide Items und Antworten wie „immer“/ “alle“/ “keiner“/ “niemals“

- *Ich gehe **nie** bei Rot über die Straße.*
- *Ich bin **immer** bereit, anderen Menschen zu helfen.*

Zustimmung zu einem Item dieser Art ist eher Hinweis auf soziale Erwünschtheit.



15. Vermeide quantifizierende Umschreibungen mit Häufigkeitsratings

- *Wie oft gehen Sie ins Kino? (nie-selten-gelegentlich-oft)*
- *Alternative Antwortkategorien: (trifft nicht zu 1-2-3-4-5 trifft zu)*

QUALTRICS

Nach Abschluss der Fragebogenkonstruktion wird der Fragebogen in ein Online-Umfrage-Tool eingegeben.

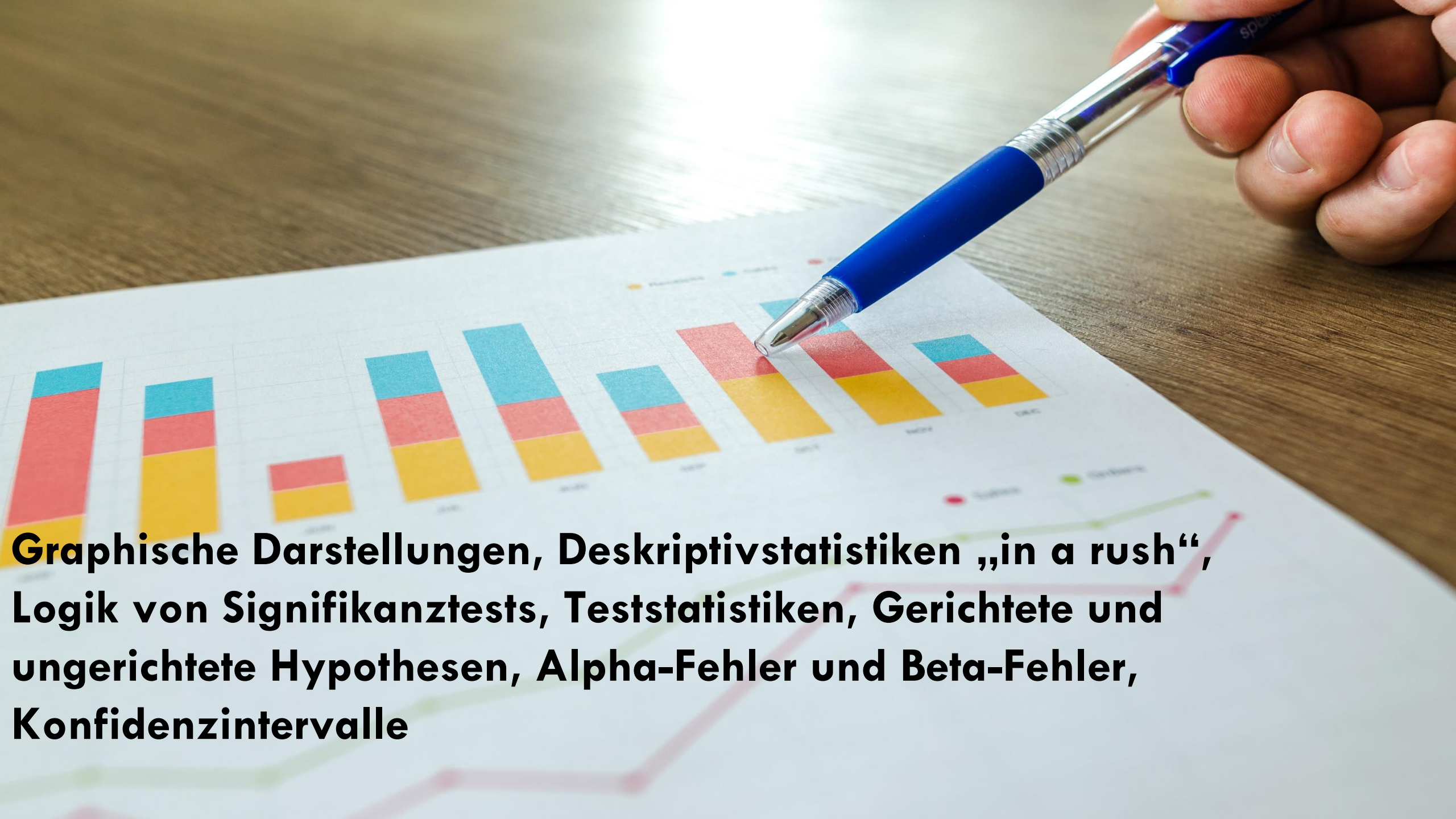
Übung 2 (25 Minuten)

Eingabe in Qualtrics

Öffnen Sie das folgende File (zu finden auf Moodle): **LV_Einheit_1_ Übung_1_Fragebogen.pdf**

Probieren Sie, den Fragebogen in Qualtrics (fhwn.qualtrics.com) anzulegen

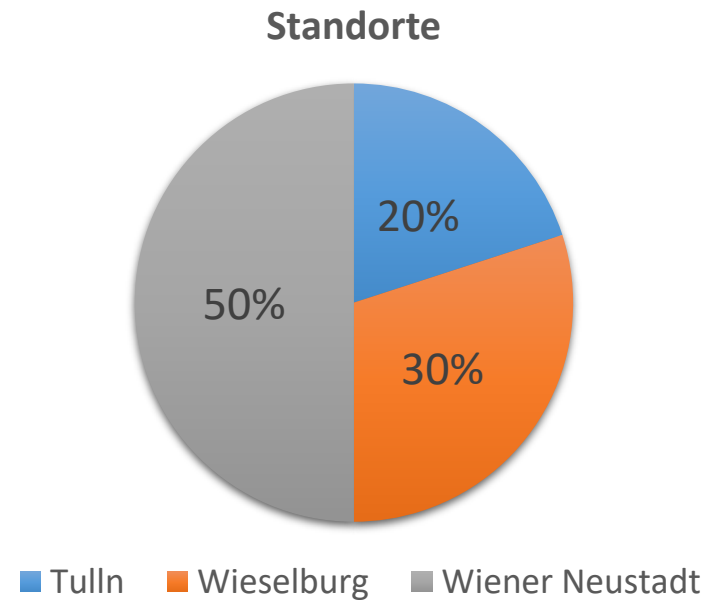




**Graphische Darstellungen, Deskriptivstatistiken „in a rush“,
Logik von Signifikanztests, Teststatistiken, Gerichtete und
ungerichtete Hypothesen, Alpha-Fehler und Beta-Fehler,
Konfidenzintervalle**

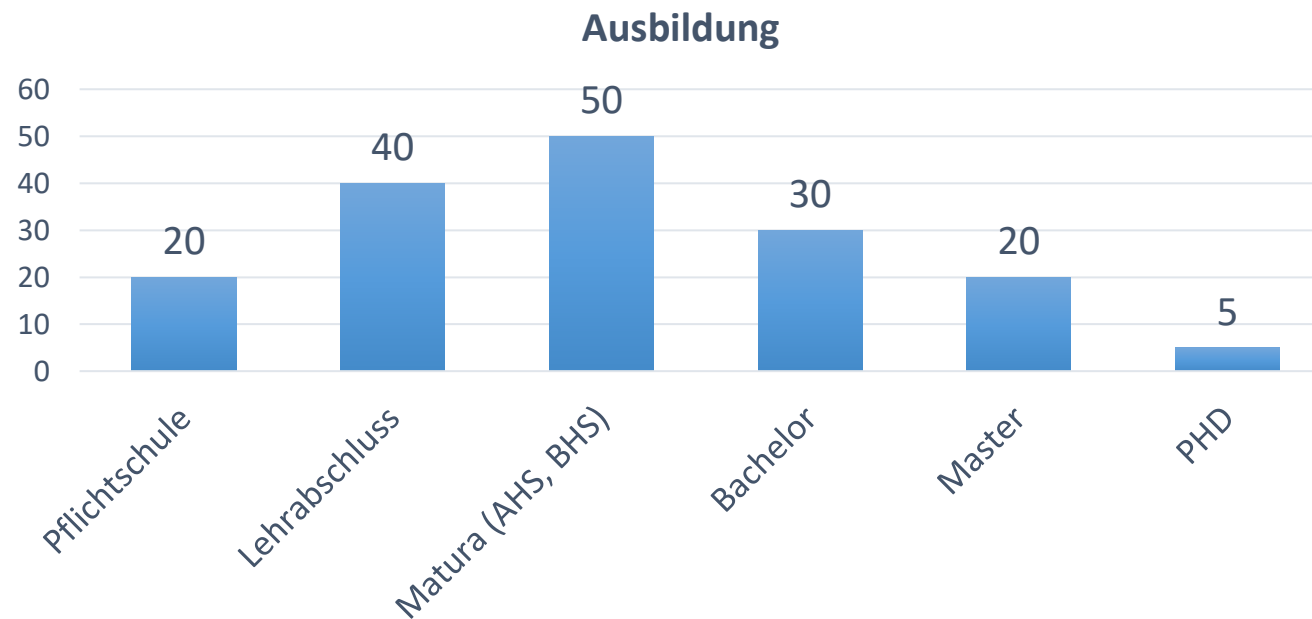
Grafische Darstellung

- Kreisdiagramm
 - Geeignet für nominal skalierte Variablen



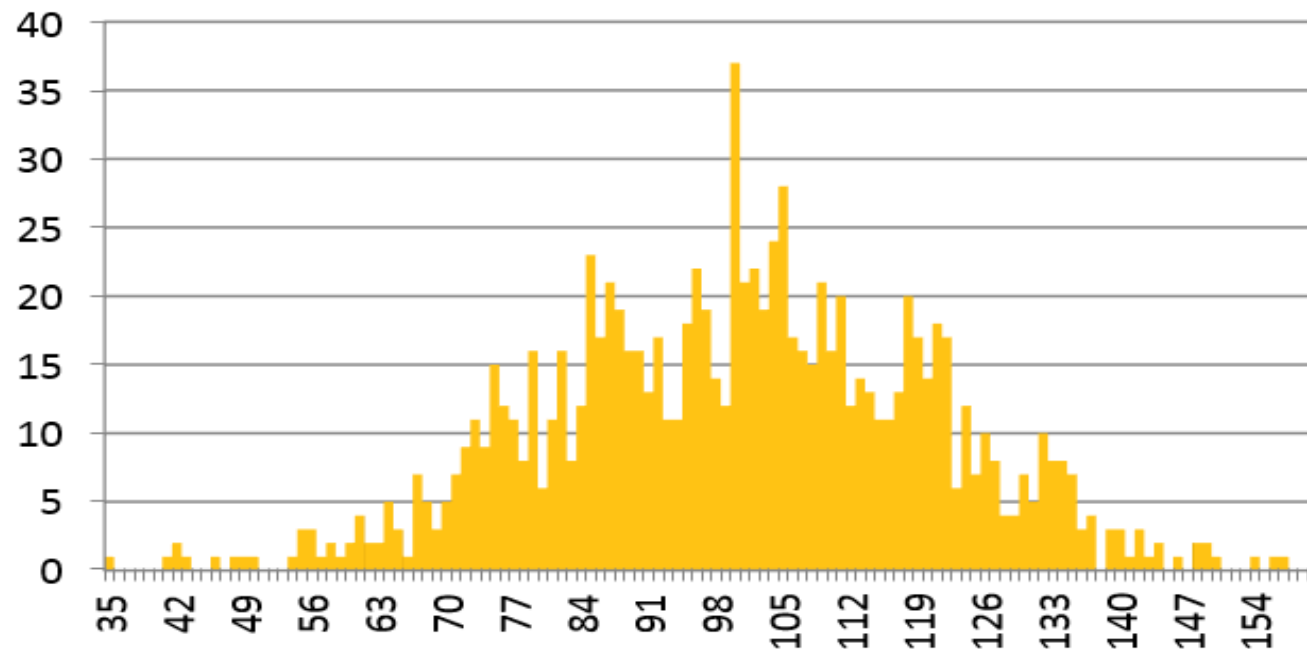
Grafische Darstellung

- Balkendiagramm
 - Geeignet für nominal oder ordinalskalierte Variablen, Reihenfolge bleibt erhalten
 - Häufigkeiten oder Prozente auf der Y-Achse



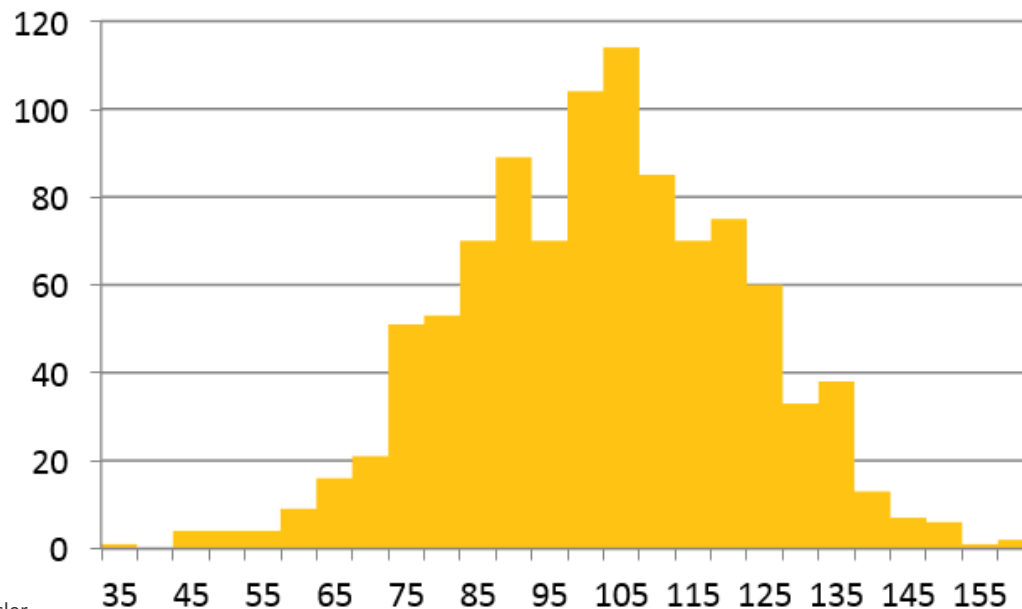
Grafische Darstellung

- Balkendiagramm
 - Für stetige Variablen ungeeignet



Grafische Darstellung

- Histogramm
 - Für die Darstellung von stetigen, metrischen Daten
 - Klassenzusammenfassung in Intervalle
 - nur bei metrischen Variablen zulässig



Lagemaße

- Mittelwert
 - Stichprobe: $\bar{x}$
 - Grundgesamtheit: μ
- Median ($\tilde{x}$)
- Modalwert (x_D)

Lagemaße/ Mittelwert

- Mittelwert

- Stichprobe: $\bar{x}$
- Grundgesamtheit: μ

- Formel:
$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

- Formel in Worten: summiere alle Werte auf die es gibt und dividiere anschließend durch die Anzahl der Werte

Lagemaße/ Mittelwert

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Berechnung:
 1. Summiere alle Werte auf, die es gibt:
 $57+40+103+234+93+53+116+98+108+121+22 = 1045$
 2. Dividiere anschließend durch die Anzahl der Werte:
 $1045/11 = 95$

Lagemaße/ Median

- Median ($\tilde{x}$)
- Der Wert einer Variable, der die der Größe nach geordnete Gesamtzahl der Werte halbiert (50% liegen unterhalb, 50% oberhalb des Medians)
- Vorgehensweise:
 1. Sortierung der Werte aufsteigend nach der Größe
 2. $(\text{Anzahl der Werte } (n) + 1)/2 \rightarrow$ ergibt die Position des Medians (etwas unterschiedlich in der Ermittlung, abhängig von der Anzahl der Werte)

Lagemaße/ Median

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22

1. Sortierung der Werte aufsteigend nach der Größe:

22; 40; 53; 57; 93; 98; 103; 108; 116; 121; 234

Lagemaße/ Median

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22

2. $(\text{Anzahl der Werte } (n) + 1)/2 = (11 + 1)/2 = 12/2 = 6 \rightarrow \text{Median}$
an 6. Position = 98

Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234

Lagemaße/ Median

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F10):
F1: 57; F2: 40; F3: 103; F4: 93; F5: 53; F6: 116; F7: 98; F8: 108; F9: 121; F10: 22

1. Sortierung der Werte aufsteigend nach der Größe:

22; 40; 53; 57; 93; 98; 103; 108; 116; 121

Lagemaße/ Median

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F10):
F1: 57; F2: 40; F3: 103; F4: 93; F5: 53; F6: 116; F7: 98; F8: 108; F9: 121; F10: 22

2. $(\text{Anzahl der Werte } (n) + 1)/2 = (10+1)/2 = 11/2 = 5,5 \rightarrow$
Median zwischen 5. und 6. Position, Ermittlung durch Bildung des Durchschnitts der zwei Werte an 5. und 6. Position

Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234

Lagemaße/ Median

- Rechenbeispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F10):
F1: 57; F2: 40; F3: 103; F4: 93; F5: 53; F6: 116; F7: 98; F8: 108; F9: 121; F10: 22

2. $(93+98)/2 = 95,5 = \text{Median}$

Lagemaße/ Modalwert

- Modalwert (x_D)
- Ist der am häufigsten vorkommende Wert in einer Variable
- Beispiel Noten

Note	Häufigkeit
1	3
2	5
3	10
4	4
5	2

Lagemaße/ Modalwert

- Beispiel Noten bei $N = 24$ (unimodal)

Note	Häufigkeit
1	3
2	5
3	10
4	4
5	2

Lagemaße/ Modalwert

- Beispiel Noten bei $N = 24$ (bimodal)

Note	Häufigkeit
1	1
2	8
3	8
4	2
5	5

Lagemaße/ Modalwert

- Beispiel Noten bei $N = 24$ (multimodal)

Note	Häufigkeit
1	1
2	2
3	7
4	7
5	7

Streuungsmaße

- Minimum (*Min*) und Maximum (*Max*); Streubereich $[x_1; x_n]$
- Spannweite (*R*, *Range*)
- Quartile (Q_p) und Interquartilsabstand (*IQR*)
- Perzentile/Quantile (Q_p)
- Varianz (σ^2)
- Standardabweichung (σ)

Streuungsmaße/ Minimum, Maximum, Streubereich

- Minimum = kleinster Wert
- Maximum = größter Wert
- Streubereich = Bereich, der vom kleinsten bis zum größten Wert reicht

Streuungsmaße/ Minimum, Maximum, Streubereich

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Minimum = 22
- Maximum = 234
- Streubereich [22;234]

Streuungsmaße/ Spannweite

- Spannweite = Größter Wert – kleinster Wert

Streuungsmaße/ Spannweite

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Spannweite = $234 - 22 = 212$

Streuungsmaße/ Quartile und Interquartilsabstand

- Quartile (Q_p)
 - Quartile sind die drei Schnittstellen, die eine geordnete Folge von Werten in vier gleich große Teile aufteilen
 - 2. Quartil = Median
 - 1. Quartil = Median der unteren 50% der Werte
 - 3. Quartil = Median der oberen 50% der Werte
- Interquartilsabstand (IQR)
 - Spannweite der mittleren 50% der Werte ohne die extremsten 25% der unteren und 25% der oberen Werte
 - Berechnung: 3. Quartil – 1. Quartil

Streuungsmaße/ Quartile und Interquartilsabstand

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Bestimmung des **2. Quartils = Median** (siehe dazu auch die Folie weiter oben):

Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234

Streuungsmaße/ Quartile und Interquartilsabstand

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Bestimmung des 1. Quartils = Median der unteren 50% der Werte = $(5+1)/2 = 3$. Position = 53

Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234

Streuungsmaße/ Quartile und Interquartilsabstand

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Bestimmung des **3. Quartils = Median der oberen 50% der Werte** = $((11-6)+1)/2 = 3$. Position (von der 6. Position aus) = 116

Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234

Streuungsmaße/ Quartile und Interquartilsabstand

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):

F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22

- Berechnung des Interquartilsabstands: 3. Quartil – 1. Quartil = $116 - 53 = 63$

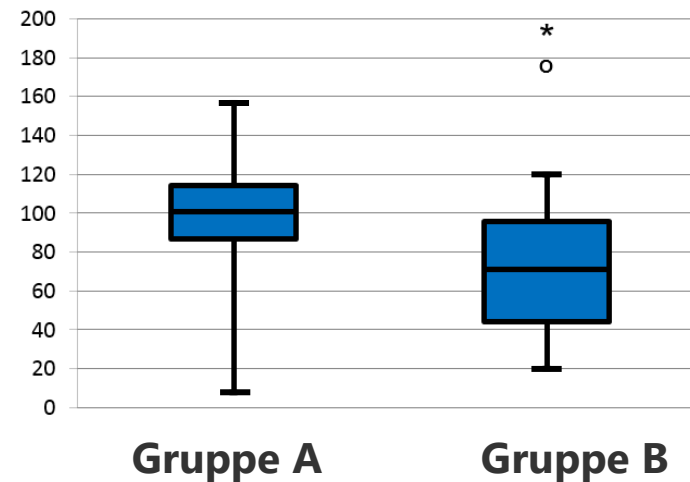


Position	1.	2.	3.	4.	5.	6.	7.	8.	9.	10.	11.
Wert	22	40	53	57	93	98	103	108	116	121	234
Quartil			1.			2.			3.		

Grafische Darstellung

■ Boxplot

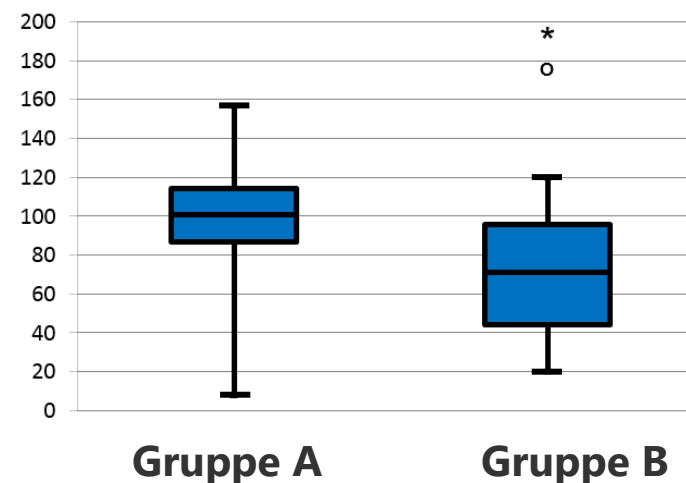
- Im Box-Plot werden verschiedenste Lage- und Streuungsmaße in einem Diagramm zusammengefasst
- Der oberste und unterste Wertepunkt gibt die Grenze der Spannweite an, wenn es keine Ausreißer gibt



Grafische Darstellung

■ Boxplot

- Die Box in der Mitte wird unten von $Q_{.25}$ und oben von $Q_{.75}$ begrenzt. Dazwischen liegen also 50% aller Werte (IQR). Der breite Querstrich innerhalb der Box gibt den Median an
- Ausreißer, die eine bestimmte Distanz von den Quartilen entfernt sind, werden mit Kreisen, noch weiter entfernte Ausreißer mit Sternen (= Extremwerte), angegeben



Grafische Darstellung

■ Boxplot

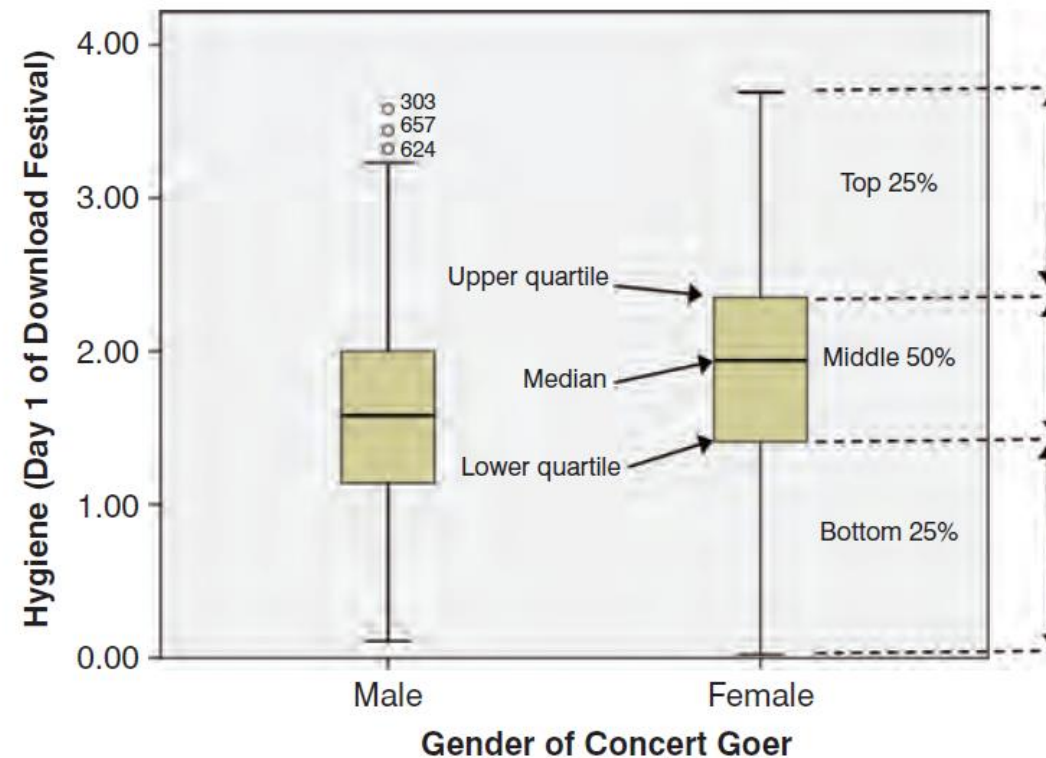


FIGURE 4.14
Boxplot of
hygiene scores
on day 1 of
the Download
Festival split by
gender

Streuungsmaße/ Perzentile, Quantile

- Perzentile/Quantile (Q_p)
- Quantile teilen eine geordnete Folge von Werten in gleich große Teile auf
 - Quartile sind damit ein Spezialfall der Quantile
 - Z.B. 10%-Quantil der Variable Gehalt: EUR 1200 → 10% aller Variablenwerte sind kleiner oder gleich 1200

Streuungsmaße/ Varianz, Standardabweichung

- Varianz (σ^2) und Standardabweichung (σ)
- Beim Interquartilsabstand als Streumaß werden nur 50% aller Werte berücksichtigt
- Kann man ein Streumaß finden, in dem alle Werte in die Berechnung eingehen?

Streuungsmaße/ Varianz, Standardabweichung

- Varianz (σ^2) und Standardabweichung (σ)
- Man könnte sich folgendes überlegen: nimmt man den Mittelwert als Lagemaß, dann kann man die Differenz jedes einzelnen Wertes vom Mittelwert berechnen, also:

$$x_i - \bar{x}$$

Streuungsmaße/ Varianz, Standardabweichung

- Beispiel: erhoben wurde die Anzahl der Freunde folgender Facebook-NutzerInnen (F1-F11):
F1: 57; F2: 40; F3: 103; F4: 234; F5: 93; F6: 53; F7: 116; F8: 98; F9: 108; F10: 121; F11: 22
- Man könnte sich folgendes überlegen: nimmt man den Mittelwert als Lagemaß, dann kann man die Differenz jedes einzelnen Wertes vom Mittelwert berechnen, also:

$$x_i - \bar{x}$$

Streuungsmaße/ Varianz, Standardabweichung

Anzahl der Freunde (x_i)	Mittelwert ($\bar{x}$)	Abweichung ($x_i - \bar{x}$)
22	95	-73
40	95	-55
53	95	-42
57	95	-38
93	95	-2
98	95	3
103	95	8
108	95	13
116	95	21
121	95	26
234	95	139

Streuungsmaße/ Varianz, Standardabweichung

- Damit man jetzt eine Idee von der gesamten Streuung hat, könnte man jetzt diese einzelnen Differenzen einfach aufsummieren...

Streuungsmaße/ Varianz, Standardabweichung

Anzahl der Freunde (x_i)	Mittelwert ($\bar{x}$)	Abweichung ($x_i - \bar{x}$)
22	95	-73
40	95	-55
53	95	-42
57	95	-38
93	95	-2
98	95	3
103	95	8
108	95	13
116	95	21
121	95	26
234	95	139
Summe		$\sum_{i=1}^n x_i - \bar{x} = 0$

Streuungsmaße/ Varianz, Standardabweichung

- ...das führt aber zum Wert 0! Demnach wäre also die gesamte Abweichung 0 – das stimmt aber so nicht
- Das Problem ist deshalb vorhanden, weil es negative und positive Abweichungen gibt
- Um dieses Problem zu lösen, quadriert man die Differenzen, sodass man nur noch positive Werte erhält...

Streuungsmaße/ Varianz, Standardabweichung

Anzahl der Freunde (x_i)	Mittelwert ($\bar{x}$)	Abweichung ($x_i - \bar{x}$)	Abweichung quadriert ($(x_i - \bar{x})^2$)
22	95	-73	5329
40	95	-55	3025
53	95	-42	1764
57	95	-38	1444
93	95	-2	4
98	95	3	9
103	95	8	64
108	95	13	169
116	95	21	441
121	95	26	676
234	95	139	19321
Summe		$\sum_{i=1}^n x_i - \bar{x} = 0$	$\sum_{i=1}^n (x_i - \bar{x})^2 = 32246$

Streuungsmaße/ Varianz, Standardabweichung

- Damit erhält man die quadrierte Abweichungssumme
- Hier hat man aber erneut ein Problem: die Größe dieses Wertes hängt von der Anzahl der in die Berechnung eingegangenen Werte ab, daher kann man diese Maßzahl nicht für Stichprobenvergleiche heranziehen
- Man wählt daher die durchschnittliche Abweichungsquadratsumme, oder anders ausgedrückt: die Varianz!

Streuungsmaße/ Varianz, Standardabweichung

- Die Populationsvarianz, an der man ja grundsätzlich interessiert ist, liest sich damit folgendermaßen:

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

- ...wobei, ganz korrekt müsste die Formel mit Blick auf die vorhin vorgestellte Notation ja eigentlich lauten...

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2$$

Streuungsmaße/ Varianz, Standardabweichung

- Da wir also hier den Populationsmittelwert nicht zur Verfügung haben, müssen wir auf den Mittelwert der Stichprobendaten zurückgreifen; da wir damit also in der Formel einen Wert „herbeibemühen“ (=den Mittelwert), muss man von der Anzahl der Beobachtungen auch den Wert von 1 abziehen und erhält die korrigierte Stichprobenvarianz:

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

Streuungsmaße/ Varianz, Standardabweichung

- Für das Beispiel:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = 32246/(11-1) = 3224,6$$

Streuungsmaße/ Varianz, Standardabweichung

- Ein letztes Problem in diesem Zusammenhang: bei der Varianz erhält man quadrierte Werte, die man nur schlecht interpretieren kann; besser ist es daher, das Quadrieren „rückgängig“ zu machen, und das erreicht man, indem man die Wurzel zieht – damit erhält man die korrigierte Standardabweichung:

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

Streuungsmaße/ Varianz, Standardabweichung

- Für das Beispiel:

$$\sqrt{s^2} = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2} = \sqrt{3224,6} = 56,79$$

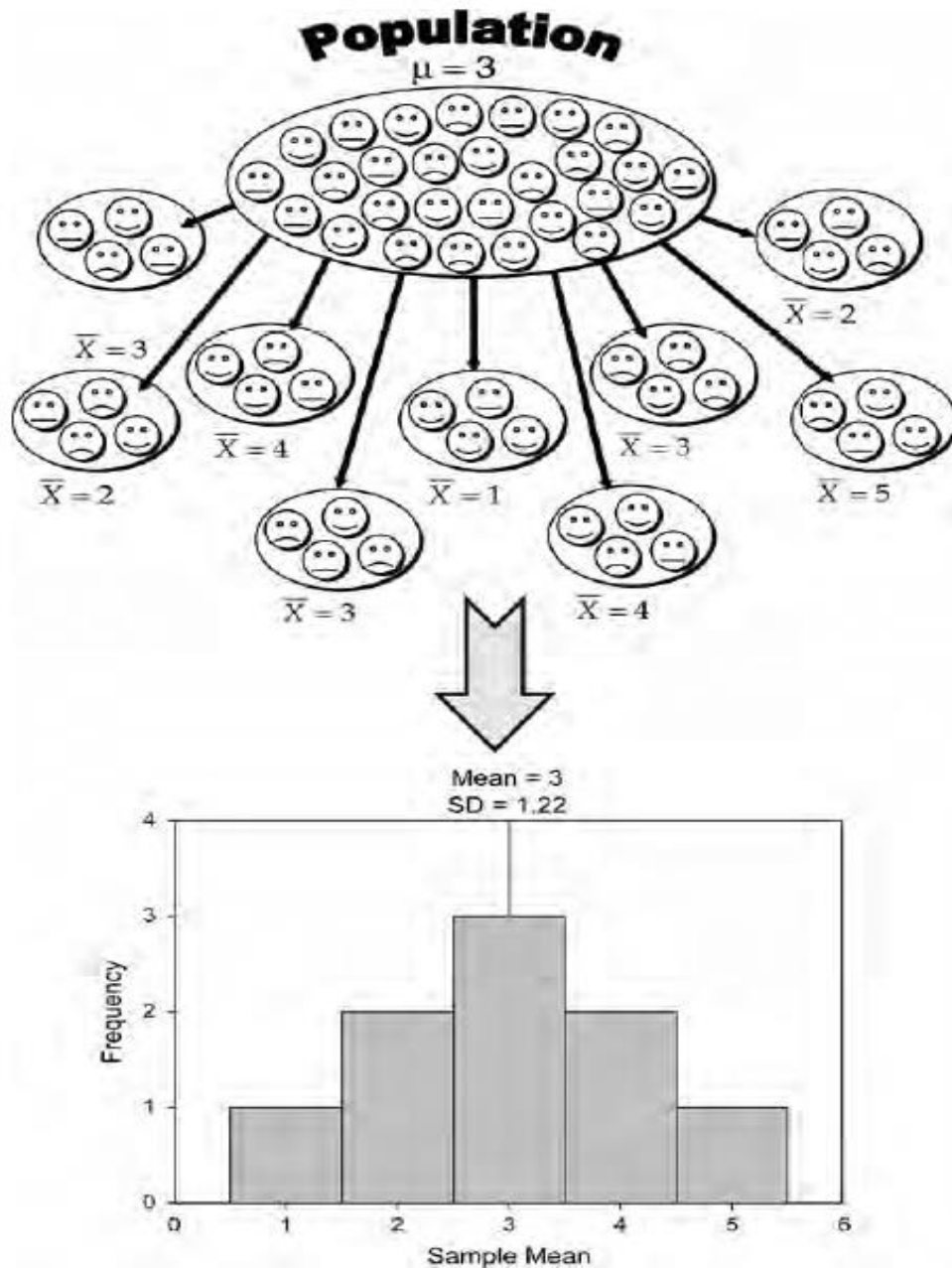
- Die nicht bekannte Standardabweichung der Grundgesamtheit wird geschrieben mit:

$$\sigma = \sqrt{\sigma^2}$$

Standardfehler des Mittelwertes (SE)

- Der Fehler bei Generalisierung eines Mittelwerts auf die Population
- Wird bei Konfidenzintervallen und statistischen Tests verwendet
- $SE = \frac{s}{\sqrt{n}}$, *je größer n , desto kleiner der SE*

FIGURE 2.7
Illustration of the
standard error
(see text for
details)



Zentrales Grenzwerttheorem:

„Die Verteilung von Mittelwerten aus Stichproben des Umfangs n , die derselben Grundgesamtheit entnommen wurden, geht mit wachsendem Stichprobenumfang in eine Normalverteilung über. [...] In der Literatur gibt es unterschiedliche Aussagen hinsichtlich des Stichprobenumfangs, der für beliebige Verteilungsformen benötigt wird, um von einer normalverteilten Mittelwertverteilung ausgehen zu können. Häufig wird $n > 30$ als notwendige Voraussetzung genannt.“ (Bortz und Schuster 2010, 86 und 87)

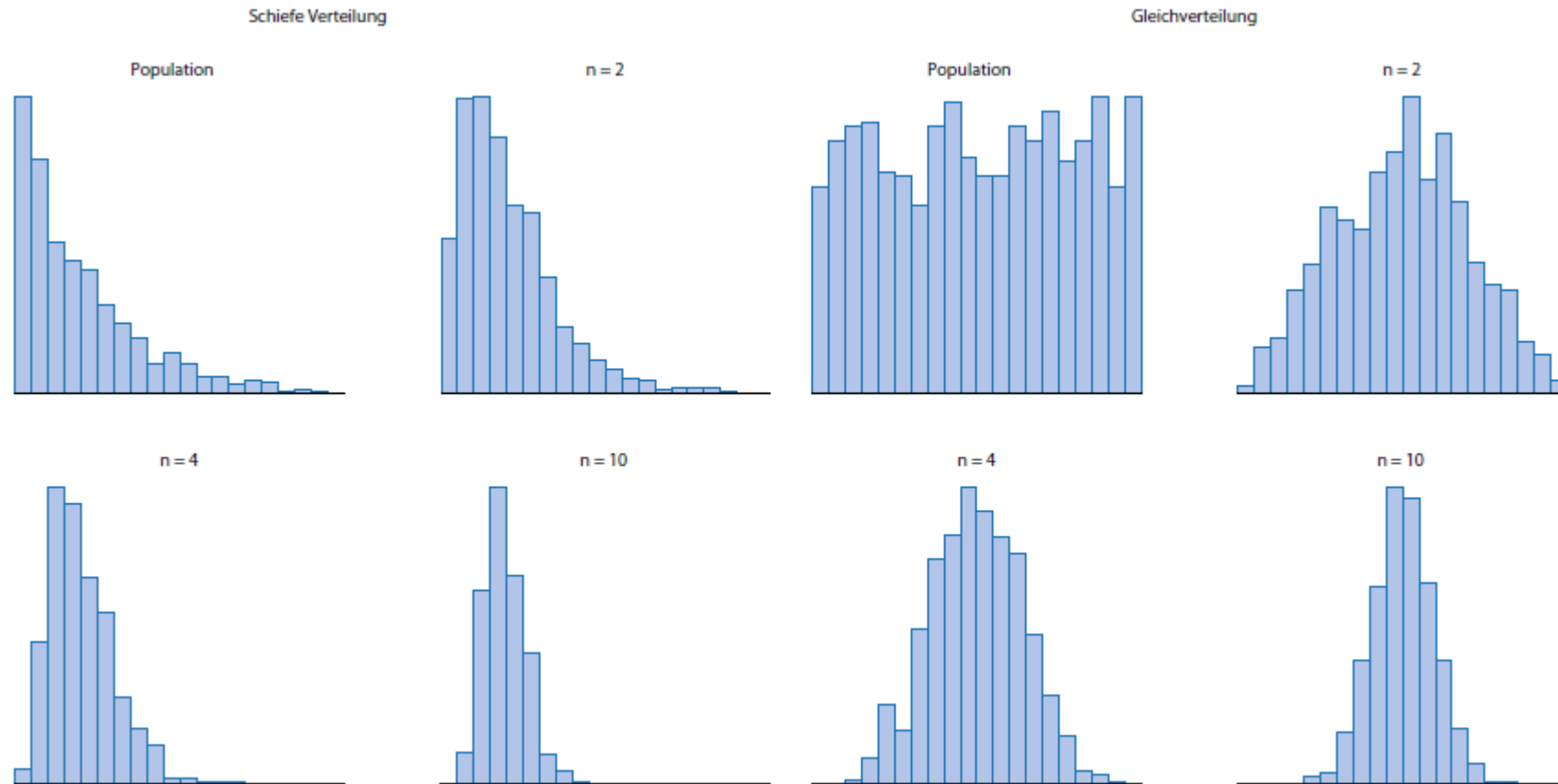
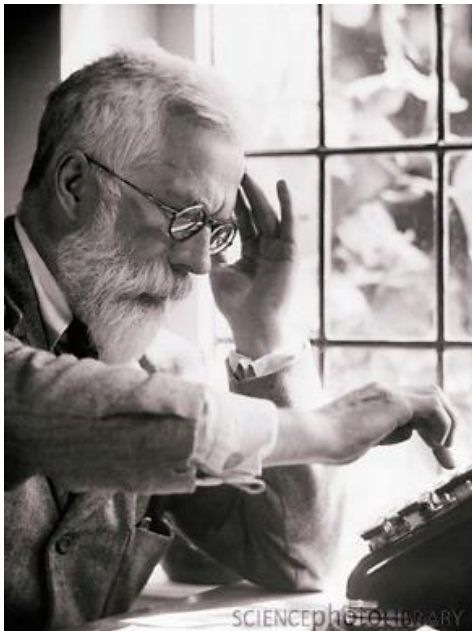


Abbildung 6.2. Mittelwertverteilungen in Abhängigkeit vom Stichprobenumfang n für schiefverteilte und gleichverteilte Rohwerte; Histogramme basieren auf 1000 simulierten Stichproben des Umfangs n

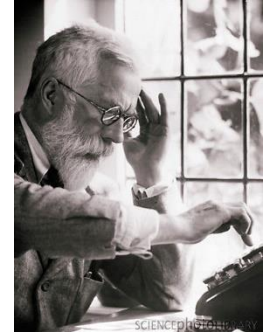
Nullhypothesen-Signifikanztest

- Der Nullhypothesen-Signifikanztest („klassischer“ Signifikanztest) ist ein Hybrid aus dem Modell von Sir Ronald A. Fisher (1925, 1926) und dem Modell von Jerzy Neyman und Egon Pearson (1928, 1933)



Nullhypothesen-Signifikanztest

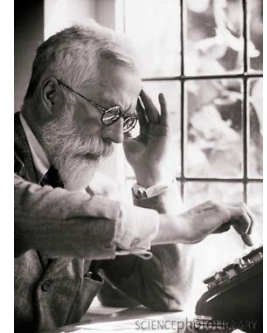
- Modell von Fisher
- „Lady tasting tea“- Experiment



- Ausgangspunkt für dieses Experiment war die Behauptung von Frau Bristol, angeben zu können, ob zuerst die Milch oder der Tee in die Tasse eingeschenkt wurde
- Also dachte sich Fisher folgendes: er servierte Frau Bristol einige Tassen, in denen zuerst die Milch, dann der Tee eingeschenkt wurde, und einige Tassen, in denen zuerst der Tee, und dann die Milch eingeschenkt wurde, und die Aufgabe von Frau Bristol bestand darin, bei jeder Tasse die Reihenfolge des Einschenkens richtig zu bestimmen
- Frau Bristol wusste, dass es eine gleich große Anzahl an Tassen gab, in denen die Milch vor bzw. nach dem Tee eingeschenkt wurde, jedoch nicht die Servier-Reihenfolge

Nullhypothesen-Signifikanztest

- Modell von Fisher
- „The Lady tasting tea“- Experiment

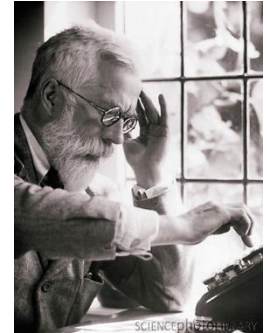


- Denkt man sich nun im einfachsten Fall zwei Tassen (1 x Milch vor Tee, 1 x Milch nach Tee), dann hat Frau Bristol eine Ratewahrscheinlichkeit von 50%; wenn sie also richtig liegt, dann ist nicht unbedingt davon auszugehen, dass ihre Behauptung stimmt, da die Ratewahrscheinlichkeit sehr hoch liegt
- Wie ist es bei 6 Tassen (3 x Milch vor Tee, 3 x Milch nach Tee)? Hier gibt es schon 20 mögliche Servier-Reihenfolgen. Hier beträgt die Ratewahrscheinlichkeit nur noch 5% ($1/20=0.05$). Schafft es also Frau Bristol hier, alle Proben korrekt zu bestimmen, dann kann man schon ziemlich sicher sein, dass ihre Behauptung auch stimmt

Nullhypothesen-Signifikanztest

- Modell von Fisher
- „The Lady tasting tea“- Experiment

- Die grundlegende Idee hinter dem Ansatz von Fisher besteht also darin, dass man von der Existenz von Effekten nur dann zuversichtlich ausgehen kann, wenn der Effekt mit einer Ratewahrscheinlichkeit (also bei bloßem Zufall) von 5% oder darunter auftritt
- Man könnte auch anders formulieren: wenn unter der Annahme, dass kein Effekt existiert, das Ergebnis, das wir erhalten, mit einer Wahrscheinlichkeit von 5% oder einer noch geringeren Wahrscheinlichkeit auftritt, dann können wir zuversichtlich sein, dass der Effekt auch existiert (wir geben also die Annahme, dass kein Effekt existiert, auf) – die Wahrscheinlichkeit selbst, mit der das Ergebnis unter dieser Annahme auftritt, wird als p -Wert bezeichnet
- In den empirischen Sozialwissenschaften legt man in diesem Zusammenhang ein sogenanntes Signifikanzniveau (üblicherweise mit 5%; das muss aber nicht immer so sein) fest – wenn der p -Wert darunter liegt (oder gleich diesem Niveau ist), spricht man von einem signifikanten Ergebnis



Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson

- Hypothesenpaar

- Demnach sollten wissenschaftliche Behauptungen in prüfbare und einander ausschließende Hypothesen so aufgeteilt werden, dass eine Hypothese besagt, dass ein Effekt existiert, und die andere Hypothese besagt, dass kein Effekt existiert
- Folgendes Hypothesenpaar sollte also aufgestellt werden:
 - H_1 = die Alternativhypothese besagt, dass ein Effekt existiert
 - H_0 = die Nullhypothese besagt, dass kein Effekt existiert



Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson

- Schokolade-Beispiel bei Field:

- Ausgangspunkt ist folgende Beobachtung: wenn man sich bloß vorstellt, Schokolade zu essen, dann isst man nachher weniger Schokolade
- Folgendes Hypothesenpaar wird dazu aufgestellt:
 - H_1 = Wenn man sich vorstellt, Schokolade zu essen, dann isst man nachher weniger davon
 - H_0 = Wenn man sich vorstellt, Schokolade zu essen, dann isst man nachher genauso viel Schokolade wie auch sonst immer



Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson

- Schokolade-Beispiel bei Field:

- Angenommen, man führt eine Studie in zwei Tagen mit 100 Testpersonen durch, und misst an Tag 1 der Erhebung, wieviel Schokolade sie essen; am zweiten Tag sollen sich die 100 Testpersonen vorstellen, Schokolade zu essen, und man misst wieder, wieviel Schokolade sie nachher essen



Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson



- Schokolade-Beispiel bei Field:

- Angenommen, das Ergebnis ist: 75% der Testpersonen essen am Tag 2 weniger Schokolade als am Tag 1
- Die Frage, die man sich angesichts des Ergebnisses stellen muss, lautet: „Angenommen, die Vorstellung, Schokolade zu essen, hat keinen Effekt – ist es dann wahrscheinlich, dass 75% der Testpersonen am Tag 2 weniger Schokolade essen?“
- Hier kann man sagen, dass die Wahrscheinlichkeit sehr gering ist – wenn die Nullhypothese wahr ist, dann sollten die Testpersonen an beiden Tagen dieselbe Menge Schokolade essen, und wenn 75% der Testpersonen am Tag 2 weniger Schokolade essen, dann ist das unter der Annahme der Nullhypothese ziemlich unwahrscheinlich

Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson



- Schokolade-Beispiel bei Field:

- Angenommen, das Ergebnis ist: 1% der Testpersonen essen am Tag 2 weniger Schokolade als am Tag 1
- Die Frage, die man sich angesichts des Ergebnisses stellen muss, lautet: „Angenommen, die Vorstellung, Schokolade zu essen, hat keinen Effekt – ist es dann wahrscheinlich, dass 1% der Testpersonen am Tag 2 weniger Schokolade essen?“
- Hier kann man sagen, dass die Wahrscheinlichkeit sehr hoch ist – wenn die Nullhypothese wahr ist, dann sollten die Testpersonen an beiden Tagen dieselbe Menge Schokolade essen, und wenn nur 1% der Testpersonen an Tag 2 weniger Schokolade essen, dann ist das unter der Annahme der Nullhypothese ziemlich wahrscheinlich

Nullhypothesen-Signifikanztest

- Modell von Neyman und Pearson



- Schokolade-Beispiel bei Field:

- Das Beispiel zeigt auch, dass es keinen Sinn macht, davon zu sprechen, dass die Nullhypothese richtig oder falsch bzw. die Alternativhypothese richtig oder falsch ist
- Man kann nur Wahrscheinlichkeiten angeben, dass ein Ergebnis unter der Annahme der Nullhypothese auftritt oder nicht auftritt (man testet also grob gesprochen die Nullhypothese!)
- Außerdem ist zu beachten, dass die Alternativhypothese nur eine von vielen möglichen Alternativhypothesen ist, sodass trotz eines für die eigens aufgestellte Alternativhypothese günstigen Ergebnisses vielleicht andere Alternativhypothesen noch zutreffender sein könnten

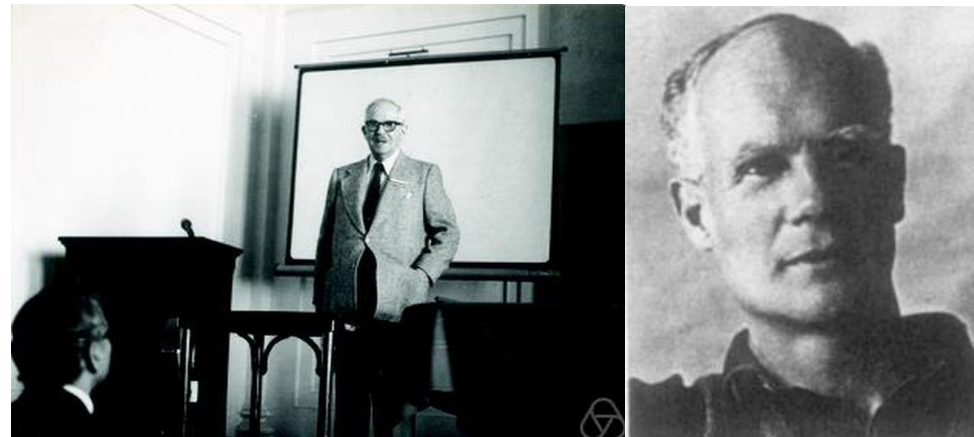
Nullhypothesen-Signifikanztest

- Der Nullhypothesen-Signifikanztest ist die Synthese aus beiden Modellen

***p*-Wert**



Hypothesenpaar



Zusammenfassung

- Signifikanz p -Wert
 - Entspricht der Wahrscheinlichkeit eines Ereignisses unter der Annahme der H_0
 - Entspricht der Wahrscheinlichkeit sich zu irren, bei der Entscheidung für die H_1
 - Wenn Irrtumswahrscheinlichkeit p sehr gering ($< \alpha$), dann entscheiden wir uns zu Gunsten der H_1

Teststatistik

- Mit Teststatistiken wird im Sinne des Nullhypothesen-Signifikanztests (grob gesprochen) erhoben, wie zuversichtlich man sein kann, dass die Alternativhypothese, die ja einen Effekt postuliert, gegenüber der Nullhypothese, in der das Vorliegen des Effekts verneint wird, zutrifft
- Obwohl es viele verschiedene Teststatistiken gibt (etwa t , χ^2 , F , etc.), repräsentieren sie alle folgendes Prinzip:

$$\text{Teststatistik} = \frac{\text{Signal}}{\text{Rauschen}} = \frac{\text{Systematische Varianz (Varianz, die durch das Modell erklärt wird)}}{\text{Unsystematische Varianz (Varianz, die nicht durch das Modell erklärt wird)}} = \frac{\text{Effekt}}{\text{Fehler}}$$

Teststatistik

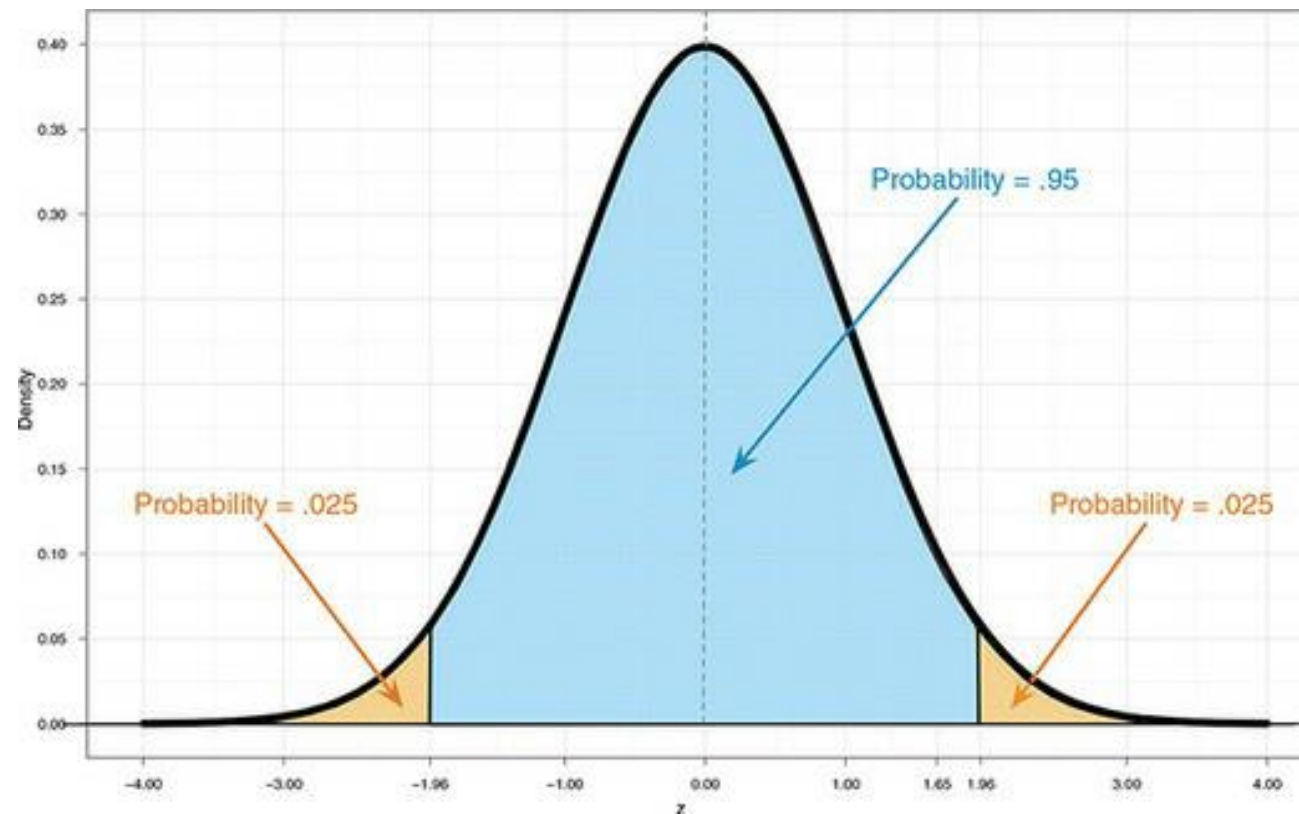
- Von diesen einzelnen Teststatistiken kennt man die Wahrscheinlichkeitsverteilung, das heißt, man weiß, mit welcher Wahrscheinlichkeit die Werte der jeweiligen Teststatistiken auftreten – es handelt sich bei dieser ermittelten Wahrscheinlichkeit um den p -Wert
- Nochmals zur Erinnerung: der p -Wert ist die bedingte Wahrscheinlichkeit, dass unter der Annahme der Gültigkeit der Nullhypothese das empirische oder ein extremeres Stichprobenergebnis auftritt
- Wenn der ermittelte p -Wert unter den als Signifikanzniveau definierten Wert (üblicherweise 5% – das muss aber wie gesagt nicht immer so sein) fällt, dann ist das Ergebnis statistisch signifikant

Normalverteilung (NV)

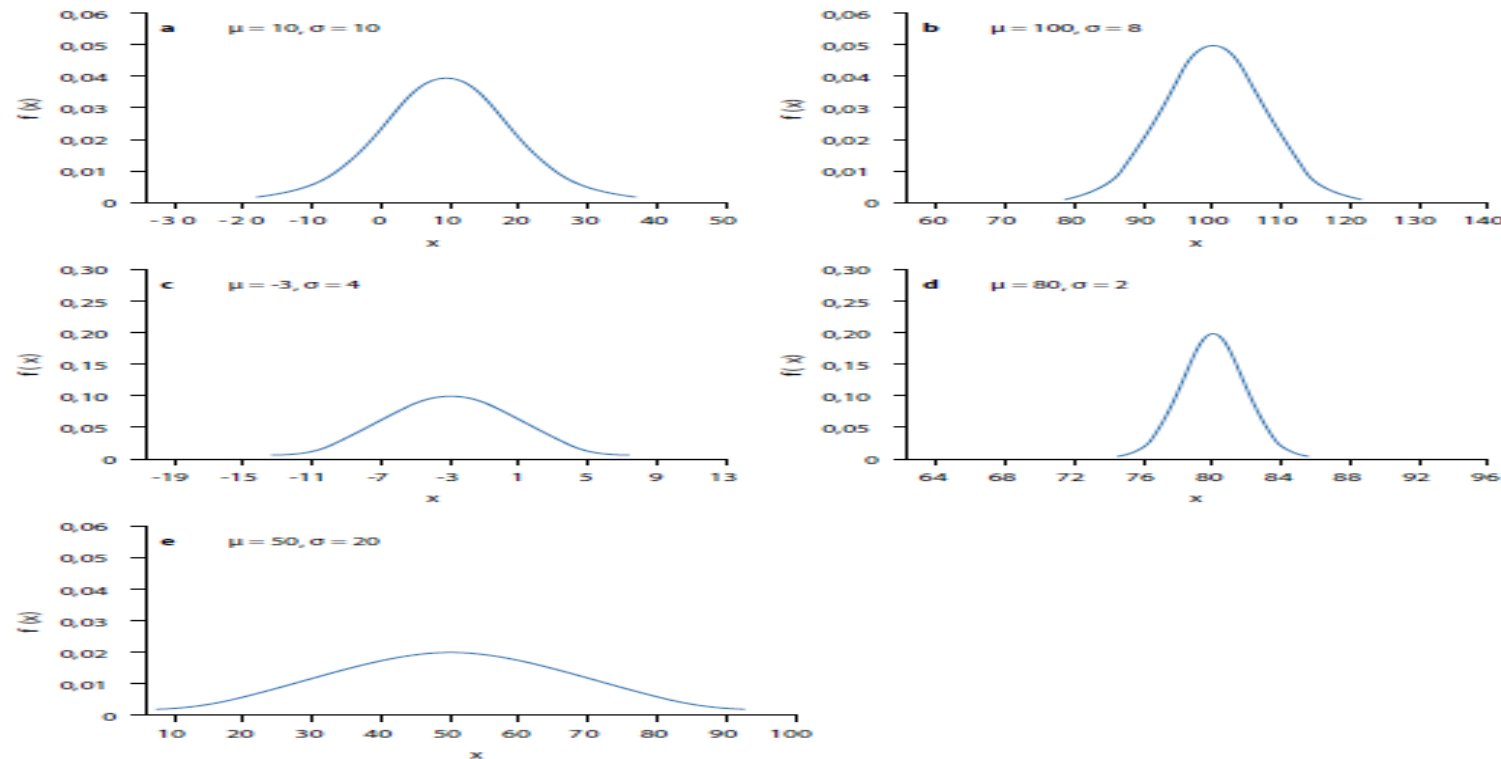
- Standardisierung (auch z-Transformation)
 - Mittelwert wird 0 gesetzt und Standardabweichung 1
 - Formel:

$$Z = \frac{X - \mu}{\sigma}$$

Teststatistik/ Wahrscheinlichkeitsverteilung Standardnormalverteilung ($\mu = 0, \sigma = 1$)

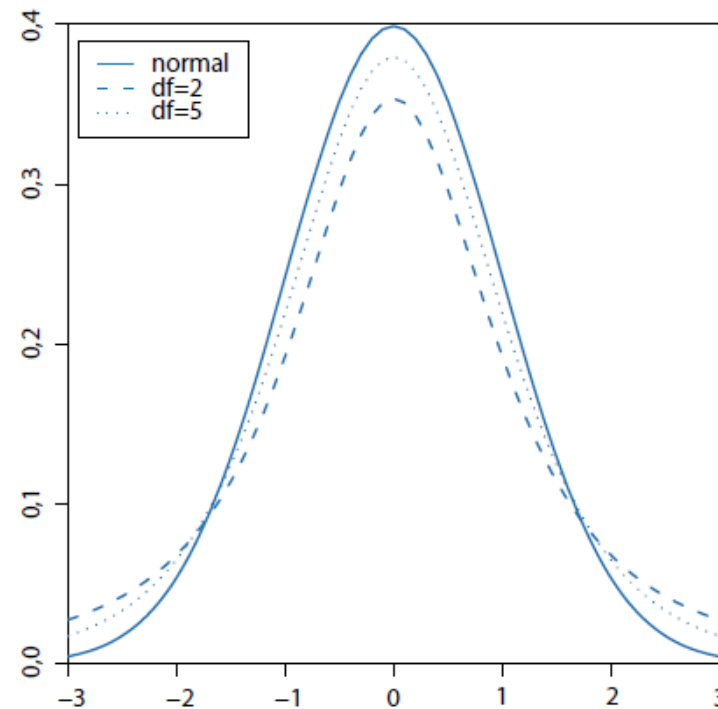


Teststatistik/ Wahrscheinlichkeitsverteilungen verschiedene Normalverteilungen



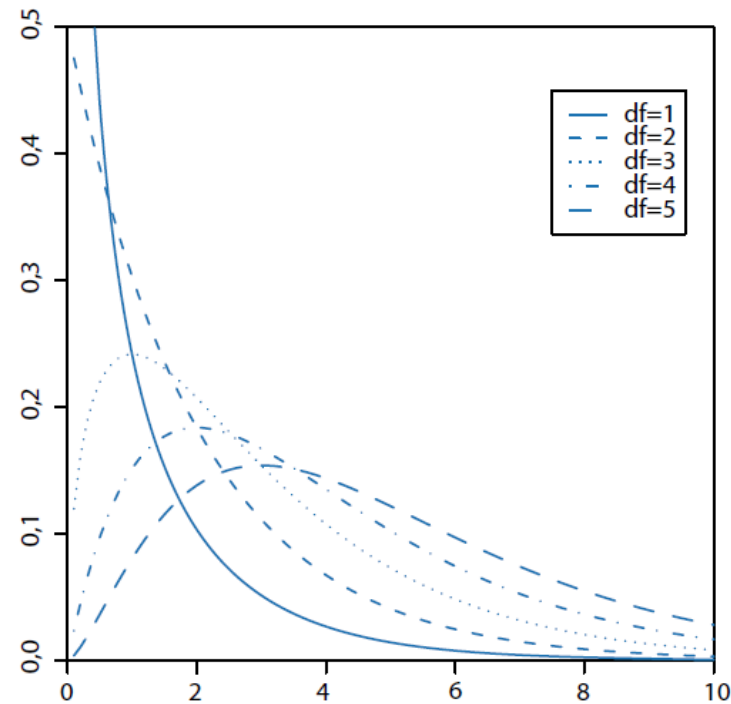
Quelle: Bortz, Jürgen und Christof Schuster. 2010. *Statistik für Human- und Sozialwissenschaftler*. 7., vollständig überarbeitete und erweiterte Auflage. Berlin, Heidelberg: Springer-Verlag, S. 71

Teststatistik/ Wahrscheinlichkeitsverteilungen t -Verteilungen



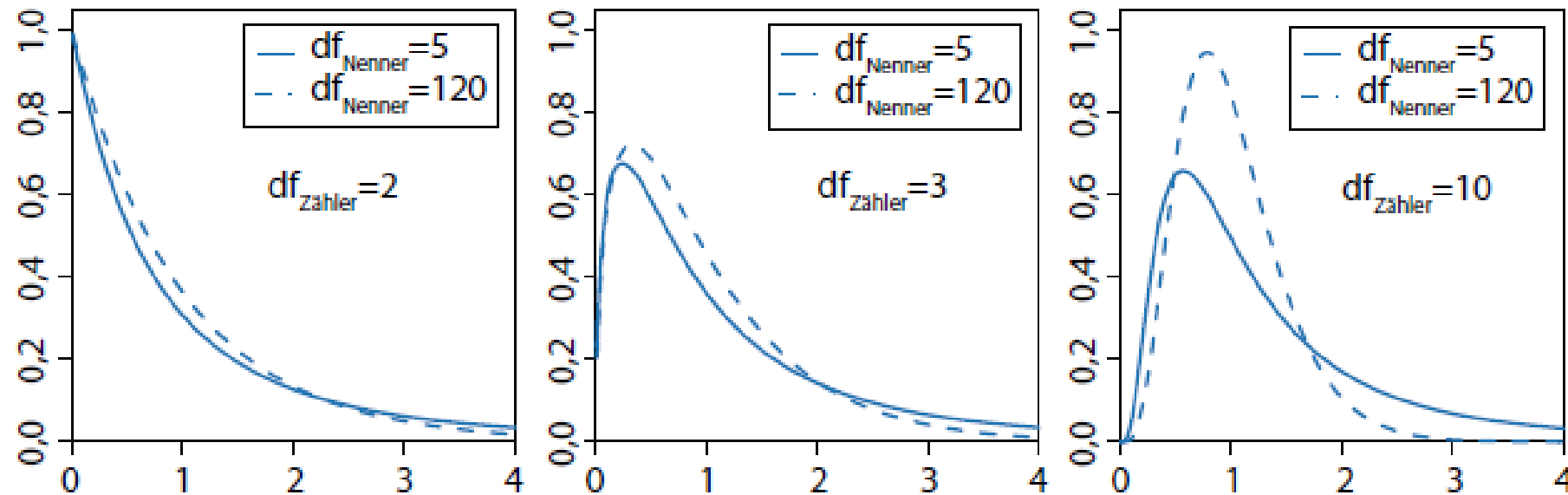
Quelle: Bortz, Jürgen und Christof Schuster. 2010. *Statistik für Human- und Sozialwissenschaftler*. 7., vollständig überarbeitete und erweiterte Auflage. Berlin, Heidelberg: Springer-Verlag, S. 75

Teststatistik/ Wahrscheinlichkeitsverteilungen χ^2 -Verteilungen



Quelle: Bortz, Jürgen und Christof Schuster. 2010. *Statistik für Human- und Sozialwissenschaftler*. 7., vollständig überarbeitete und erweiterte Auflage. Berlin, Heidelberg: Springer-Verlag, S. 75

Teststatistik/ Wahrscheinlichkeitsverteilungen F -Verteilungen



Quelle: Bortz, Jürgen und Christof Schuster. 2010. *Statistik für Human- und Sozialwissenschaftler*. 7., vollständig überarbeitete und erweiterte Auflage. Berlin, Heidelberg: Springer-Verlag, S. 76

Statistische Schlussfolgerung

- Ungerichtete/zweiseitige Hypothesen
 - AbsolventInnen der FHWN unterscheiden sich von AbsolventInnen der FH-OÖ bzgl. des Erfolgs
 - $H_0: \mu_1 = \mu_2$
 - $H_1: \mu_1 \neq \mu_2$
- Gerichtete/einseitige Hypothesen
 - AbsolventInnen der FHWN sind erfolgreicher als AbsolventInnen der FH-OÖ
 - $H_0: \mu_1 = \mu_2$
 - $H_1: \mu_1 > \mu_2$

Testung einseitig und zweiseitig

- Da eine zweiseitige Hypothese in beide Richtungen zu prüfen ist, ist der jeweilige Ablehnungsbereich der Nullhypothese kleiner

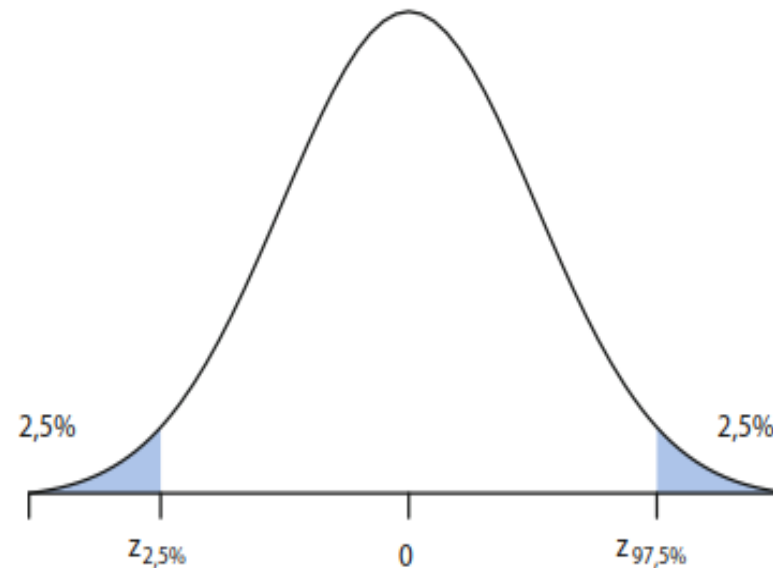


Abbildung 7.2. Ablehnungsbereich für den zweiseitigen z-Test für $\alpha = 0,05$; der Ablehnungsbereich wird durch die kritischen Werte $z_{2,5\%} = -1,96$ und $z_{97,5\%} = 1,96$ begrenzt

Testung einseitig und zweiseitig

- Da eine einseitige Hypothese in ein einzige Richtung zu prüfen ist, ist der jeweilige Ablehnungsbereich der Nullhypothese größer

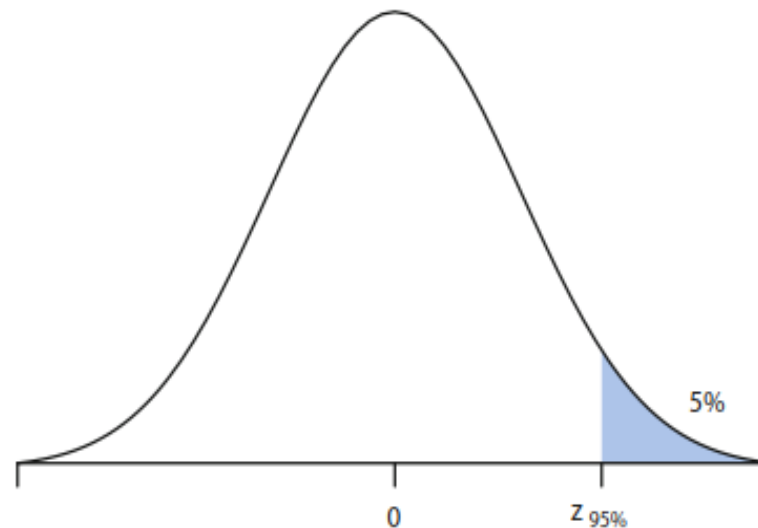


Abbildung 7.1. Ablehnungsbereich des einseitigen z-Tests zur Überprüfung der $H_1 : \mu > \mu_0$ für $\alpha = 0,05$; der Ablehnungsbereich wird durch den kritischen Wert $z_{95\%} = 1,65$ begrenzt

Testung einseitig und zweiseitig

TAKE HOME MESSAGES

1) BEI ZWEISEITIGEN TESTS ERHÄLT MAN SCHWERER EIN SIGNIFIKANTES ERGEBNIS



2) EINSEITIGE TESTS SOLLTEN NUR NACH SORGFÄLTIGEM ABWÄGEN VERWENDET WERDEN

Abbildung 7.2: Ablehnungsbereich für den zweiseitigen z-Test
für $\alpha = 0,05$. Die kritischen Werte $z_{2,5\%} = -1,96$ und $z_{97,5\%} = 1,96$ begrenzen
Abbildung 7.1: Ablehnungsbereich des einseitigen z-Tests zur
Prüfung der Hypothese $H_0: \mu = \mu_0$ gegen $H_1: \mu > \mu_0$. Der
Ablehnungsbereich wird durch den kritischen Wert $z_{95\%} = 1,65$ begrenzt

Testung einseitig und zweiseitig

- Was bedeutet „sorgfältiges Abwägen“?
- Field zitiert in diesem Zusammenhang John Tukey's Antwort auf folgende Frage: „Do you mean to say that one should *never* do a one-tailed test?“
- Tukeys Antwort: „Not at all. It depends upon to whom you are speaking. *Some people will believe anything*“

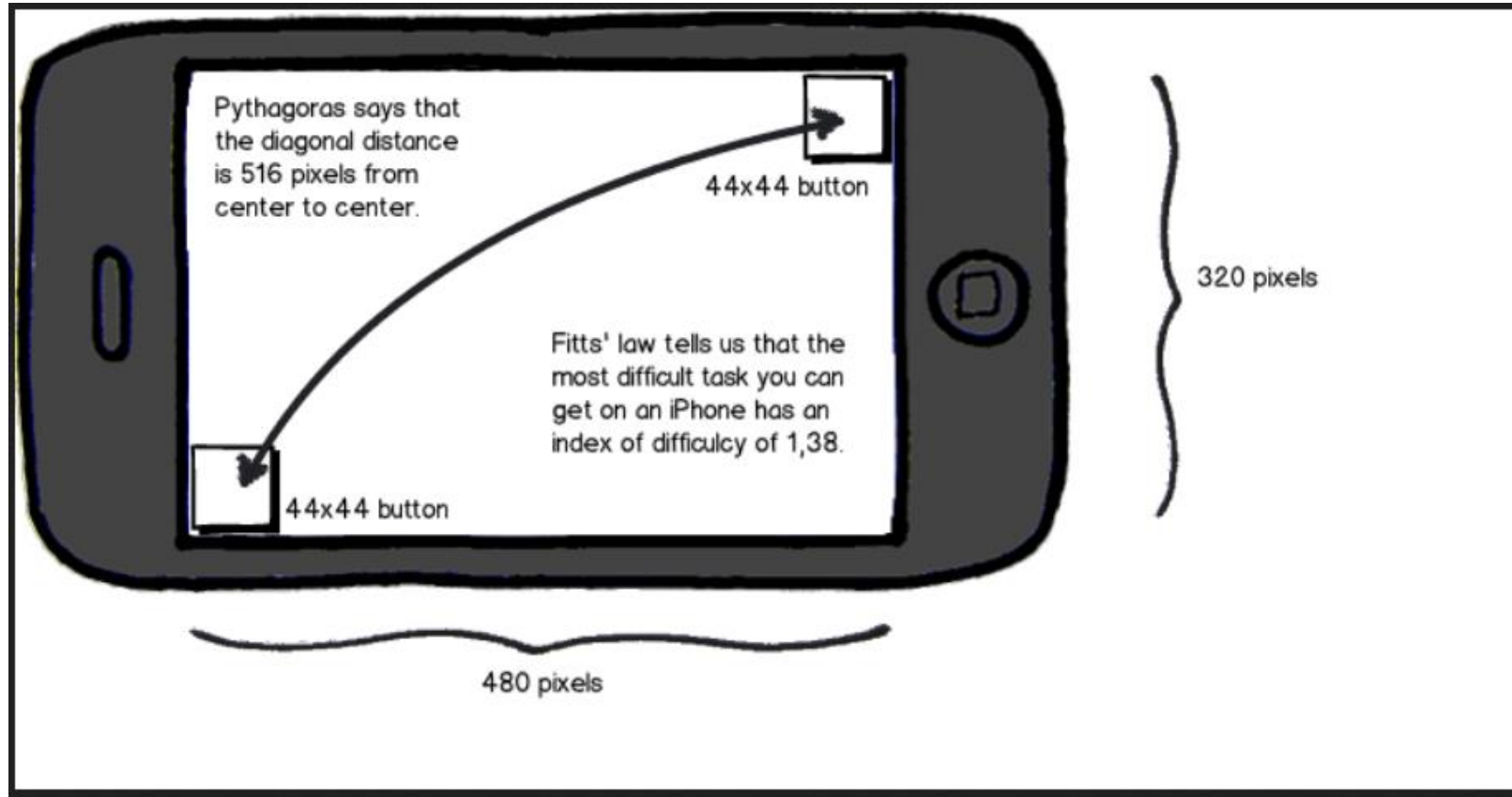
Testung einseitig und zweiseitig

- Was bedeutet „sorgfältiges Abwägen“?
- Einseitige Hypothesen können dann aufgestellt werden, wenn
 - aufgrund zahlreicher empirischer Evidenz aus bisherigen Studien ein Effekt in eine gewisse Richtung erwartet werden kann und daher die Ableitung einer einseitigen Hypothese dadurch gestützt wird
 - wenn das Ergebnis eines einseitigen Hypothesentests zu denselben Entscheidungen führt wie ein nicht signifikantes Ergebnis

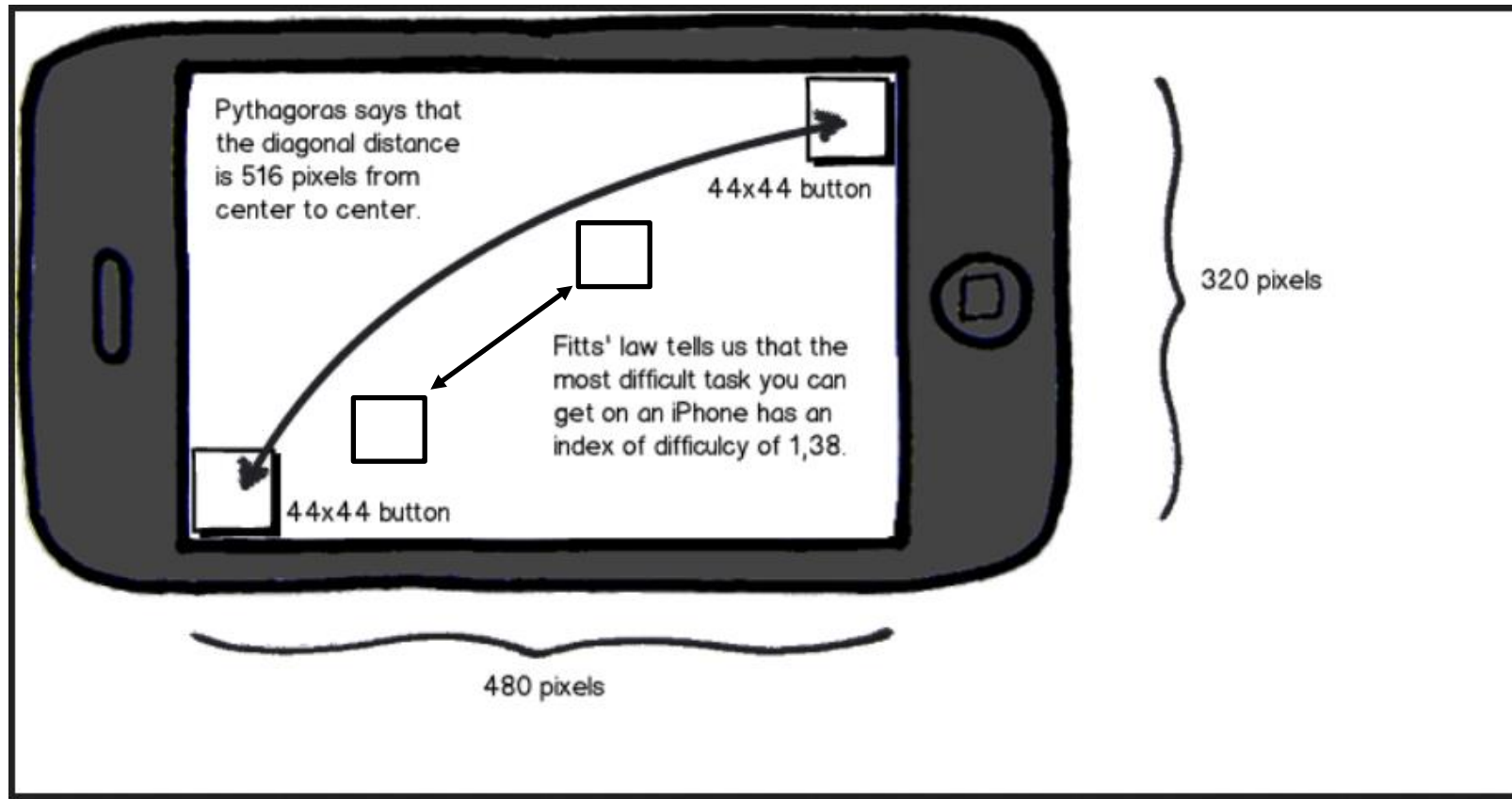
Testung einseitig und zweiseitig

- Einseitige Hypothesen können dann aufgestellt werden, wenn aufgrund zahlreicher empirischer Evidenz aus bisherigen Studien ein Effekt in eine gewisse Richtung erwartet werden kann und daher die Ableitung einer einseitigen Hypothese dadurch gestützt wird
 - Anwendungsbeispiel Usability-Testung – Das Fitts'sche Gesetz
 - ...besagt, dass in Abhängigkeit von Zielgröße g und dem Abstand d zwischen zwei Zielgrößen die Dauer der Bewegung zwischen den Zielgrößen von der Gesetzmäßigkeit $t = a + b \times \log_2(2d/g)$ abhängig ist (Fitts, 1954)
 - Einseitige Hypothese: Je geringer der Abstand zwischen zwei Zielgrößen ist, desto kürzer ist die Dauer der Bewegung

Testung einseitig und zweiseitig



Testung einseitig und zweiseitig



Testung einseitig und zweiseitig

- Einseitige Hypothesen können dann aufgestellt werden, wenn das Ergebnis eines einseitigen Hypothesentests zu derselben Entscheidung führt wie ein nicht signifikantes Ergebnis
- Wann ist das der Fall?
 - Beispiel: eine neue Absatzstrategie soll gegenüber einer bewährten Absatzstrategie getestet werden – stellt sich heraus, dass die neue Absatzstrategie nicht besser ist als die bewährte Absatzstrategie (nicht signifikantes Ergebnis), wird man die neue Absatzstrategie nicht anwenden; ist die neue Absatzstrategie signifikant schlechter als die bewährte Absatzstrategie (signifikanter Effekt, aber in die entgegengesetzte Richtung), dann würde man die Absatzstrategie ebenfalls nicht anwenden

Alpha-Fehler und Beta-Fehler

- Alpha-Fehler = Fehler 1. Art → wenn eine richtige Nullhypothese zugunsten der Alternativhypothese abgelehnt wird (wird für gewöhnlich mit 5% festgelegt)
- Beta-Fehler = Fehler 2. Art → wenn eine falsche Nullhypothese beibehalten wird (wird für gewöhnlich mit 20% festgelegt)

Tabelle 7.1. Fehlerarten bei statistischen Entscheidungen

		Entscheidung	
		für H_0	gegen H_0
Population	H_0 gilt		Fehler 1. Art
	H_0 gilt nicht	Fehler 2. Art	

Alpha-Fehler und Beta-Fehler

- Alpha-Fehler = Fehler 1. Art → wenn eine richtige Nullhypothese zugunsten der Alternativhypothese abgelehnt wird (wird für gewöhnlich mit 5% festgelegt)
 - Angenommen, es gibt in Wahrheit keinen Effekt in der Population – wiederholt man eine Studie 100x, kann man bei einem Signifikanzniveau von 5% 5x erwarten, eine Teststatistik zu erhalten, die so groß ist, dass man einen Effekt annimmt, obwohl in Wahrheit kein Effekt existiert

Alpha-Fehler und Beta-Fehler

- Beta-Fehler = Fehler 2. Art → wenn eine falsche Nullhypothese beibehalten wird (wird für gewöhnlich mit 20% festgelegt)
 - Angenommen, es gibt in Wahrheit einen Effekt in der Population – wiederholt man eine Studie 100x, kann man bei einem Beta-Fehler von 20% 20x erwarten, eine Teststatistik zu erhalten, die nicht groß genug ist, um den in Wahrheit existierenden Effekt in der Population aufzudecken

Alpha-Fehler und Beta-Fehler

- Es existiert zwischen den zwei Fehlerarten ein Trade-off
- Wenn man den Alpha-Fehler kleiner macht und damit die Wahrscheinlichkeit, einen Fehler 1. Art zu begehen, minimiert...
- ...dann erhöht man damit die Wahrscheinlichkeit, einen Fehler 2. Art zu begehen, das heißt, einen Effekt abzulehnen, der in Wahrheit existiert
- Diese Beziehung ist allerdings nicht geradlinig bzw. wie Field es formuliert, „straightforward“, weil die beiden Fehlerarten auf zwei unterschiedlichen Annahmen basieren (Effekt existiert nicht bzw. Effekt existiert in Wahrheit)

Testmacht (statistical power)

- Ausgangspunkt ist der Beta-Fehler = Fehler 2. Art → wenn eine falsche Nullhypothese beibehalten wird (wird für gewöhnlich mit 20% festgelegt)
- Setzt man den Beta-Fehler sehr hoch an, dann ist die Wahrscheinlichkeit hoch, einen Effekt, der in Wahrheit existiert, zu übersehen
- Setzt man den Beta-Fehler eher gering an, dann ist die Wahrscheinlichkeit hoch, einen Effekt, der in Wahrheit existiert, auch aufzudecken
- → die Teststärke ist die Fähigkeit eines Tests, einen Effekt aufzudecken, von dem angenommen wird, dass er in Wahrheit existiert und ist definiert mit 1-Beta-Fehler ($=1-\beta$; üblich ist $1-0.2 = 0.8$)

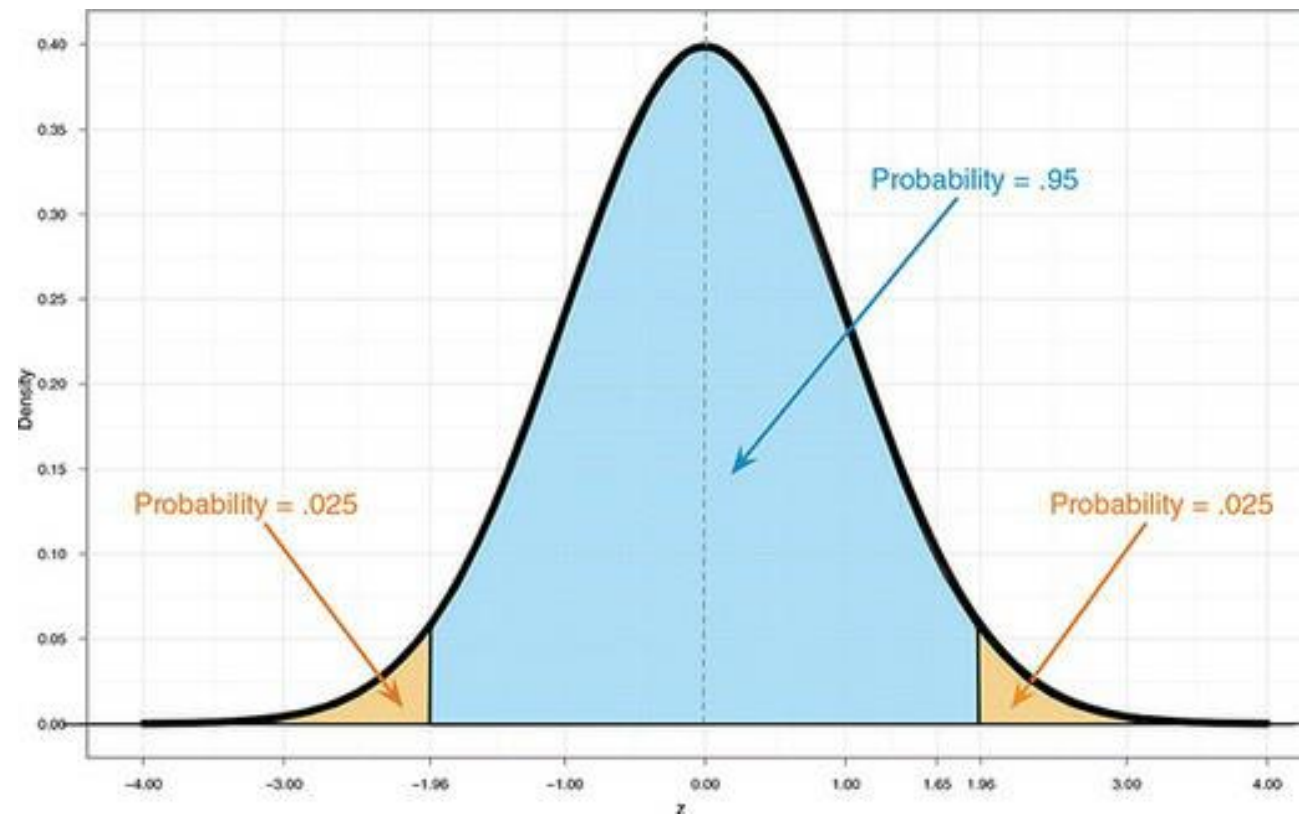
Statistische Schlussfolgerung

- Einflüsse auf die Testmacht
 - Größe des Effekts (Effektstärke)
 - Streuung der Werte
 - Wahl des α -Niveaus (und der Alpha-Fehler-Kumulierung)
 - Stichprobengröße
 - Richtige Verfahrenswahl
 - Einseitige vs. zweiseitige Hypothesentestung
 - Abhängige/unabhängige Stichproben

Statistische Schlussfolgerung

- Hypothesen testen
 - Stichprobe -> Grundgesamtheit
 - H_0
 - H_1
 - Annahme- & Ablehnungsbereich

Teststatistik/ Wahrscheinlichkeitsverteilung Standardnormalverteilung ($\mu = 0, \sigma = 1$)



Konfidenzintervalle

- *KI für μ bei bekannter Populationsvarianz σ^2 hier mit Wahrscheinlichkeit von 95%*
 - $\mu_{1,2} = \bar{x} \pm z_{97,5} \frac{s}{\sqrt{n}}$
 - $\mu_{1,2} = \bar{x} \pm 1,96 \frac{s}{\sqrt{n}}$
 - Intervall in dem der wahre Parameter mit einer bestimmten Wahrscheinlichkeit $(1-\alpha)$ liegt

Konfidenzintervalle

■ Beispiel

- In welchem Intervall liegt mit 95 prozentiger Wahrscheinlichkeit der wahre Mittelwert μ , wenn aus einer Stichprobe von 100 Studierenden ein Mittelwert von 115 (Standardabweichung=15) ermittelt wurde?
- $n = 100$ StudentInnen, $\bar{x} = 115$, $s = 15$

$$\mu_{1,2} = 115 \pm 1,96 \frac{15}{\sqrt{100}}; \mu_1 = 112,06 \quad \mu_2 = 117,94$$

$$KI = [112,06; 117,94]$$

Konfidenzintervalle

- Eigenschaften
 - Je größer die Sicherheit $(1-\alpha)$, desto breiter KI
 - Je größer n , desto schmaler KI
 - Je größer die Streuung desto breiter KI

THE RESEARCH PROCESS



**Beispiele für inferenzstatistische Tests,
quantitativer Forschungsprozess (+ Übung)**

Inferenzstatistische Verfahren

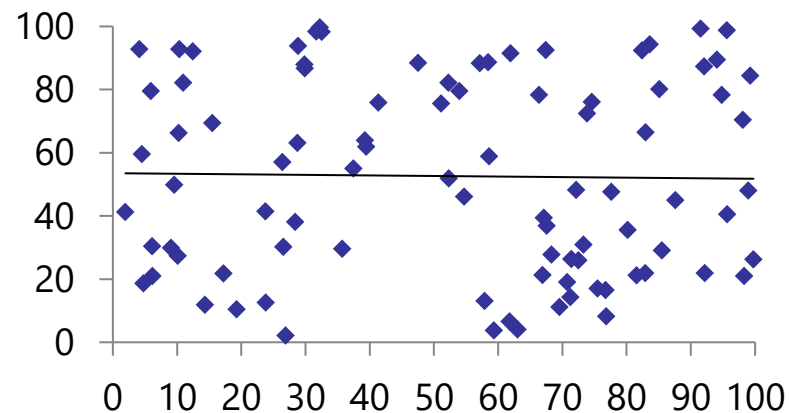
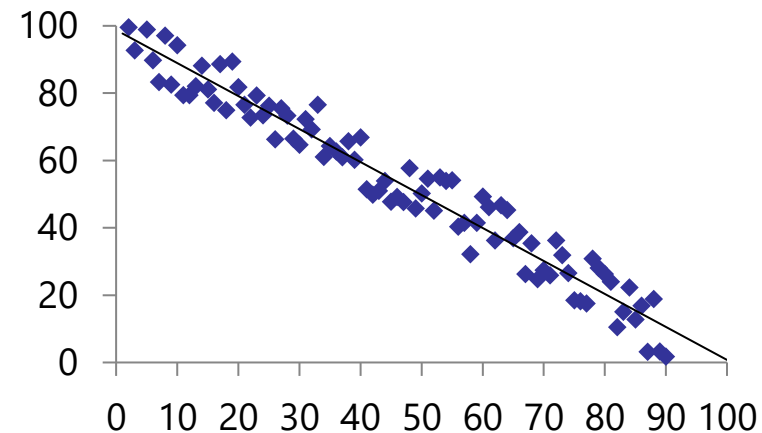
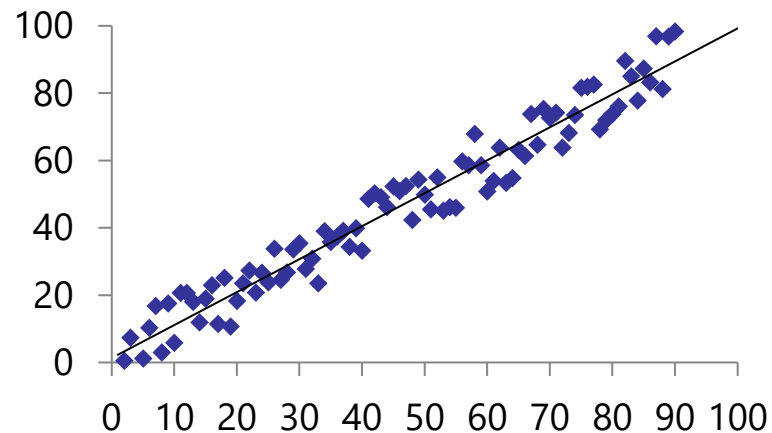
- Prüfung auf Zusammenhänge
 - Pearson-Korrelationskoeffizient/Regression
 - Spearman
 - Chi-Quadrat-Test

- Prüfung auf Unterschiede
 - Mittelwertsvergleich (T-Test)

Grafische Darstellung

■ Streudiagramm

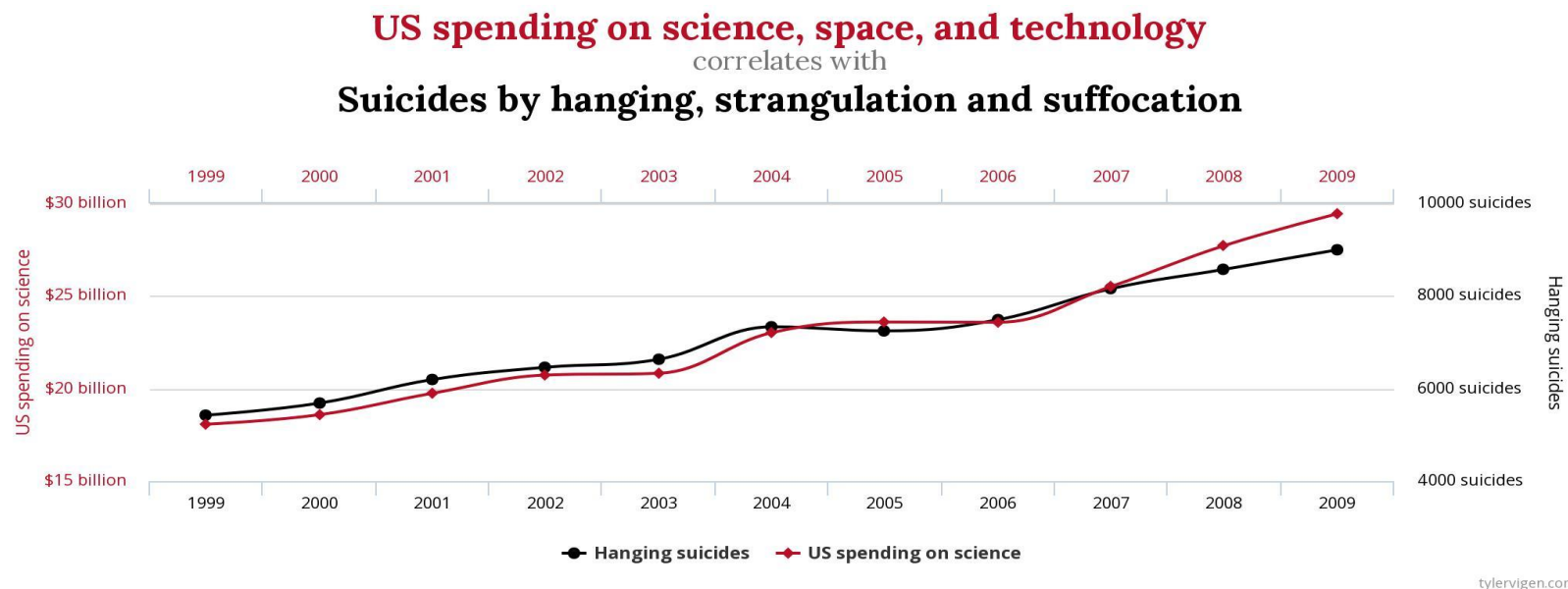
Darstellung des Zusammenhangs zweier Variablen



Grafische Darstellung

- Streudiagramm

Darstellung des Zusammenhangs zweier Variablen



Zusammenhangsmaße

- Kovarianz (cov)
- Pearson Korrelation (r)
- Bestimmtheitsmaß (B, R^2, r^2)
- Spearman-Korrelation (r_{sp})

Zusammenhangsmaße

- Kovarianz (*cov*)

$$\text{variance}(s^2) = \frac{\sum (x_i - \bar{x})^2}{N - 1} = \frac{\sum (x_i - \bar{x})(x_i - \bar{x})}{N - 1} \quad \Rightarrow \quad \text{cov}(x, y) = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{N - 1}$$

- Pearson Korrelation (*r*)
- Bestimmtheitsmaß (*B*, R^2 , r^2)
- Spearman-Korrelation (r_{sp})

Zusammenhangsmaße

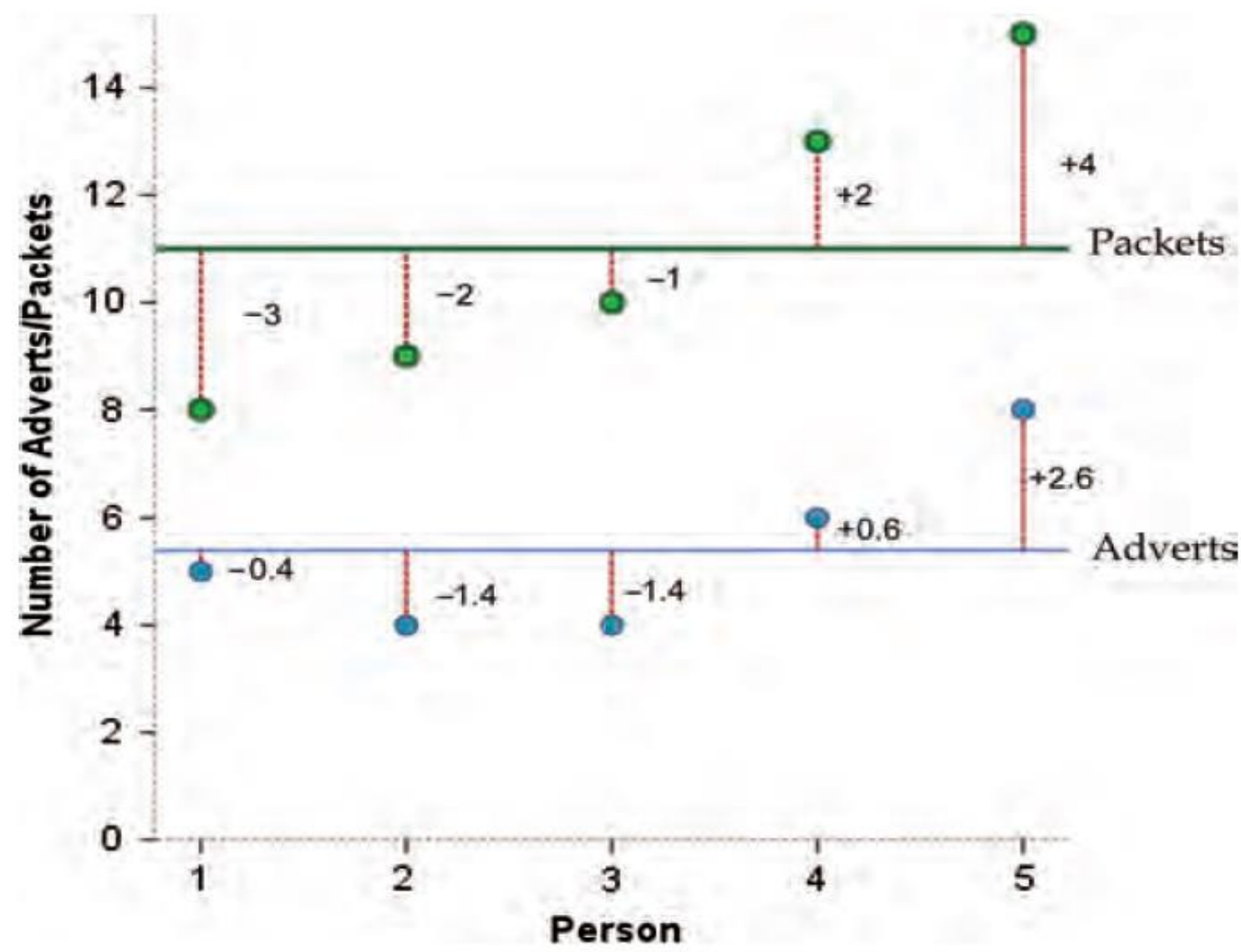
- Kovarianz (*cov*)

$$\text{variance}(s^2) = \frac{\sum (x_i - \bar{x})^2}{N - 1} = \frac{\sum (x_i - \bar{x})(x_i - \bar{x})}{N - 1} \quad \Rightarrow \quad \text{cov}(x, y) = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{N - 1}$$

Testperson	1	2	3	4	5	Mittelwert $\bar{x}$	s
Anzahl betr. Werbungen	5	4	4	6	8	5,4	1,67
Anzahl Käufe beworbenes Produkt	8	9	10	13	15	11,0	2,92

FIGURE 6.2

Graphical display of the differences between the observed data and the means of two variables



Zusammenhangsmaße

- Kovarianz (*cov*)

$$\text{variance}(s^2) = \frac{\sum (x_i - \bar{x})^2}{N - 1} = \frac{\sum (x_i - \bar{x})(x_i - \bar{x})}{N - 1} \quad \Rightarrow \quad \text{cov}(x, y) = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{N - 1}$$

$$\begin{aligned} \text{cov}(x, y) &= \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{N - 1} \\ &= \frac{(-0.4)(-3) + (-1.4)(-2) + (-1.4)(-1) + (0.6)(2) + (2.6)(4)}{4} \\ &= \frac{1.2 + 2.8 + 1.4 + 1.2 + 10.4}{4} \\ &= \frac{17}{4} \\ &= 4.25 \end{aligned}$$

Zusammenhangsmaße

- Kovarianz (cov)

$$\text{variance}(s^2) = \frac{\sum (x_i - \bar{x})^2}{N-1} = \frac{\sum (x_i - \bar{x})(x_i - \bar{x})}{N-1} \quad \Rightarrow \quad \text{cov}(x, y) = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{N-1}$$

- Pearson Korrelation (r)

$$r = \frac{\text{cov}_{xy}}{s_x s_y} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{(N-1)s_x s_y}$$

- Bestimmtheitsmaß (B, R^2, r^2)
- Spearman-Korrelation (r_{sp})

Zusammenhangsmaße

- Pearson Korrelation (r)

$$r = \frac{\text{COV}_{xy}}{s_x s_y} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{(N - 1) s_x s_y}$$

Testperson	1	2	3	4	5	Mittelwert $\bar{x}$	s
Anzahl betr. Werbungen	5	4	4	6	8	5,4	1,67
Anzahl Käufe beworbenes Produkt	8	9	10	13	15	11,0	2,92

- $r = 4,25 / (1,67 * 2,92) = 0,87$

Zusammenhangsmaße

- Pearson Korrelation (r)

$$r = \frac{\text{COV}_{xy}}{s_x s_y} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{(N - 1) s_x s_y}$$

- Pearson-Korrelationskoeffizient = standardisiertes Zusammenhangsmaß
- $r = +1 \rightarrow$ perfekter linearer positiver Zusammenhang zwischen 2 Variablen,
 $r = -1 \rightarrow$ perfekter linearer negativer Zusammenhang zwischen 2 Variablen,
 $r = 0 \rightarrow$ kein linearer Zusammenhang zwischen 2 Variablen

Zusammenhangsmaße

- Pearson Korrelation (r)

$$r = \frac{\text{COV}_{xy}}{s_x s_y} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{(N - 1) s_x s_y}$$

- Pearson-Korrelationskoeffizient = standardisiertes Zusammenhangsmaß = Effektstärkenmaß
- $r = \pm 0.10 \rightarrow$ kleiner Effekt (der Effekt erklärt 1% der Varianz)
 $r = \pm 0.30 \rightarrow$ mittlerer Effekt (der Effekt erklärt 9% der Varianz)
 $r = \pm 0.50 \rightarrow$ großer Effekt (der Effekt erklärt 25% der Varianz)

Zusammenhangsmaße

- Bestimmtheitsmaß (B, R^2, r^2)
 - Wert für gemeinsam geteilte Varianz (r quadriert) = Ausmaß der Varianz, welche eine Variable mit der anderen teilt
 - in der linearen Regression ist es der Anteil der durch die Regression erklärten Varianz in der abhängigen Variable
- Spearman-Korrelation (r_{sp}) → nichtparametrisches Ersatzverfahren zur Pearson-Korrelation (bei Ordinalskalierung und/oder bei Verletzung der Normalverteilungsannahme → siehe aber zentrales Grenzwerttheorem!)

Lineare Regression

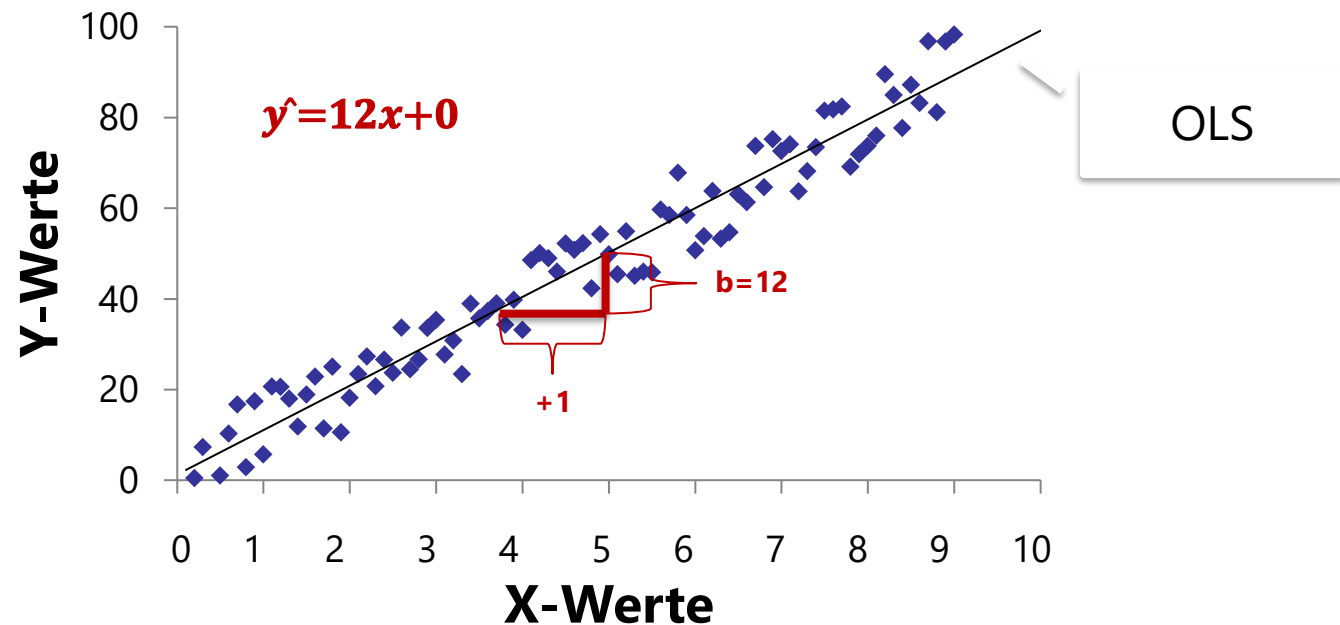
- Formel:

- $\hat{y} = bx + a$

- $B = r^2 = \frac{s_{\hat{y}}^2}{s_y^2}$

- Interpretationen

- Eine Einheit (+1) mehr in der X-Variable führt zu einem Anstieg/einer Veränderung b in der Y-Variable
 - z.B. eine Stunde mehr lernen resultiert in einer um 12 Punkte besseren Leistung...



Inferenzstatistik – Häufigkeitsvergleiche

- Chi-Quadrat Test (χ^2)
 - Testet Zusammenhänge zwischen kategorialen/nominalskalierten Variablen

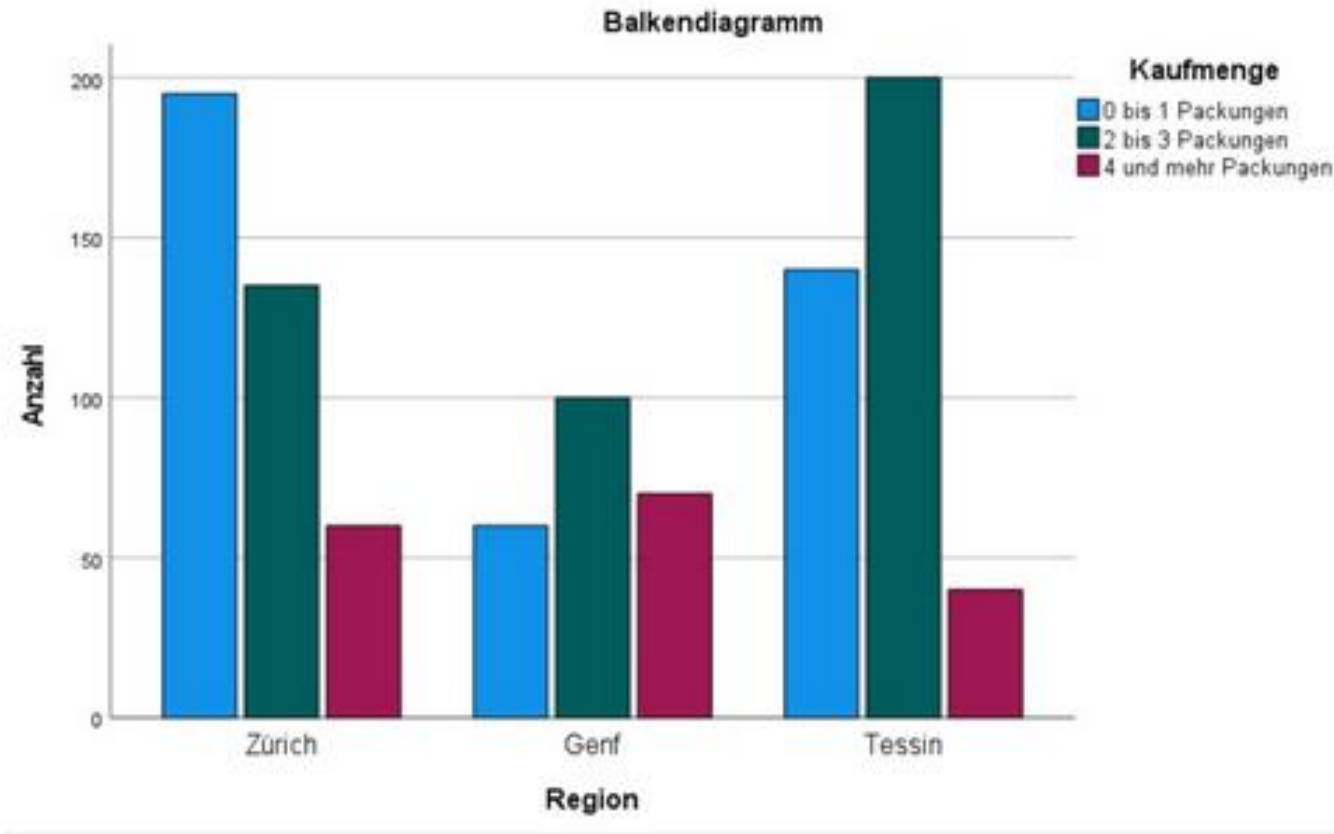
Inferenzstatistik – Häufigkeitsvergleiche

- Beispiel:

Beobachtete Anzahl		Kaufmenge			Gesamt	Prozent
		0-1 Packungen	2-3 Packungen	4 + Packungen		
Region	Zürich	195	135	60	390	39.0%
	Genf	60	100	70	230	23.0%
	Tessin	140	200	40	380	38.0%
Gesamt		395	435	170	1000	
Prozent		39.5%	43.5%	17.0%		100.0%

Inferenzstatistik – Häufigkeitsvergleiche

- Beispiel:



Inferenzstatistik – Häufigkeitsvergleiche

■ Beispiel:

Chi-Quadrat-Tests

	Wert	df	Asymptotisch e Signifikanz (zweiseitig)
Pearson-Chi-Quadrat	70,788 ^a	4	,000
Likelihood-Quotient	68,284	4	,000
Zusammenhang linear- mit-linear	2,706	1	,100
Anzahl der gültigen Fälle	1000		

a. 0 Zellen (0,0%) haben eine erwartete Häufigkeit kleiner 5.
Die minimale erwartete Häufigkeit ist 39,10.

Symmetrische Maße

		Wert	Näherungsw eise Signifikanz
Nominal- bzgl. Nominalmaß	Phi	,266	,000
	Cramer-V	,188	,000
	Kontingenzkoeffizient	,257	,000
Anzahl der gültigen Fälle		1000	

Inferenzstatistik – Mittelwertsvergleiche

- Ein-Stichproben t -Test
- t -Test für unabhängige Stichproben
- t -Test für abhängige Stichproben

Inferenzstatistik – Mittelwertsvergleiche

- Voraussetzungen des t -Tests
 - Gleiche Grundgesamtheit
 - Metrisches Skalenniveau
 - Normalverteilung in beiden Gruppen
 - Gleichheit der Varianzen (Varianzhomogenität)

Inferenzstatistik – Mittelwertsvergleiche

- Beispiel:

Gruppe	N	Resultate Gedächtnistest	Mittelwert	Standardabweichung
Schulklasse A	22	79 73 78 80 49 71 77 ... 74 70 60	74.32	9.67
Schulklasse B	25	88 91 84 70 75 91 73 ... 66 89 75	81.56	10.20

Inferenzstatistik – Mittelwertsvergleiche

- Beispiel:

Gruppenstatistiken

Schulklassen		N	Mittelwert	Standardabweichung	Standardfehler des Mittelwertes
Gedächtnistest	1	22	74.32	9.668	2.061
	2	25	81.56	10.198	2.040

Inferenzstatistik – Mittelwertsvergleiche

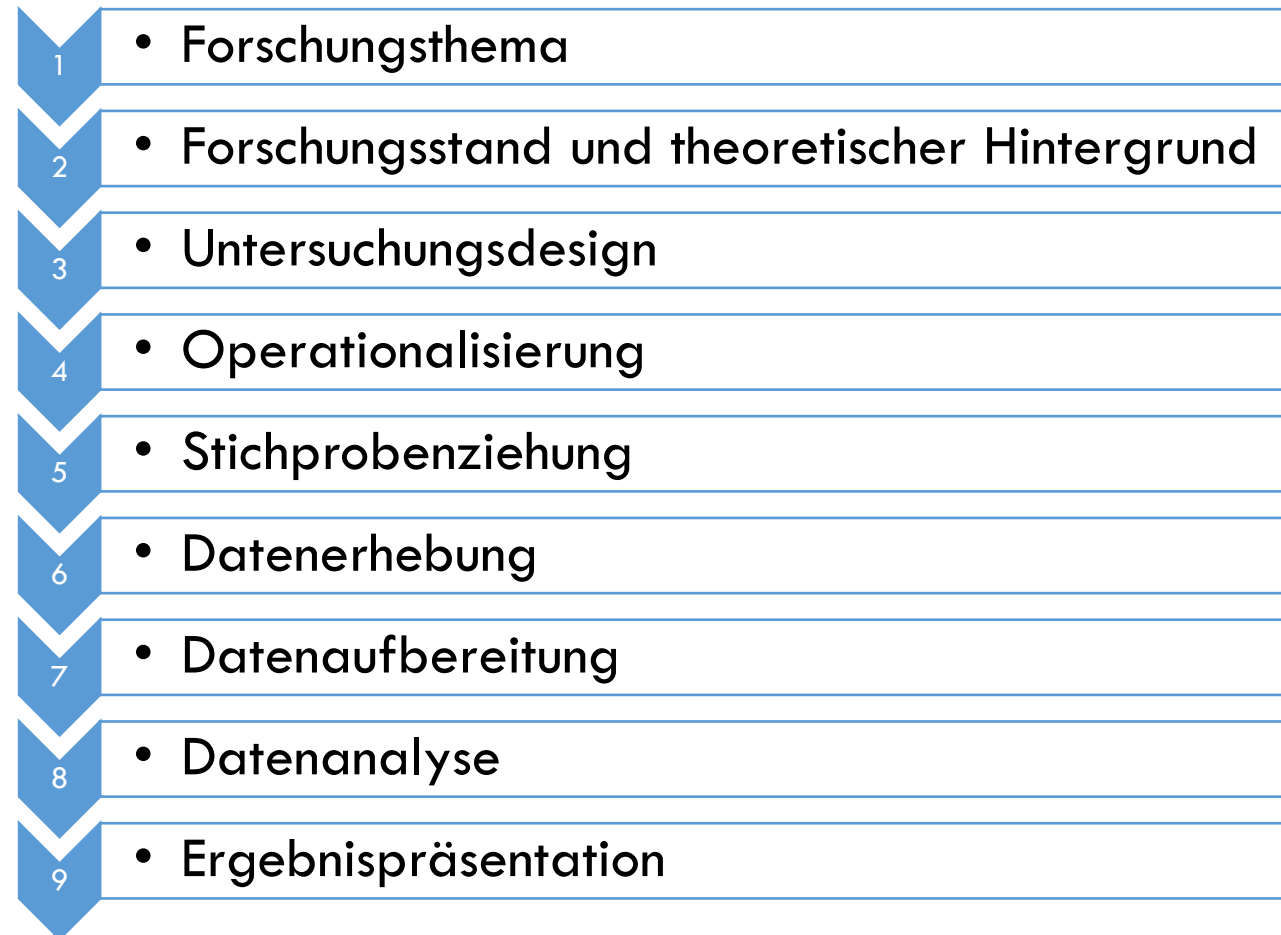
■ Beispiel:

Test bei unabhängigen Stichproben

		Levene-Test der Varianzgleichheit		T-Test für die Mittelwertgleichheit						
		F	Signifikanz	T	df	Sig. (2-seitig)	Mittlere Differenz	Standardfehler der Differenz	95% Konfidenzintervall der Differenz	
									Untere	Obere
Gedächtnistest	Varianzen sind gleich	1.157	.288	-2.489	45	.017	-7.242	2.910	-13.103	-1.381
	Varianzen sind nicht gleich			-2.497	44.733	.016	-7.242	2.900	-13.083	-1.400

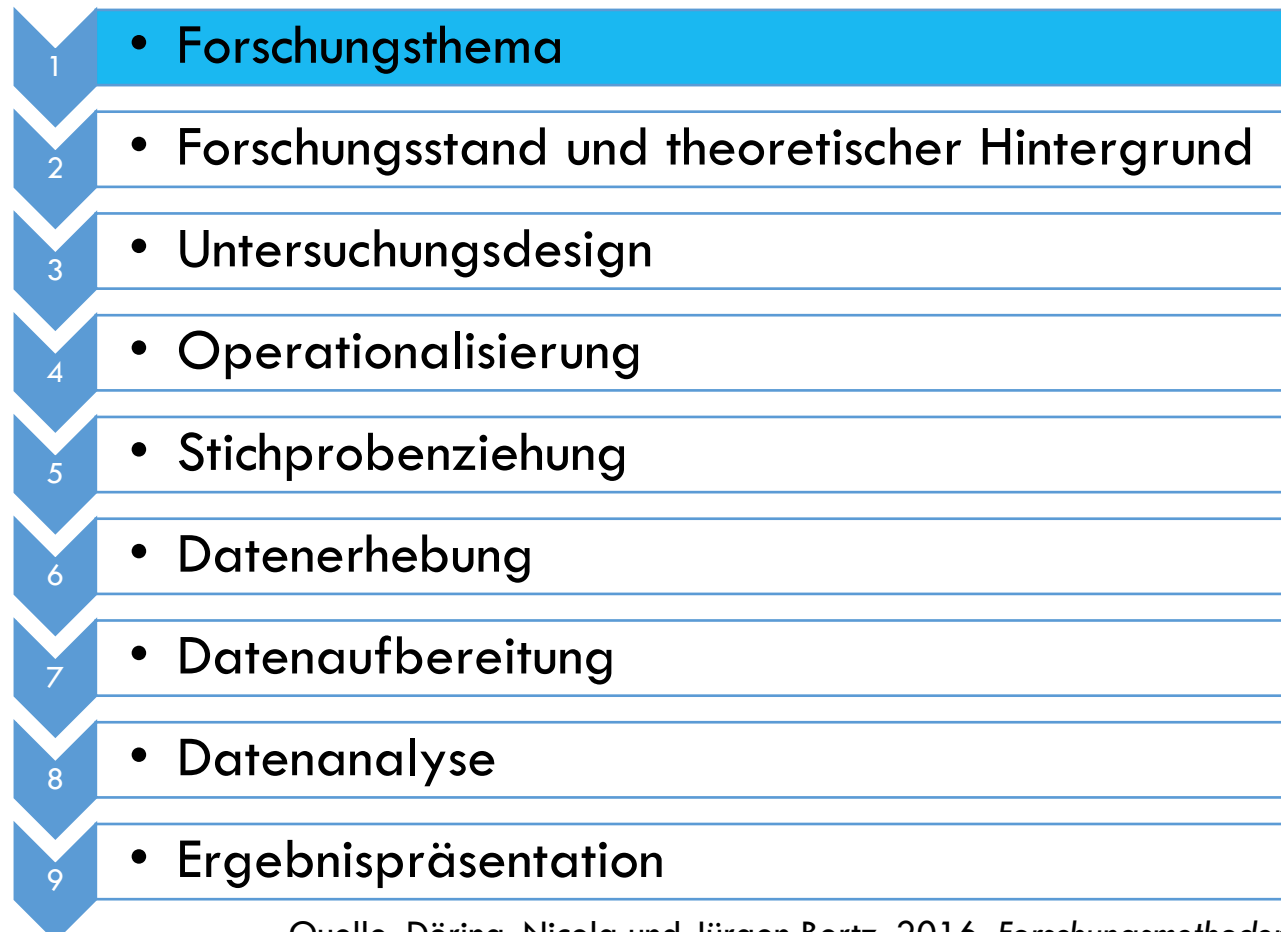
Was bedeutet es, quantitativ zu forschen?

Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



Was bedeutet es, quantitativ zu forschen?

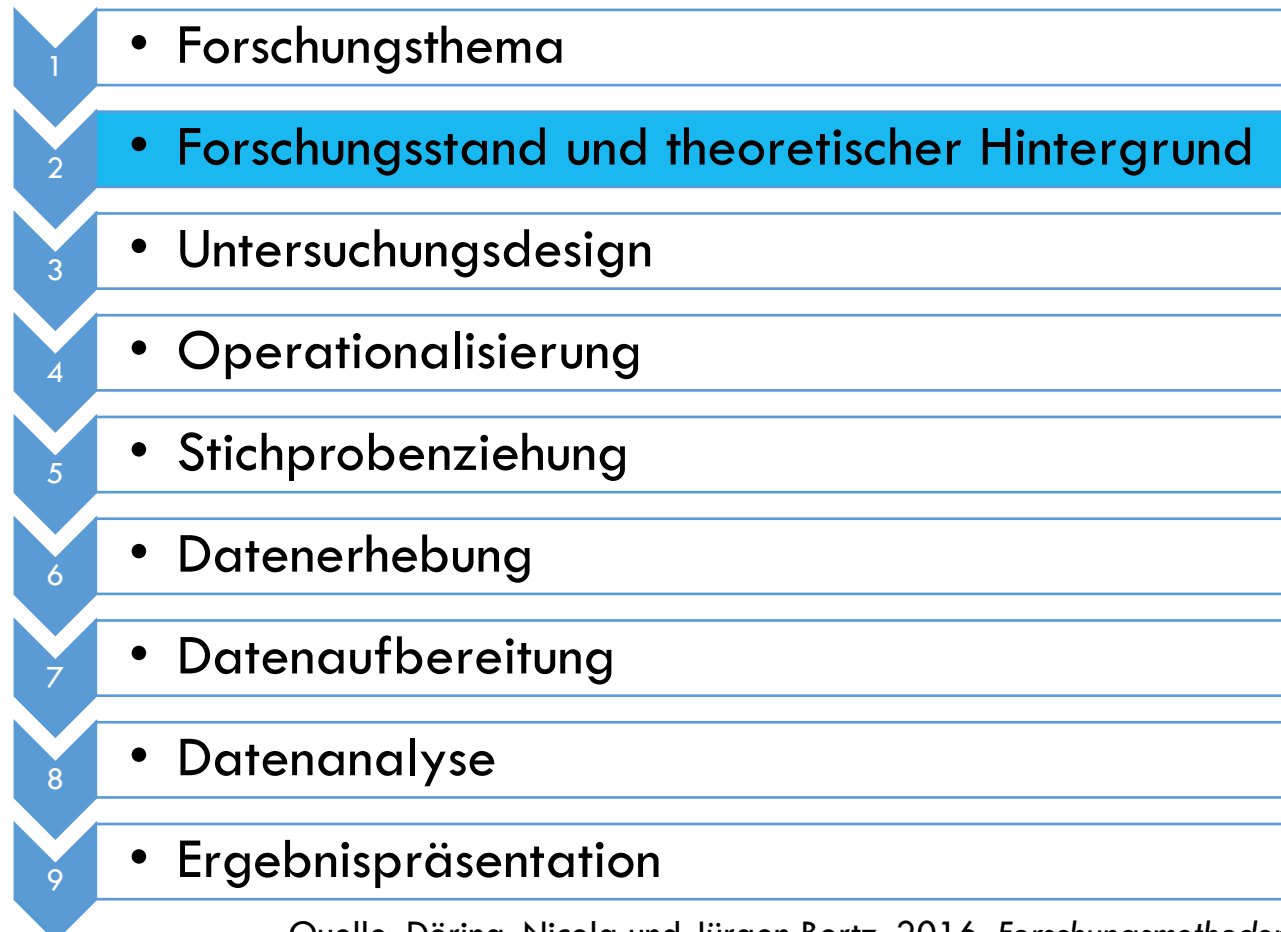
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Konkretisierung des Forschungsproblems
- Formulierung von Hypothesen
- Forschungsstand/ Theorien

Was bedeutet es, quantitativ zu forschen?

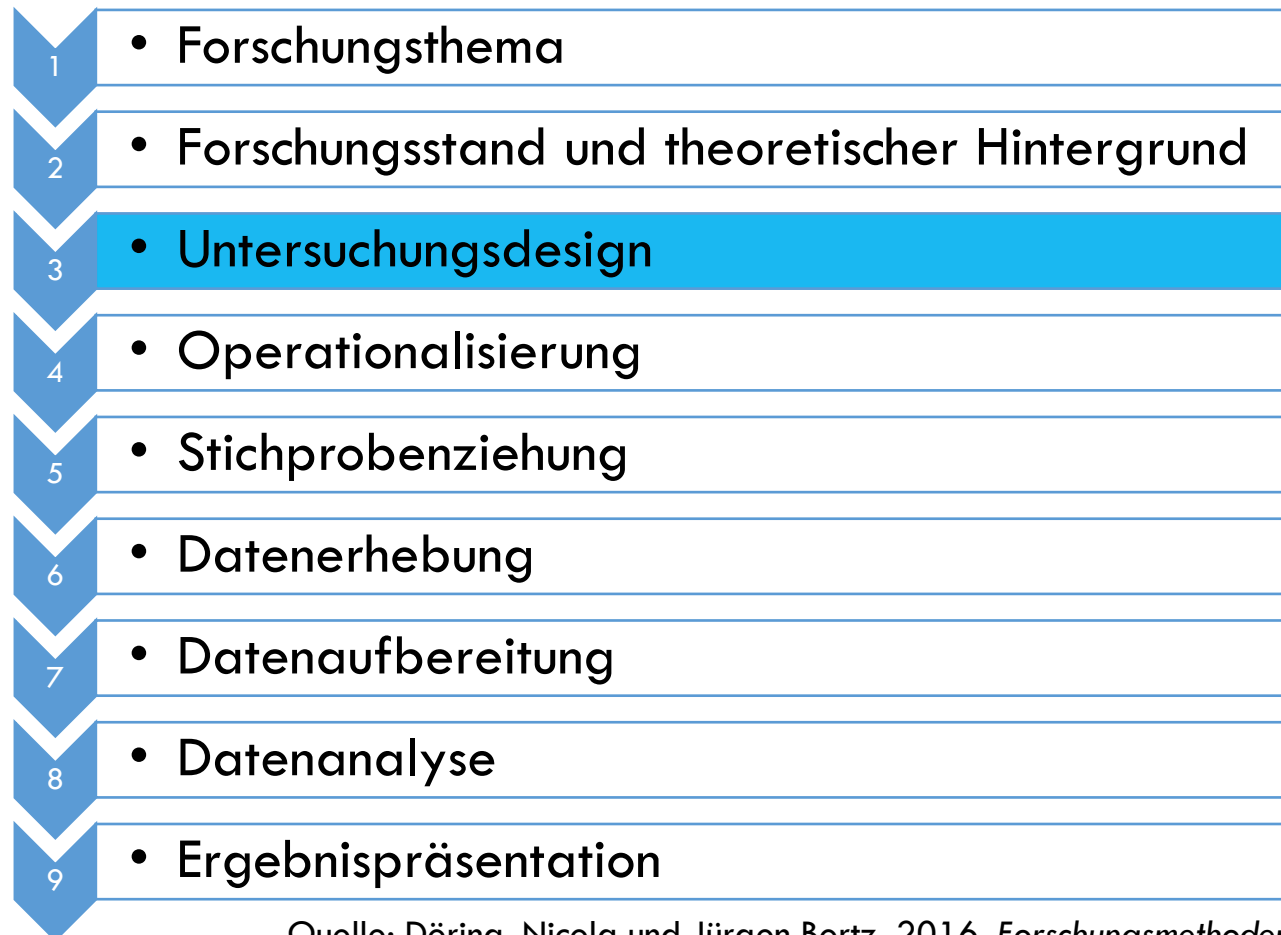
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Anknüpfung an bisherigen Forschungsstand
- Voraussetzung: gründliche Literaturrecherche
(!siehe die ZUB-Einheit auf Moodle!)
- Theoriearbeit

Was bedeutet es, quantitativ zu forschen?

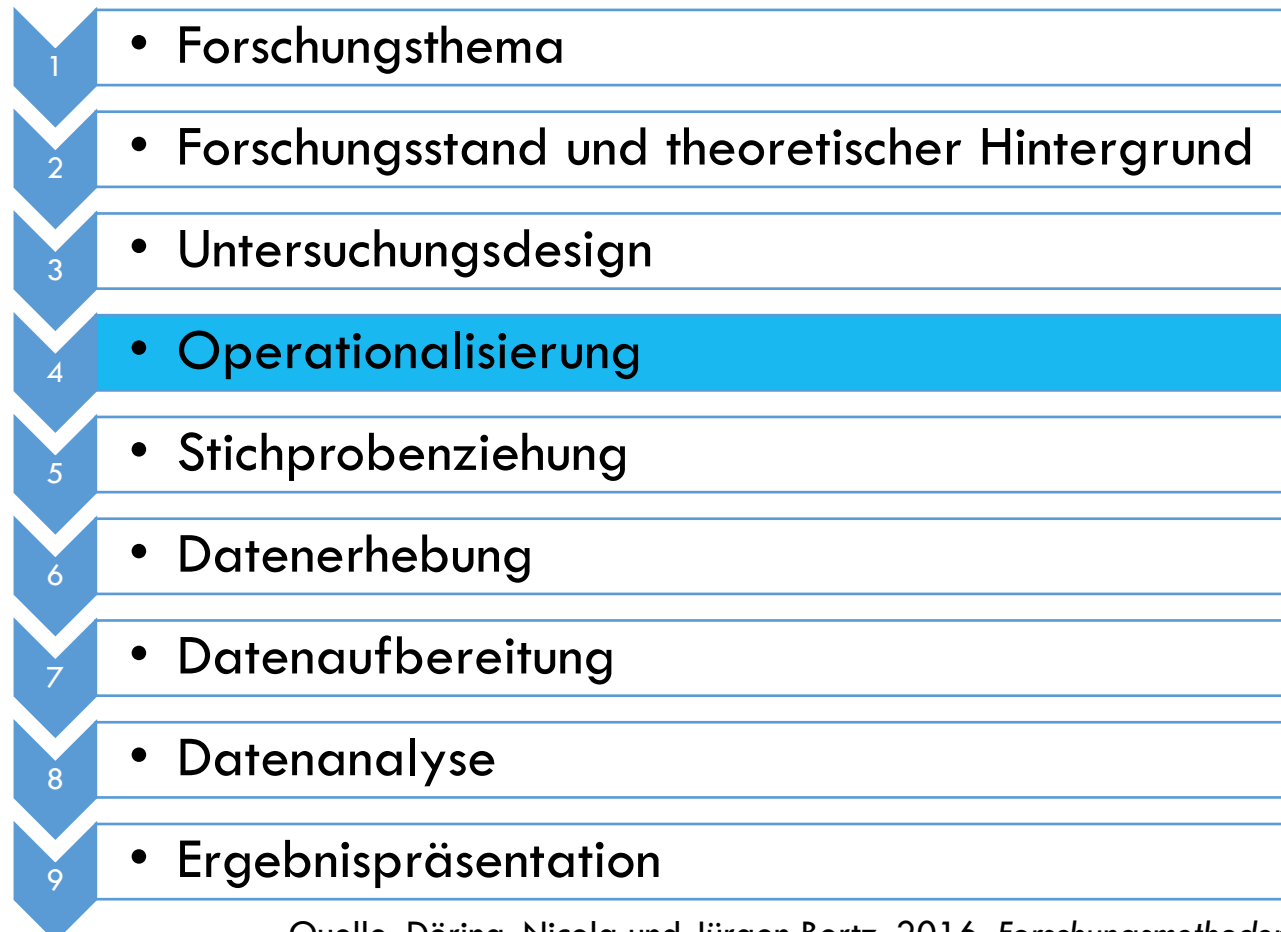
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Hypothesengenerierende
bzw.
gegenstandserkundende
Arbeiten
- Populationsbeschreibende
(deskriptive) Arbeiten
- Hypothesenprüfende
(explanative) Arbeiten

Was bedeutet es, quantitativ zu forschen?

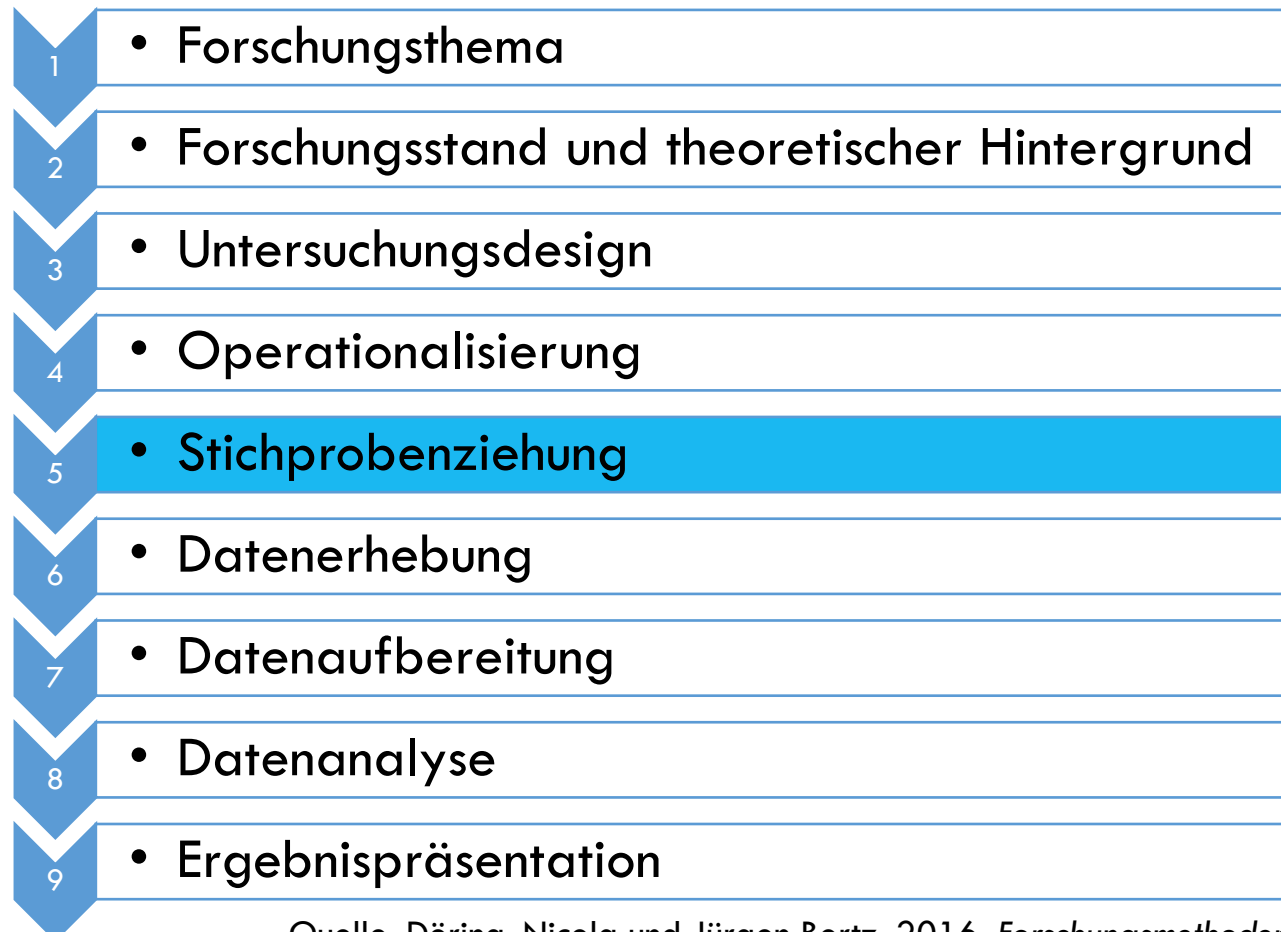
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Präzise Definition aller relevanten Merkmale (Konzeptspezifikation)
- Festlegung des Skalenniveaus jedes Merkmals

Was bedeutet es, quantitativ zu forschen?

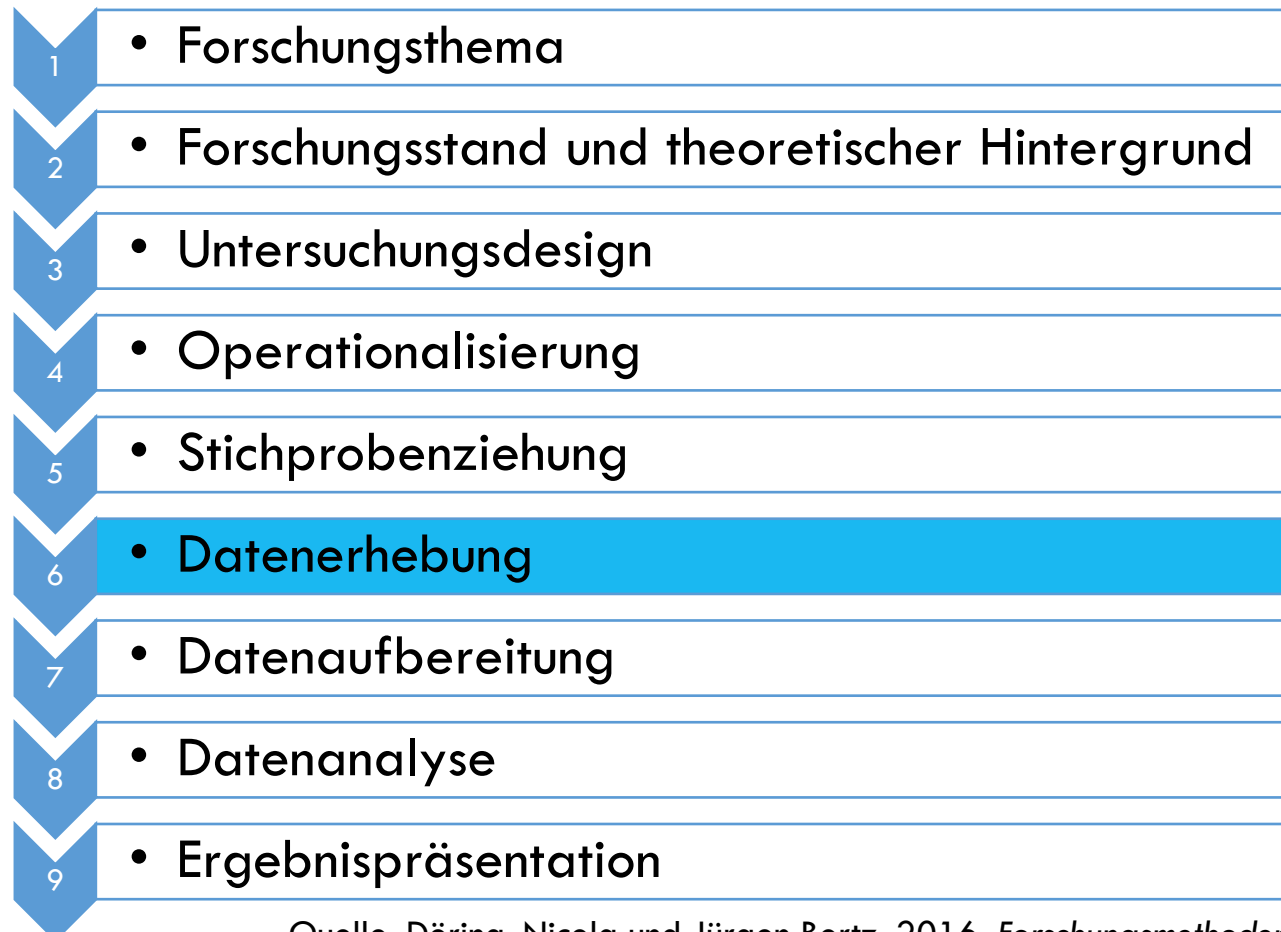
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Vollerhebung/
Stichprobenerhebung
- Stichprobenart
(zufällig/nicht zufällig)
- Stichprobenumfang

Was bedeutet es, quantitativ zu forschen?

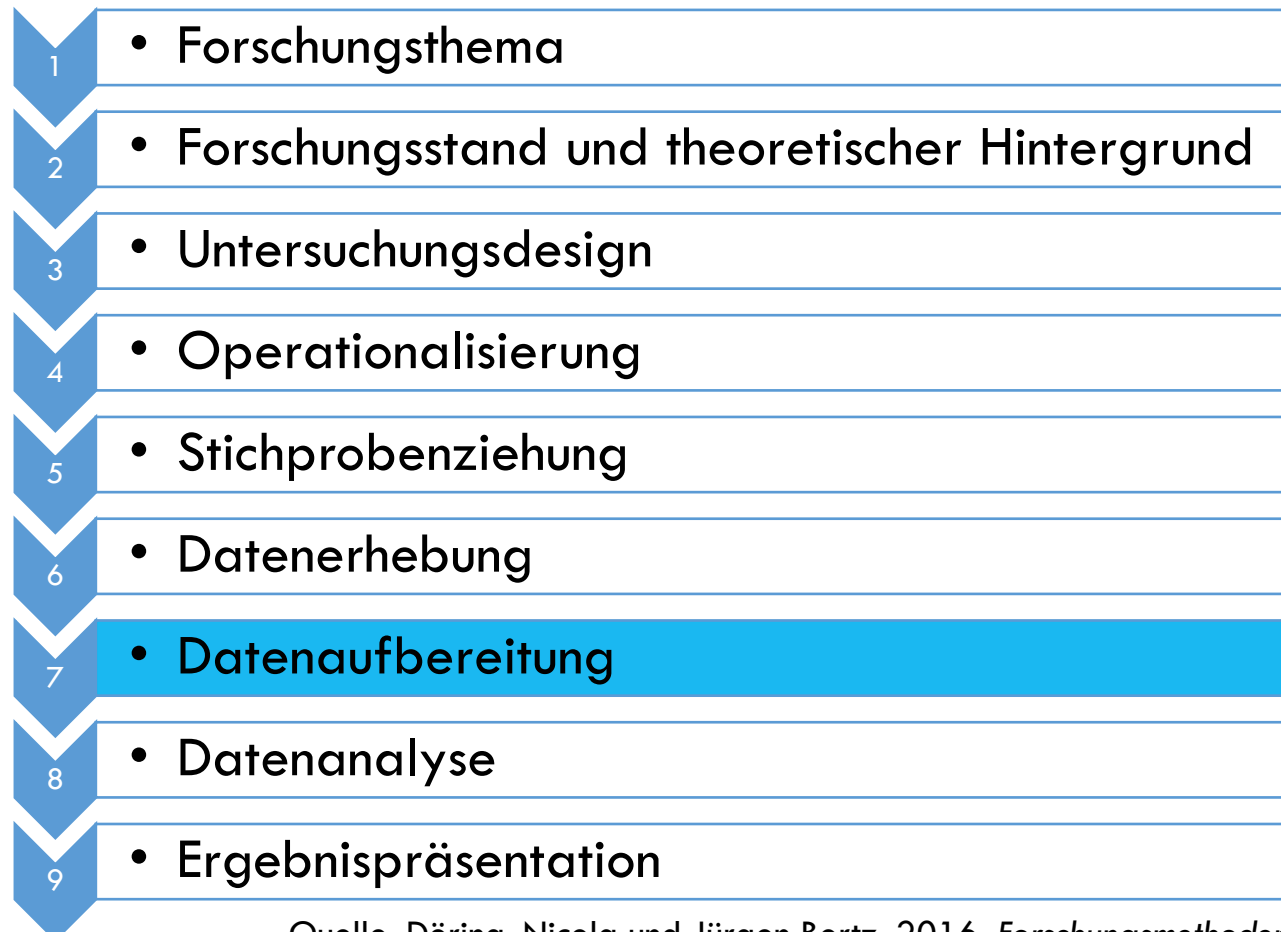
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Strukturierte Beobachtung
- Strukturierte mündliche Befragung (Interview)
- Strukturierte schriftliche Befragung (Fragebogen)
- Psychologischer Test
- Physiologische Messung
 - Quantitative Dokumentenanalyse bzw. Inhaltsanalyse

Was bedeutet es, quantitativ zu forschen?

Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell

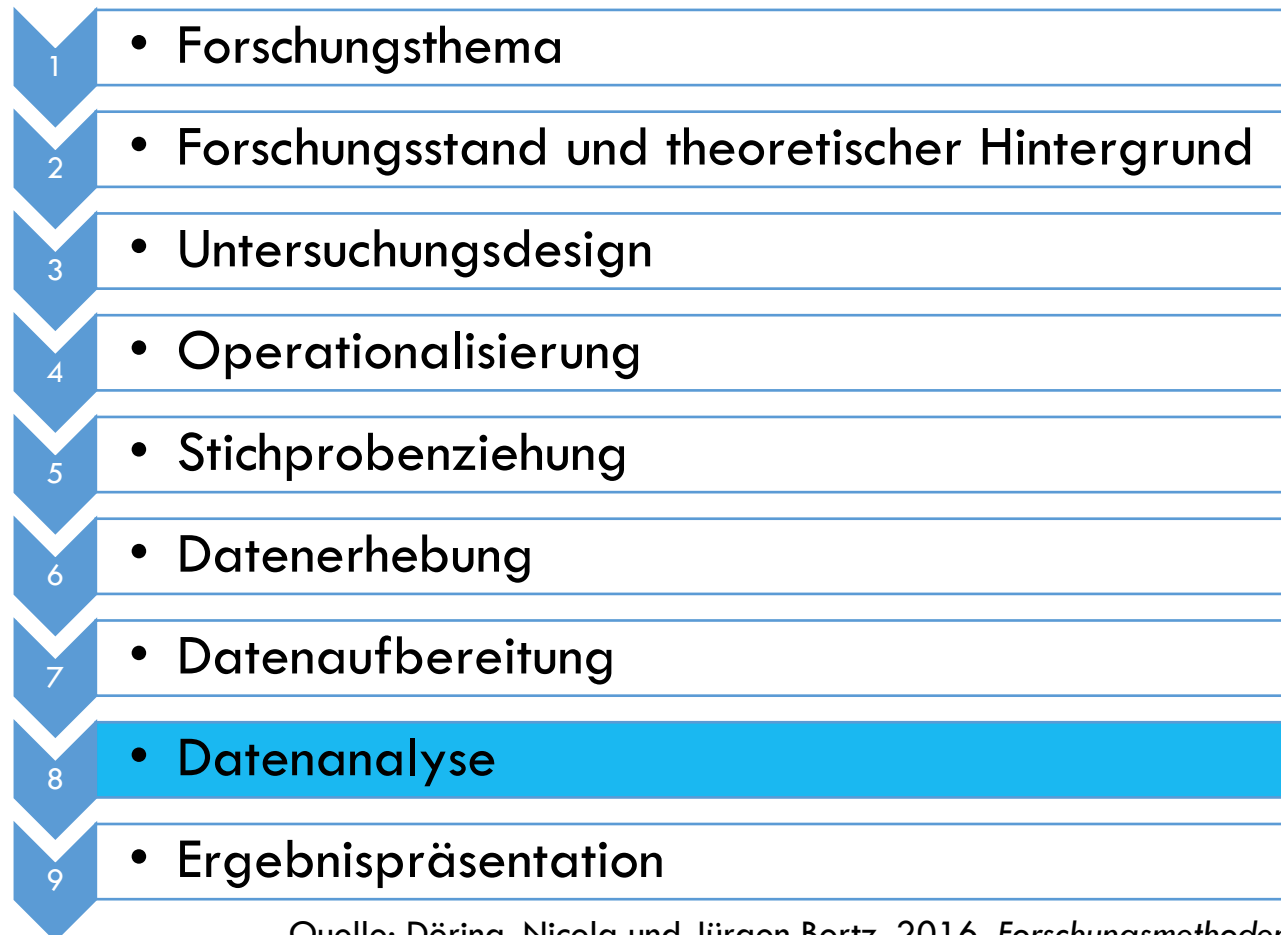


- Aufbereitung des Rohdatenmaterials (etwa Sortierung, Bereinigung um Fehler, Anonymisierung, etc.)

- Ziel: fehlerfreier, vollständiger Datensatz mit validen Werten

Was bedeutet es, quantitativ zu forschen?

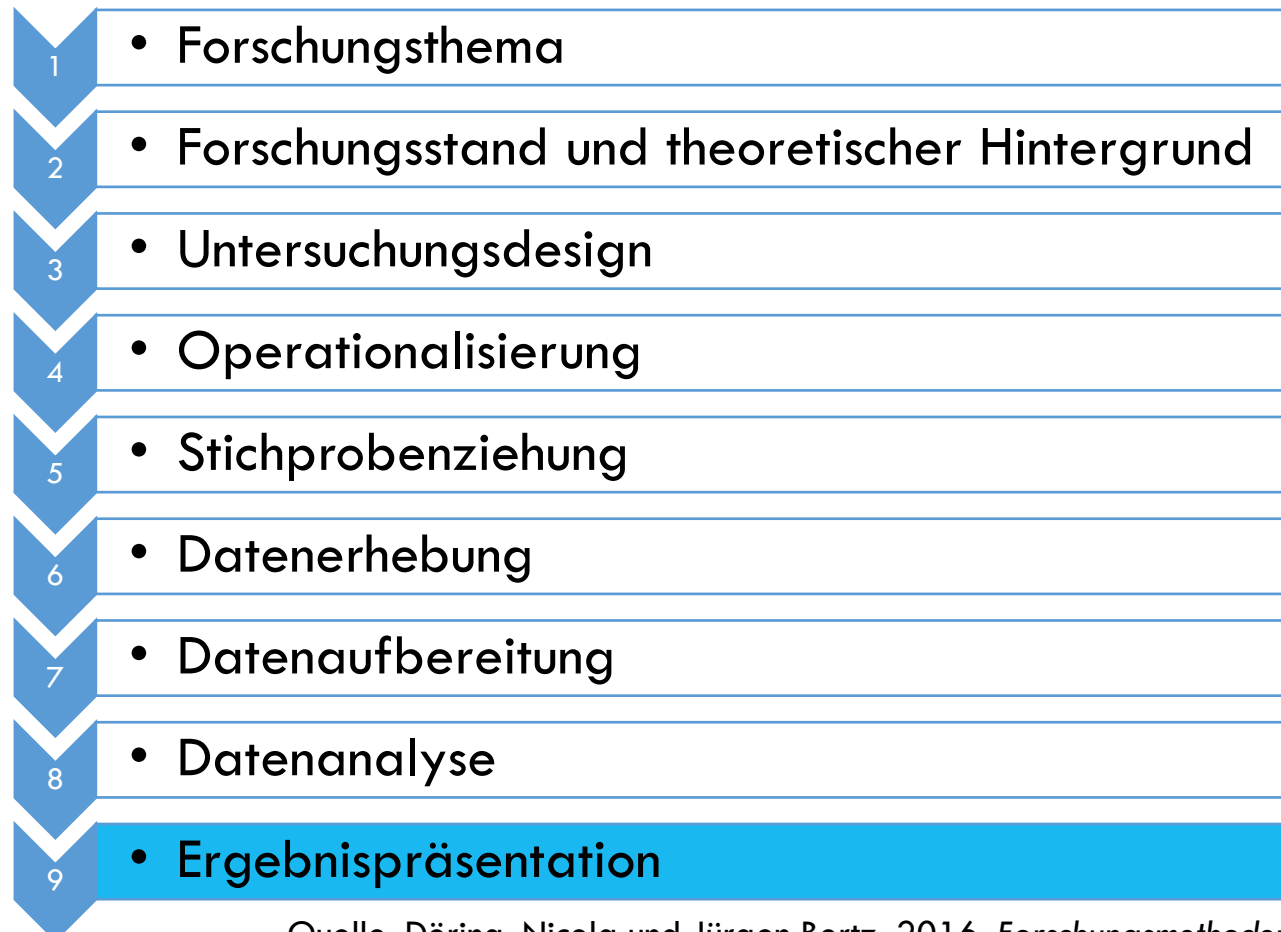
Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



- Hypothesengenerierende bzw. gegenstandserkundende Arbeiten → deskriptive und exploratorische (z.B. grafische) Auswertungen
- Populationsbeschreibende (deskriptive) Arbeiten → inferenzstatistische Schätzung von Populationsparametern mit Verfahren der Punkt- und Intervallschätzung
- Hypothesenprüfende (explanative) Arbeiten → klassische statistische Signifikanztests, Minimum-Effektgrößen-Tests, Strukturgleichungsmodelle

Was bedeutet es, quantitativ zu forschen?

Der quantitative Forschungsprozess im Detail/ Neun-Phasen-Modell



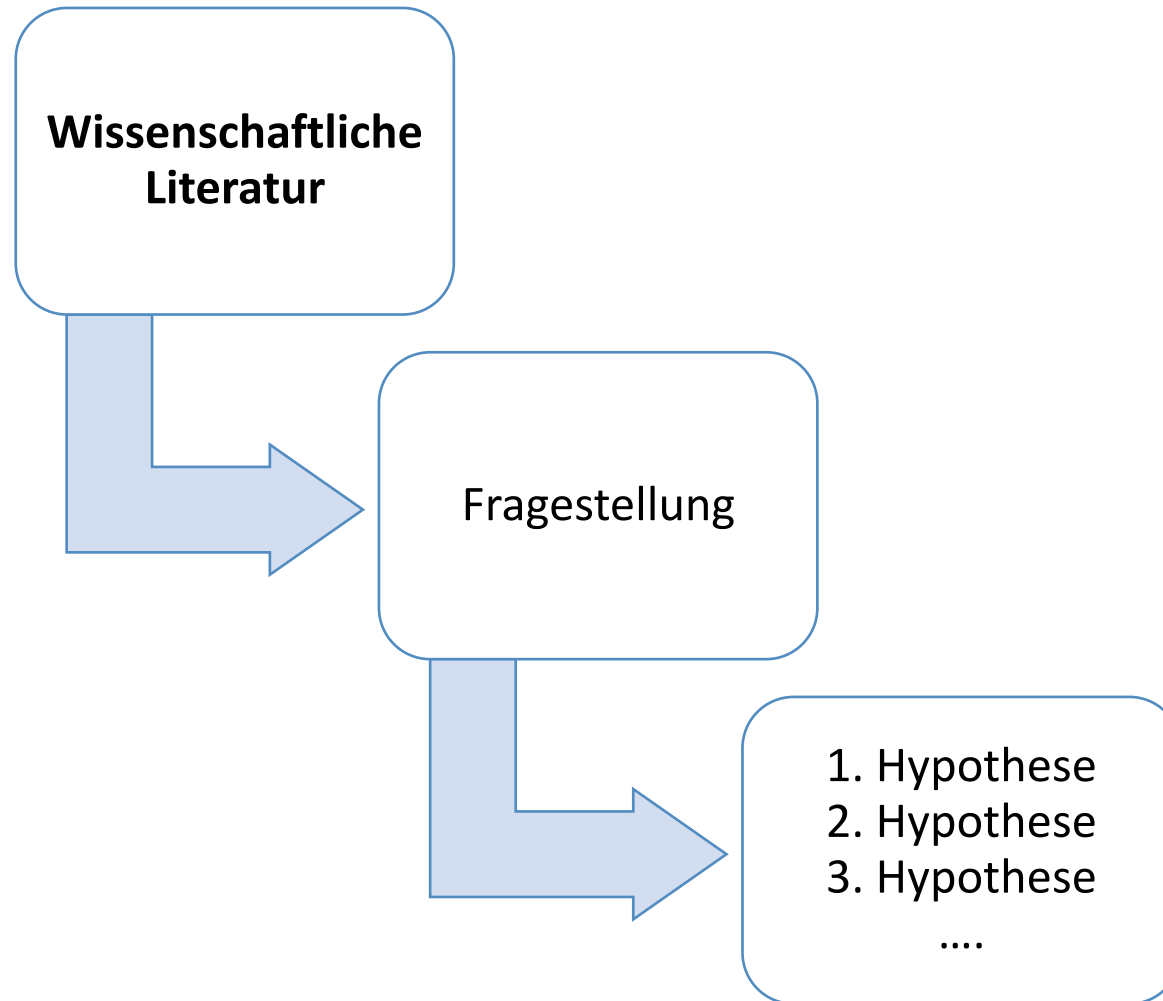
- Zeitschriftenartikel
- Konferenzbeiträge
- Poster

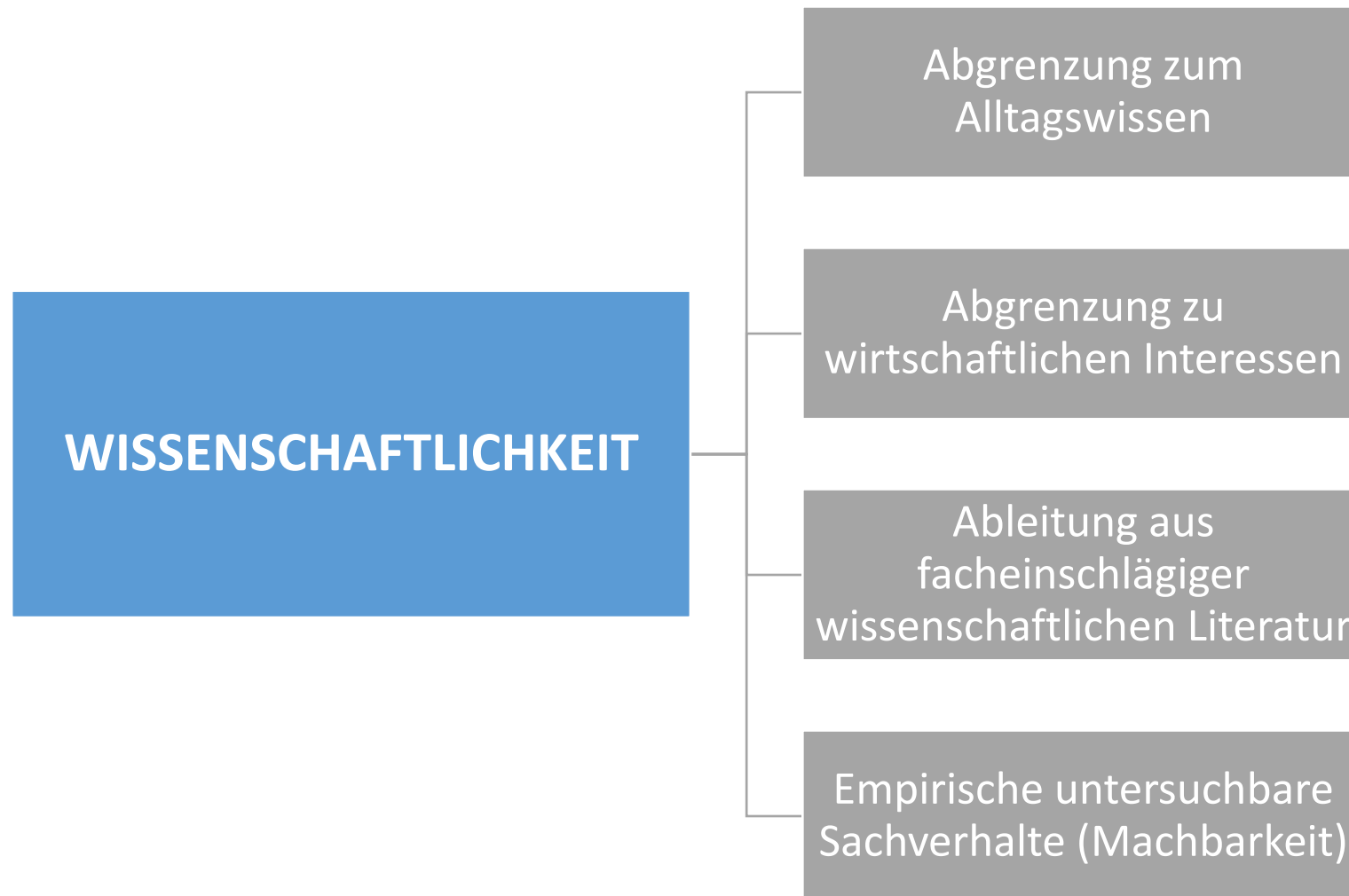
Wissenschaftliche
Abschlussarbeiten → Defensio

FRAGESTELLUNG UND HYPOTHESEN



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education





Exkurs: Was ist Wissenschaft?

Begriffsverständnis

Wissenschaftlicher Erkenntnisgewinn:

„Wissenschaftlicher Erkenntnisgewinn („scientific knowledge gain“) basiert in Erfahrungswissenschaften wie den Sozial- und Humanwissenschaften auf der systematischen Sammlung, Aufbereitung und Analyse von empirischen Daten im Rahmen eines geordneten und dokumentierten Forschungsprozesses. Dabei kommen sozialwissenschaftliche Methoden der Untersuchungsplanung, Stichprobenziehung, Datenerhebung, Datenaufbereitung und Datenanalyse zum Einsatz. Des Weiteren ist der empirische Forschungsprozess theoriebasiert, d. h. in seinem Verlauf werden wissenschaftliche Theorien über den Forschungsgegenstand (sowie über die Forschungsmethodik) angewendet und geprüft oder gebildet und weiterentwickelt. Erst mit Bezug auf Theorien sind empirische Daten sinnvoll interpretierbar.“ (Döring und Bortz 2016, 5).



Exkurs: Was ist Wissenschaft?

Begriffsverständnis

Wissenschaftliche Forschung:

„Wer wissenschaftliche Forschung („scientific research“) betreibt, sucht mithilfe anerkannter wissenschaftlicher Methoden und Methodologien auf der Basis des bisherigen Forschungsstandes (d. h. vorliegender Theorien und empirischer Befunde) zielgerichtet nach gesicherten neuen Erkenntnissen, dokumentiert den Forschungsprozess sowie dessen Ergebnisse in nachvollziehbarer Weise und stellt die Studie in Vorträgen und Publikationen der Fachöffentlichkeit vor.“ (Döring und Bortz 2016, 7).

Exkurs: Was ist Wissenschaft?

Begriffsverständnis



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Neben den empirischen Wissenschaften (dazu gehören neben den Sozial- und Humanwissenschaften auch die Natur- und Technikwissenschaften) gibt es auch nicht-empirische Wissenschaften (dazu gehören die Formalwissenschaften und die Geisteswissenschaften → siehe dazu die Tabelle auf der nächsten Folie) (Döring und Bortz 2016, 12f.)

Übung 3 (10 Minuten)

Einteilung Wissenschaften

Öffnen Sie das folgende Excel-File (zu finden auf Moodle): **LV_Einheit_1_Uebung_3.xlsx**

Ordnen Sie die in Tabellenblatt 1 (benannt mit „Wissenschaften“) aufgelisteten Wissenschaften einer der Kategorien in Tabellenblatt 2 zu (benannt mit „Ordne aus Tabellenbl. 1 zu“)



Exkurs: Was ist Wissenschaft?



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Begriffsverständnis

Nicht-empirische Wissenschaften		Empirische Wissenschaften/ Erfahrungswissenschaften		
Formalwissenschaften („formal sciences“)	Geisteswissenschaften („humanities“)	Sozialwissenschaften: auch Humanwissenschaften, Gesellschafts- wissenschaften („social sciences“)	Naturwissenschaften („natural sciences“)	Technikwissenschaften („engineering sciences“)
Philosophie	Theologie	Psychologie	Physik	Maschinenbau
Mathematik	Rechtswissenschaft	Medizin	Chemie	Elektrotechnik
etc.	Geschichte	Erziehungswissenschaft	Biologie	Bauingenieurwesen
	Literaturwissenschaft	Soziologie	Geowissenschaften	Verfahrenstechnik
	Sprachwissenschaft	Wirtschaftswissenschaft	Astronomie	Informatik
	Medienwissenschaft	Kommunikations- wissenschaft	etc.	etc.
	etc.	etc.		

Exkurs: Was ist Wissenschaft?

Ansprüche

Wissenschaft kann auch definiert werden, indem aufgezählt wird, was ihre grundlegenden Elemente sind und welche Anforderungen ein Text, eine Rede oder eine Behauptung erfüllen muss, um als „wissenschaftlich“ zu gelten (vgl. dazu noch einmal die Definition von Döring und Bortz zum Begriff der „wissenschaftlichen Forschung“):

- Präzise Eingrenzung des Geltungsbereichs
- Methodische Herangehensweise: zielgerichtete Untersuchung
- Intersubjektivität: Nachvollziehbarkeit der einzelnen Untersuchungsschritte
- Öffentlichkeit: Wissen wird publiziert und lässt sich so kritisch überprüfen

(Dahinden, Sturtzenegger & Neuroni, 2014)

Exkurs: Was ist Wissenschaft?

Abgrenzung



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Nachdem nun einige Vorschläge genannt wurden, wie man den Begriff „Wissenschaft“ definiert, muss jetzt noch geklärt werden, wie sich die „Wissenschaft“ bzw. „wissenschaftlicher Erkenntnisgewinn“ von der „Nicht-Wissenschaft“ (Autoritätspersonen, Religion, Tradition, etc.) bzw. „nicht-wissenschaftlichem Erkenntnisgewinn“ abgrenzen lässt...

Exkurs: Was ist Wissenschaft?

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Autoritätspersonen	Man stützt sich auf die Aussagen von Autoritätspersonen bzw. Experten	Unterschiedlicher Grad an tatsächlicher Expertise; subjektive Meinungen/Interessen von Experten, widersprüchliche Aussagen unterschiedlicher Experten	Experte A: Vitamin C stärkt das Immunsystem und schützt vor Erkältungen Experte B: Vitamin-tabletten sind bei normaler Ernährung völlig überflüssig	Wissenschaftliche Studien zeigen, dass Vitamin C (täglich mind. 0,2 g) bei der Normalbevölkerung zwar nicht die Häufigkeit, aber die Dauer und Schwere von Erkältungen reduziert (Hemilä, Chalker & Douglas, 2007)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Religion	Man stützt sich auf religiöse Dogmen und Schriften	Anerkennung göttlicher Offenbarung setzt entsprechenden Glauben voraus und steht im Widerspruch zur verbreiteten säkularen Weltanschauung	Kreationismus, wie er in den USA zum Teil an Schulen gelehrt wird; laut biblischer Schöpfungslehre wurden das Universum und die Erde, der Mensch und alle anderen Lebewesen tatsächlich durch Gott erschaffen	Wissenschaftliche Theorien und Befunde aus der Physik (Urknalltheorie) und der Biologie (Evolutionstheorie) widersprechen dem Kreationismus und werden mittlerweile auch vom Mainstream der christlichen Kirchen anerkannt

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Tradition	Man stützt sich auf überliefertes Wissen früherer Generationen	Tradiertes Wissen basiert oft auf Missverständnissen, Fehlern, Mythen etc.	Seit Generationen wird überliefert, dass man beim Husten und Niesen die Hand vor den Mund halten soll, um andere nicht anzustecken	Aus wissenschaftlicher Sicht wird heute abgeraten, in die Hand zu niesen oder zu husten, da Erreger an der Hand lange überleben und durch Händedruck, Berührung von Gegenständen etc. ständig weitergegeben werden (stattdessen sollte ein Taschentuch oder die Ellenbeuge genutzt werden)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Gesunder Menschenverstand	Man stützt sich auf den gesunden Menschenverstand („common sense“) als geteilte Überzeugung einer Gruppe	Was als „gesunder Menschenverstand“ angesehen wird, variiert zwischen sozialen Gruppen sehr stark und ist oft auch von Vorurteilen, Gruppeninteressen etc. geprägt	In einer Talkshow zum Thema Adoption durch gleichgeschlechtliche Paare werden verschiedene Positionen vertreten – alle berufen sich auf den gesunden Menschenverstand Position A: Es wird doch niemand bestreiten wollen, dass Kinder immer Vater und Mutter brauchen Position B: Es ist doch allgemein bekannt, dass Kinder stabile Bezugspersonen brauchen – unabhängig von Geschlecht, sexueller Orientierung oder Verwandtschaftsgrad	Ein empirischer Vergleich von $n = 106$ Adoptivkindern, die bei Frauen-, Männer- oder heterosexuellen Paaren aufwachsen, zeigte keinen Einfluss der sexuellen Orientierung der Eltern auf die Entwicklung der Kinder. Diese hing von anderen Faktoren ab (z.B. Erziehungsstil sowie Beziehungsqualität der Eltern; Farr, Forssell, & Patterson, 2010)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Intuition	Man stützt sich auf seine eigene Intuition („Bauchgefühl“, „Instinkt“)	Das „Bauchgefühl“ ist durch viele Einflussfaktoren, insbesondere auch durch Vorurteile, Wunschdenken, etc. beeinflusst	Das Bauchgefühl sagt einem, ein internet- öffentliches Sexualstraftäter- Verzeichnis würde die Sicherheit im Wohnbezirk erhöhen	Evaluationsstudien zu Effekten auf die Sicherheit fehlen, allerdings liegen Befragungsstudien zur Einstellung von Laien und Experten gegenüber den Verzeichnissen vor, denen eine sehr skeptische Haltung von Experten zu entnehmen ist (Malesky & Keim, 2001; Salerno et al., 2010)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education

Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Anekdotische Evidenz	Man stützt sich auf die Lebenserfahrungen und/oder Beispiele aus dem Umfeld oder den Medien	Persönliche Lebenserfahrungen sind sehr stark verzerrt durch Merkmale der eigenen Person sowie des eigenen kulturellen und sozialen Umfeldes	Nach fünf enttäuschenden Verabredungen mit Internet-Bekanntschaften schlussfolgert man, dass Online-Partnersuche in Wirklichkeit gar nicht funktioniert	Wissenschaftliche Studien zum Online-Dating zeigen, dass und inwiefern es sowohl Vorteile als auch Nachteile birgt (Finkel et al., 2012)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

Exkurs: Was ist Wissenschaft?

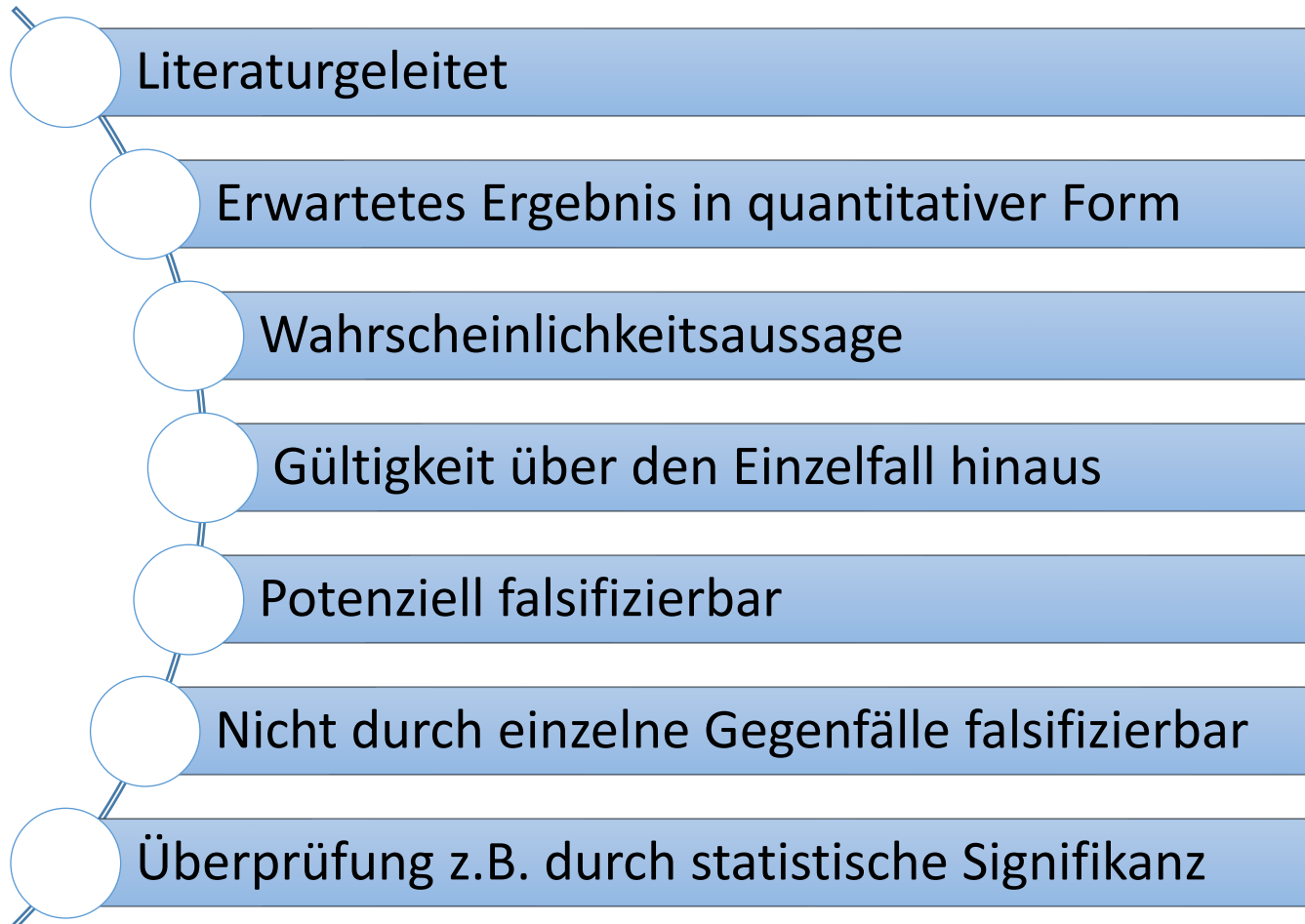
Abgrenzung

Nicht-wissenschaftliche Produktion und Begründung von Wissen durch Berufung auf...	Vorgehensweise	Grenzen	Beispiele	Wissenschaftliche Erkenntnisse
Logik	Man stützt sich auf logische Argumente	In sozialen Zusammenhängen handeln Menschen oft nicht logisch oder rational	Gegner von Frauenförderprogrammen behaupten: Bei der Besetzung von Führungspositionen gibt es keine geschlechtsbezogene Benachteiligung von Frauen, denn es wäre doch unlogisch, wenn Unternehmen sich gute Bewerberinnen entgehen ließen	Der wissenschaftliche Forschungsstand zur Situation weiblicher Führungskräfte weist verschiedene Benachteiligungen von Frauen im Berufsleben nach (Eagly & Carlib, 2003)

Varianten nicht-wissenschaftlicher Produktion und Gegenüberstellung zu wissenschaftlichem Erkenntnisgewinn.

In Anlehnung an Döring und Bortz 2016, 6f.

HYPOTHESEN in quantitativen Untersuchungssettings



HYPOTHESENPAARE

NULLHYPOTHESE (H_0)

ALTERNATIVHYPOTHESE (H_1)

- gerichtet (einseitig)
- ungerichtet (zweiseitig)

HYPOTHESEN-TYPEN

Zusammenhangshypothesen

Unterschiedshypothesen

Veränderungs-/
Vorhersagehypothesen

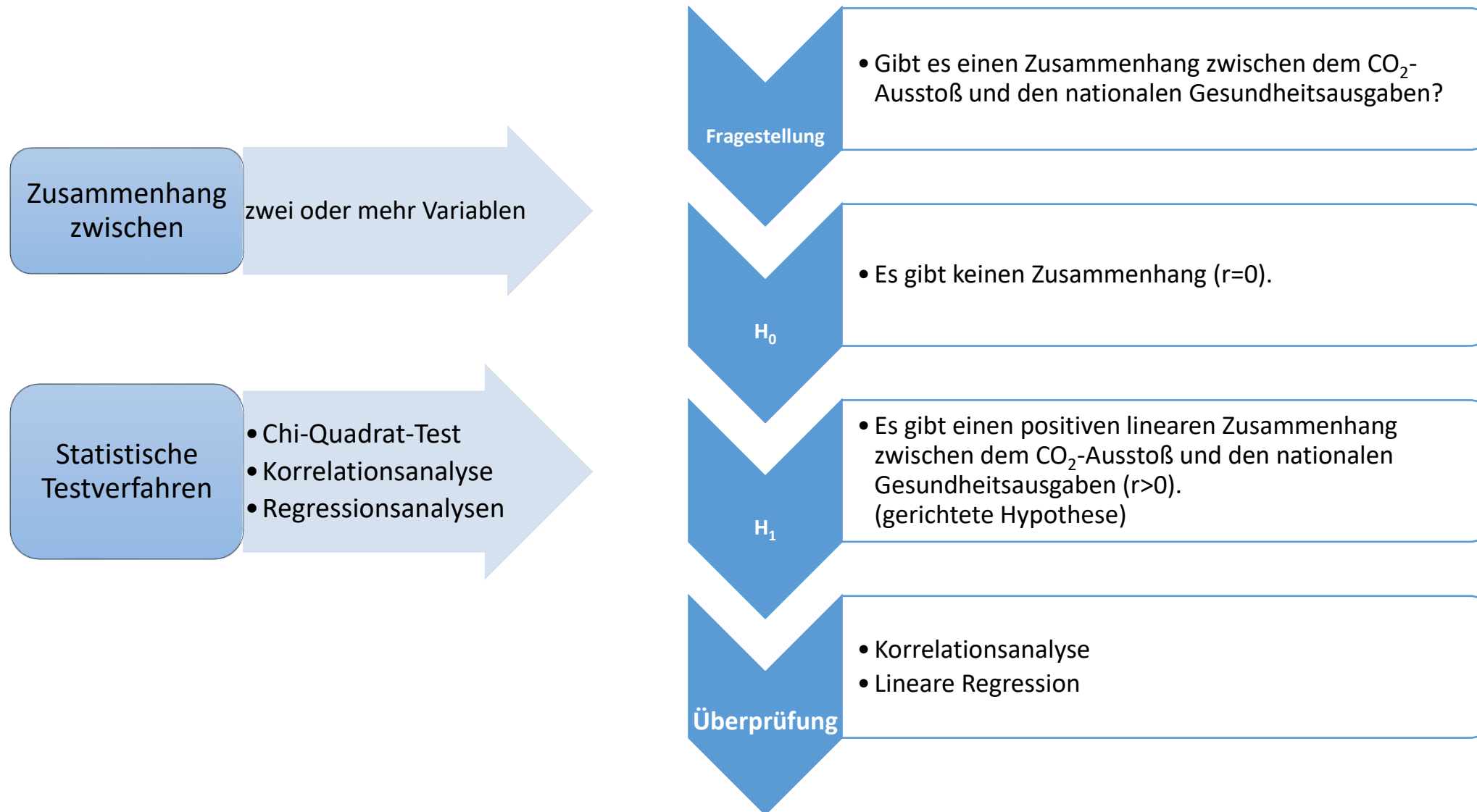
Explorativ/strukturendeckende
Hypothesen

Zusammenhangshypothesen

Beispiel

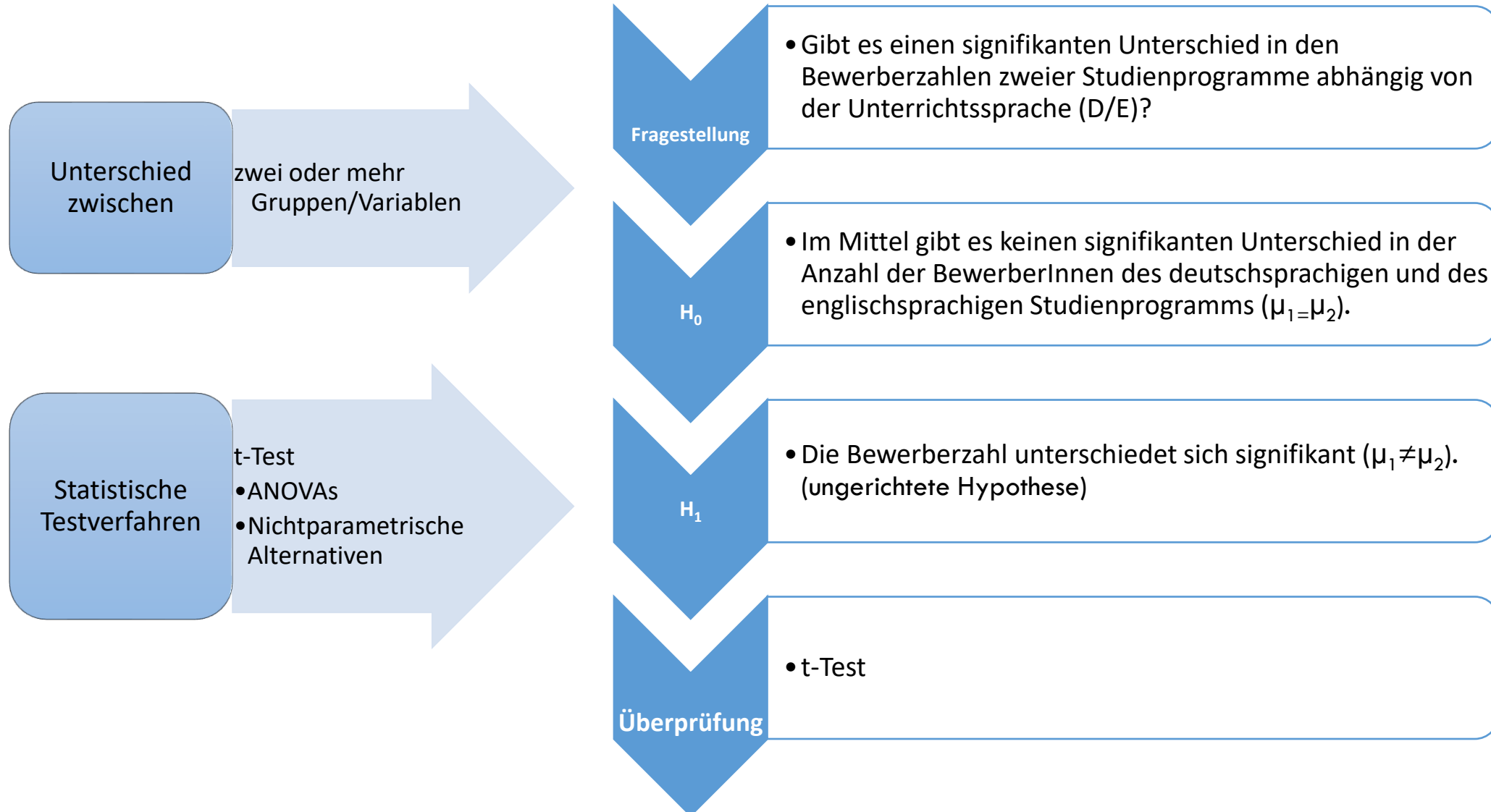


**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education



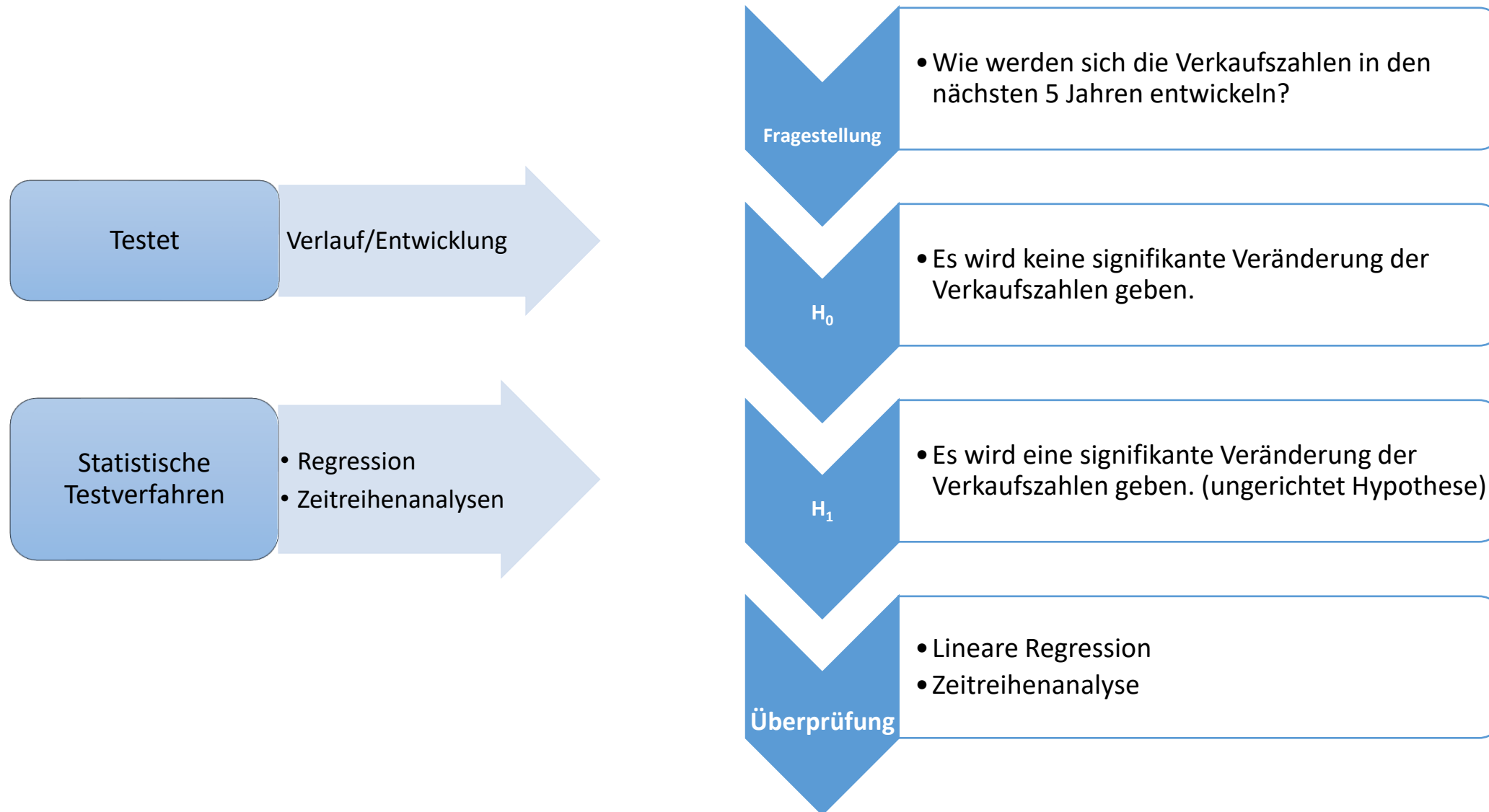
Unterschiedshypothesen

Beispiel

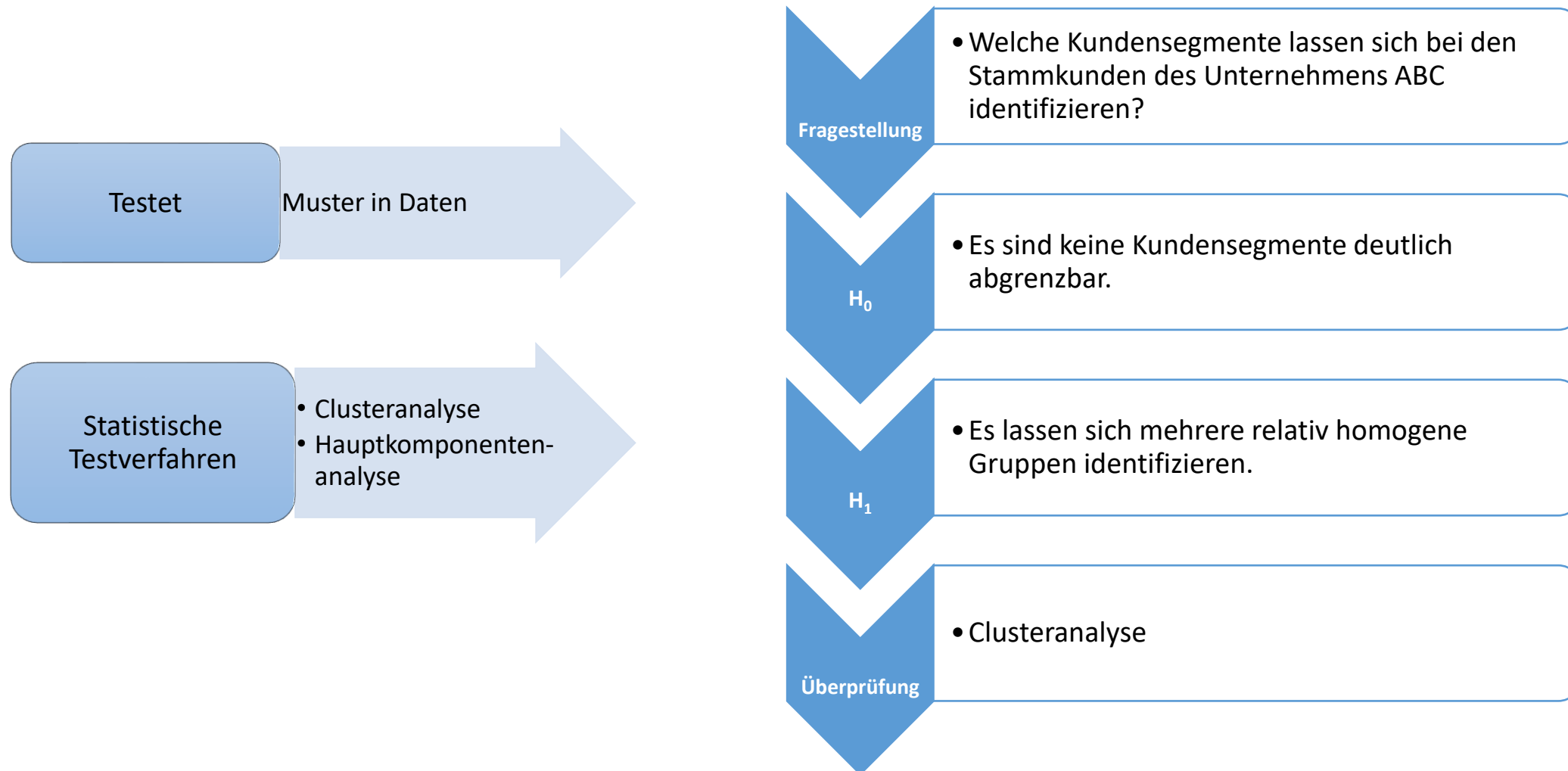


Veränderungs-/Vorhersagenhypothesen

Beispiel



Explorative/strukturendeckende Hypothesen Beispiel





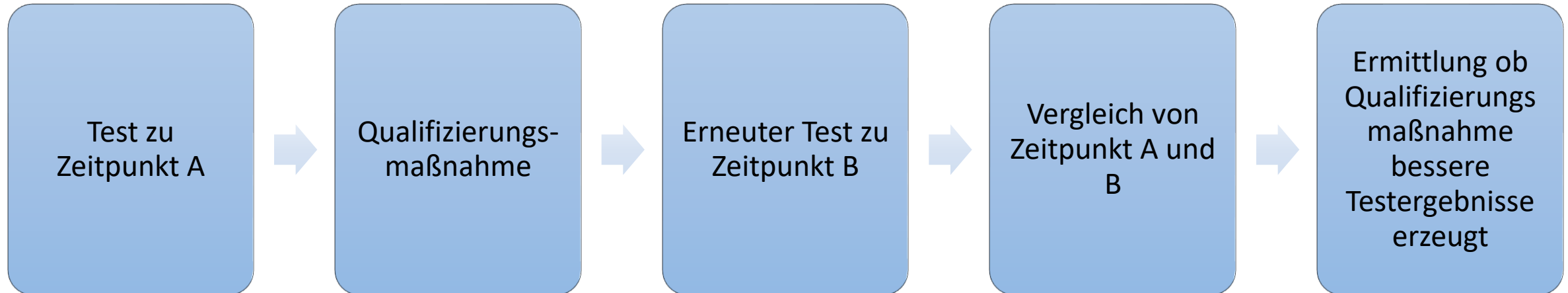
Querschnittstudie

- **einmalige** Datenerhebung zu einem Zeitpunkt / in **abgegrenztem** Zeitraum

Längsschnittstudie

- Erhebung von Vergleichswerten zu **unterschiedlichen** Zeitpunkten

Längsschnittstudie – Beispiel



Datenquellen

Unternehmens- Daten

- Datenbanken
- Flatfiles

Online-Datenquellen

(Wirtschaftsdaten,
demografische Daten,
Unternehmensdaten, etc.)

- <https://www.google.com/publicdata>
- <http://www.gapminder.org/>
- <http://de.statista.com>
- <http://www.global-markets-companies.com/Default.aspx>
- <http://data.un.org/>
- <http://hdr.undp.org/en/data>
- <http://data.worldbank.org/>
- http://www.statistik.at/web_en/statistics/index.html
- <http://ec.europa.eu/eurostat>

Eigene Datenerhebung

- Beobachtung (Visitor Counts, Traffic Count)
- Messungen (Sensordaten, Systemdaten)
- Befragungen (Fragebogenerhebung, Telefonumfrage)

Bestehende Fragebögen

- Gibt es zur gegebenen Fragestellung bereits ein Erhebungsinstrument?
- Bereits **entwickelte und validierte** Fragebögen/Skalen einsetzen
- **EIGENTUMSRECHTE** des Erhebungsinstruments sind zu beachten!
(Erlaubnis des Autors / der Autorin einholen)

Datenstruktur

- **Datenerhebung und mögliche Auswertungsmethode schon bei Erstellung des Fragebogens mitbedenken.**
- Daten müssen gewissen Kriterien entsprechen
- Variable als numerischer Wert -> Regression, ANOVA, etc.
- Variable als Kategorie -> Gruppierungsvariable, Chi-Quadrat-Test, Korrelation, etc.

Mehrere Testläufe durch unterschiedliche Personen durchführen

- Sind alle Fragen verständlich?
- Sind alle benötigten Antwortoptionen verfügbar / alle Eingaben möglich?
- Ist die technische Umsetzung gelungen (Filter testen!)?

Die Testpersonen sollten nicht an der späteren realen Befragung teilnehmen.

Testdaten $\neq$ Daten der realen Befragung!

Validierung

Erste Befragungsrunde im Feld

Überprüfung der Fragebogenqualität durch unterschiedliche
Analyseverfahren

- Item-Analyse
- Reliabilitätsanalyse

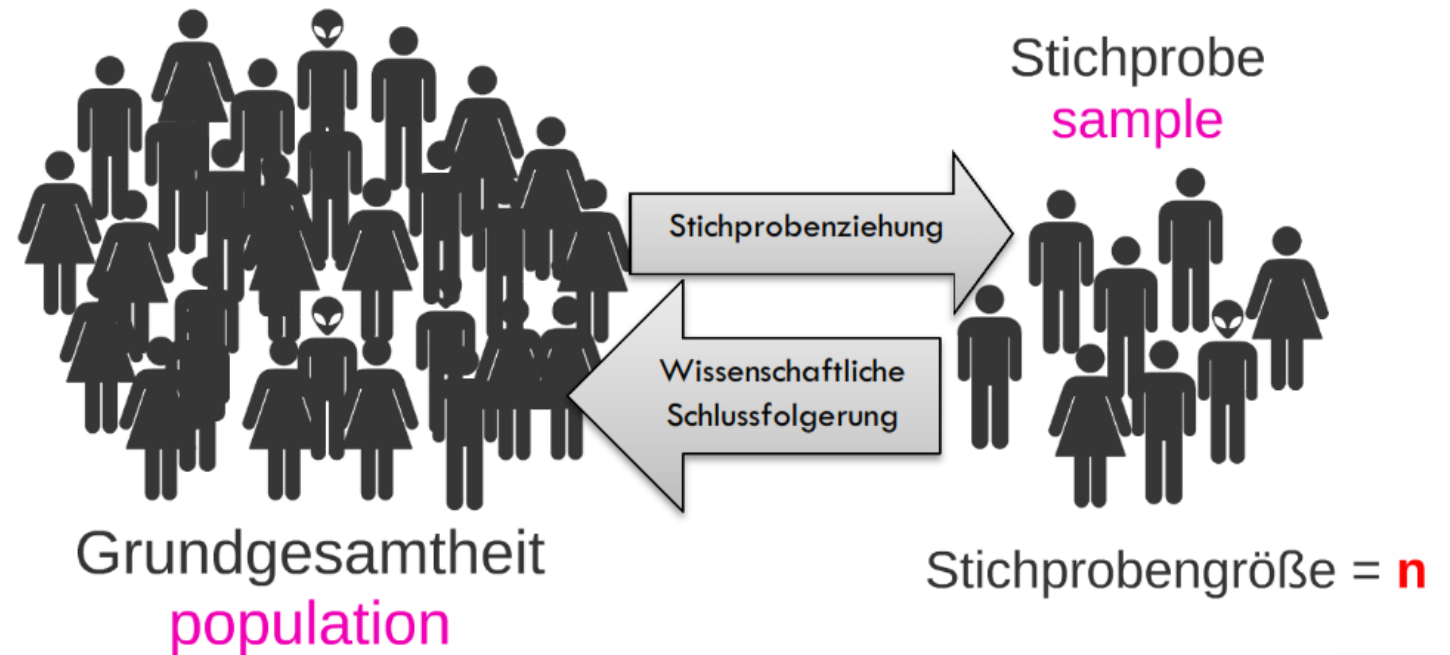
Validierung nicht möglich

- > Rückgriff auf bereits validierte Instrumente
- > Expertenrat

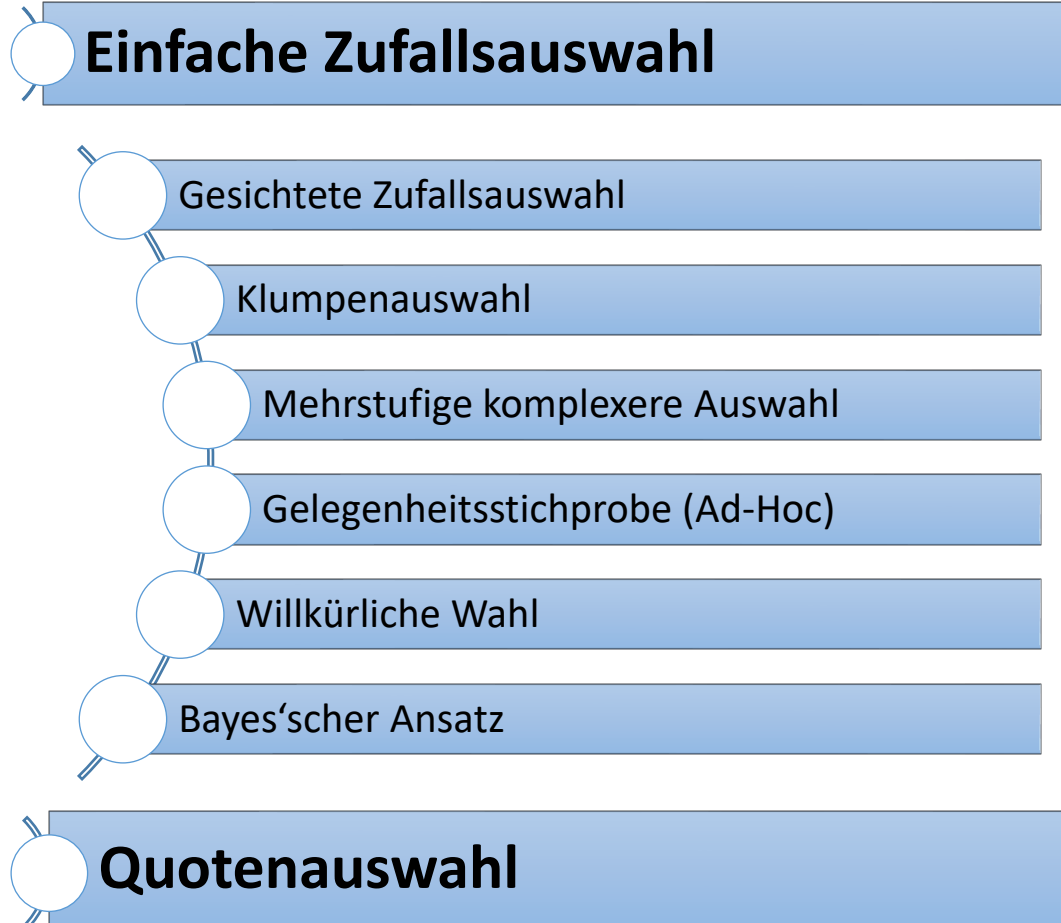
Repräsentativität

Stichprobe =
repräsentatives Abbild

Verallgemeinerung der
Schlüsse



Methoden der Stichprobenziehung (Bortz & Döring 2006, Kapitel 7)



Einfache Zufallsauswahl

**Zufällige Auswahl der Probanden
aus der Population**

*z.B.: Wie denken die ÖsterreicherInnen über
die EU?*

Auswahl der Probanden zufällig aus allen
ÖsterreicherInnen (Melderegister,
Telefonbuch o.ä.)

Problem: Ist Zufälligkeit sichergestellt?

Quotenauswahl

Soziodemografische Faktoren

- Beispiel:
Online-Shopping Nutzungsverhalten -> Alter, Geschlecht;
Konsum von Luxusgütern -> Einkommen

Für die Repräsentativität sind diese Faktoren **bereits bei der Zusammenstellung** der Stichprobe **proportional** zu berücksichtigen.

Quote

- Anzahl der Personen **jeder** Gruppe
- Entspricht dem **Verhältnis** in der Population
- **vorab** festzulegen!

Quotenauswahl - Beispiel

Absolute Häufigkeiten der Bevölkerung ab 18 Jahren zu Jahresbeginn 2016 aufgeteilt nach Alter und Geschlecht

Alter	Geschlecht	
	Frauen	Männer
18 bis 30 Jahre	710368	748380
31 bis 40 Jahre	570046	579475
41 bis 50 Jahre	661965	662971
51 bis 60 Jahre	638864	631773
61 bis 70 Jahre	466178	422765
71 und mehr Jahre	647678	447221

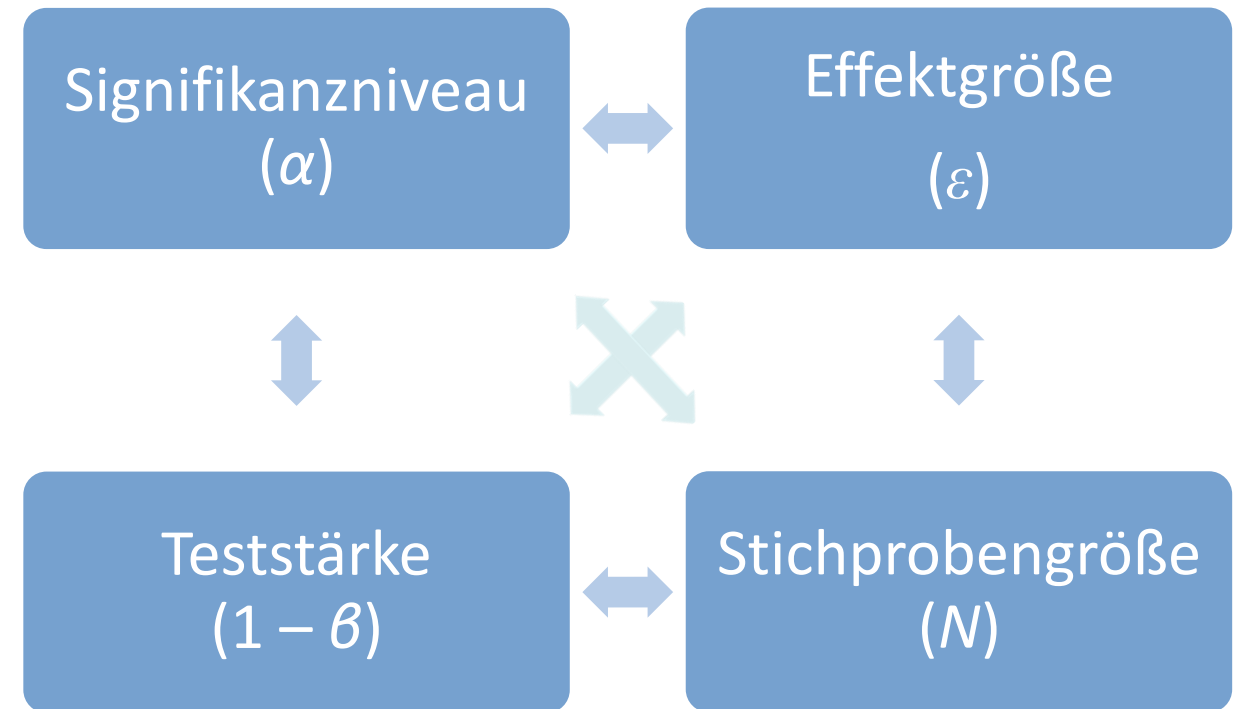
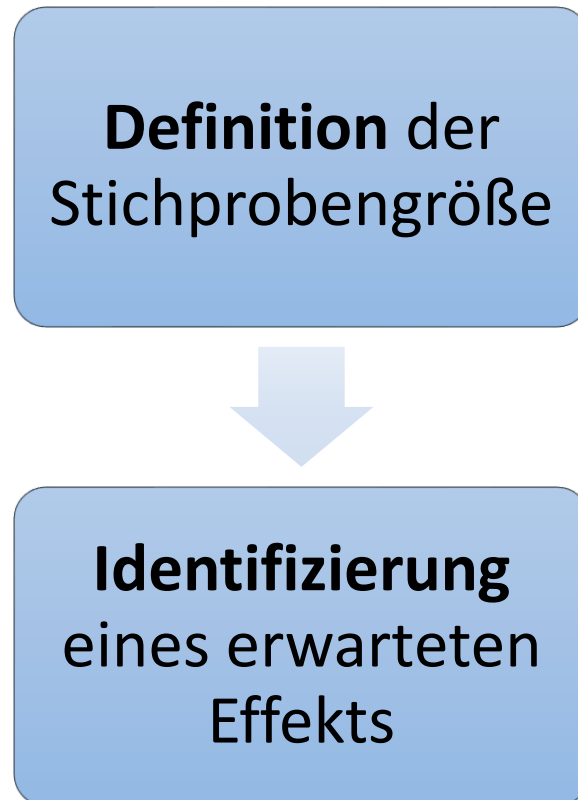
(Statistik Austria, 2016)

Quoten für die Stichprobe

Alter	Geschlecht	
	Frauen	Männer
18 bis 30 Jahre	5	5
31 bis 40 Jahre	4	4
41 bis 50 Jahre	5	5
51 bis 60 Jahre	4	4
61 bis 70 Jahre	3	3
71 und mehr Jahre	5	3

Statistik Austria (2016). Bevölkerung nach Alter und Geschlecht. Zugriff am 29.10.2016. Abgerufen von http://www.statistik.at/web_de/statistiken/menschen_und_gesellschaft/bevoelkerung/bevoelkerungsstruktur/bevoelkerung_nach_alter_geschlecht/index.html

Stichprobengröße



Berechnung der Stichprobengröße

Die Berechnung der optimalen Stichprobengröße kann auf Grundlage der Werte für

- $\alpha = 0,05$

- $1 - \beta = 0,80$ und

- ϵ (Effektgröße) = klein/mittel/groß

mittels G*Power erfolgen: <http://www.gpower.hhu.de>

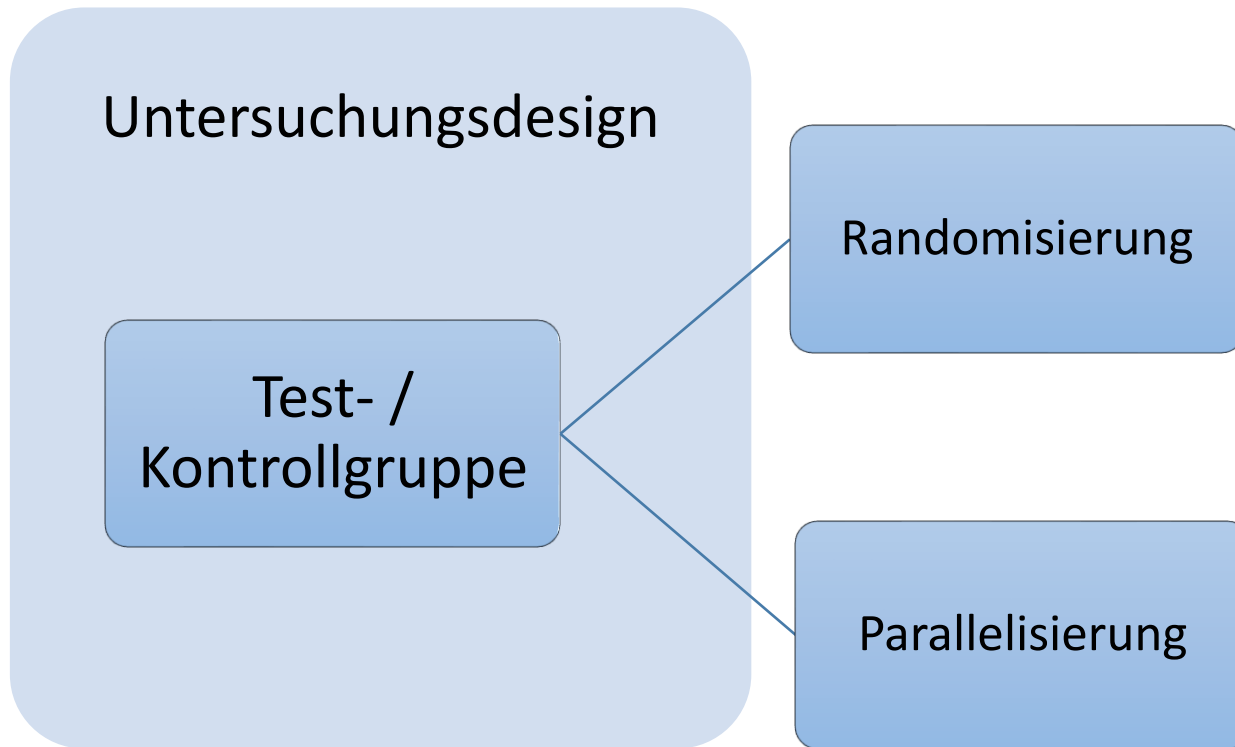
Die zu erwartenden Effektgrößen können aus nachfolgender Tabelle abgelesen werden:

Effektgröße	Cohen's d	Pearson's r	Pearson's r adaptiert	Eta ²
klein	0,2	0,1	0,10	0,01
mittel	0,5	0,3	0,24	0,06
groß	0,8	0,5	0,37	0,14

Test- und Kontrollgruppe



**FACHHOCHSCHULE
WIENER NEUSTADT**
Austrian Network for Higher Education



- Zufällige Zuteilung von Personen zu Untersuchungsgruppen
- Vergleichbare Untersuchungsgruppen hinsichtlich des interessierenden Merkmals
- Pilottestung und anschließende Zuordnung
- Bei kleinen Stichproben: matched samples