

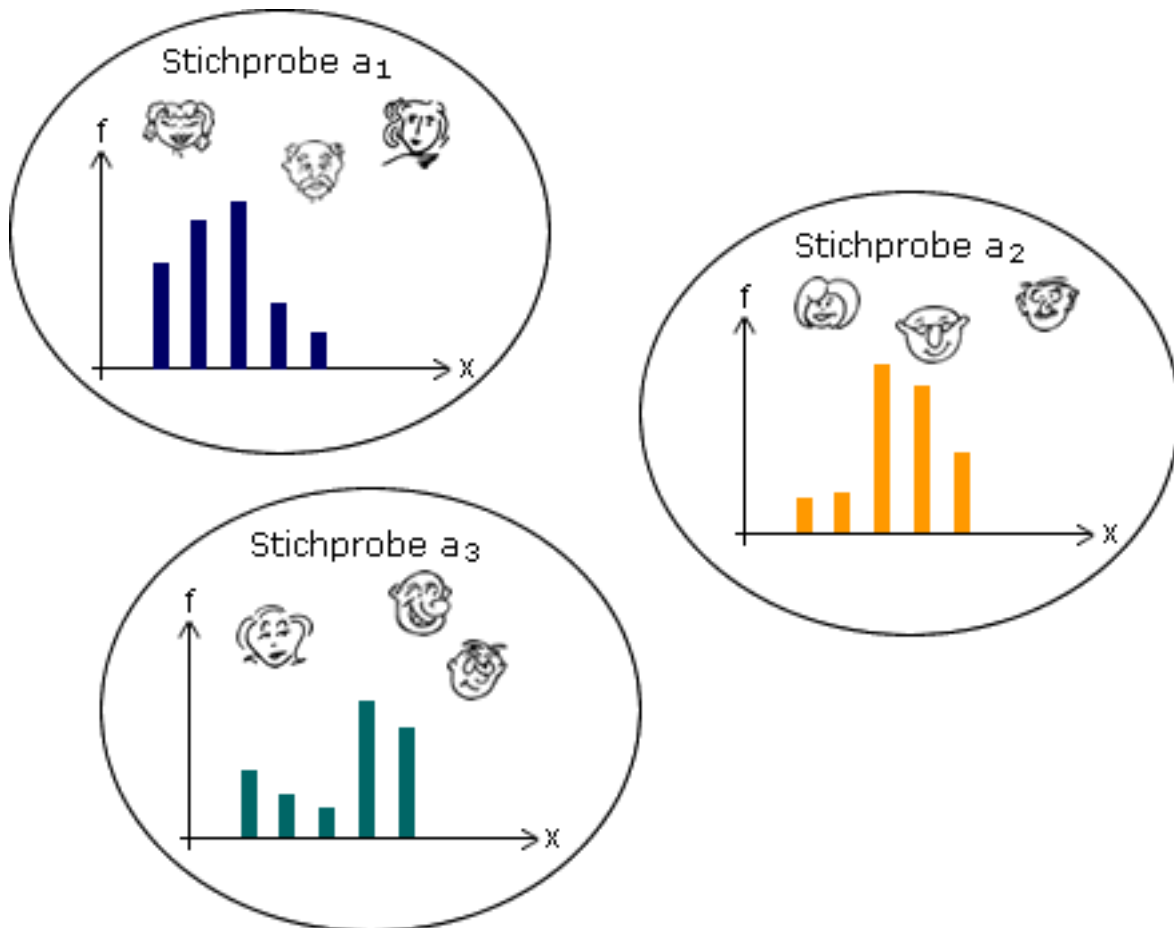
Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen

Inhaltsverzeichnis

Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen	2
Lernhinweise	2
Einführung	3
Teil 1 - Zum Grundprinzip parametrischer Varianzanalysen	
Teil 2 - Zum Grundprinzip parametrischer Varianzanalysen	
Teil 3 - Zum Grundprinzip parametrischer Varianzanalysen	
Voraussetzungen parametrischer Varianzanalysen ohne Messwiederholungen	
Theorie (1 - 3)	6
Teil 1 - Theorie	
Teil 2 - Theorie	
Teil 3 - Theorie	
Fallbeispiel	11

Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen

Lernhinweise



Im Rahmen univariater, einfaktorieller Varianzanalysen ohne Messwiederholungen definieren die Ausprägungsgrade der unabhängigen Variablen (Faktor A) eine Reihe unabhängiger Stichproben ($a_1, a_2, \dots, a_p$), in denen das interessierende Merkmal (abhängige Variable X) erhoben wird. Dieser Lernschritt gilt dem simultanen Vergleich der Stichprobendaten bezüglich ihrer zentralen Tendenz resp. der Prüfung des stochastischen Zusammenhangs zwischen der unabhängigen Variablen (Faktor A) und der zentralen Tendenz in der abhängigen Variablen X .

Benötigte Vorkenntnisse

- Grundprinzip entscheidungsstatistischer Verfahren
- Zielsetzungen und Grundprinzip varianzanalytischer Verfahren

Geschätzte Bearbeitungszeit

Nach Konsultation der in Ihrem Curriculum vorgesehenen Vorbereitungslektüre können Sie diesen Lernschritt in ca. 30 Minuten bearbeiten.

Den Studierenden im Grundkurs zur Statistik wird empfohlen, den Lernschritt nach der vorgegebenen Struktur zu bearbeiten.

Einführung

Als erstes wollen wir uns ganz kurz das Grundprinzip parametrischer Varianzanalysen und die damit verbundenen Voraussetzungen in Erinnerung rufen, welche die auszuwertenden Daten zu erfüllen haben. Vor diesem Hintergrund wird dann leicht einsehbar, in welchen Situationen nichtparametrische Varianzanalysen, sog. Rangvarianzanalysen, angezeigt sind.

Wie Sie sich erinnern mögen, wurde die Varianzanalyse von R. A. Fisher in der Absicht erfunden, mehrere Stichproben bezüglich des Mittelwertes eines intervall oder proportional skalierten Merkmals simultan zu vergleichen.

Dies ist, so konnte Fisher zeigen, anhand des Vergleichs von zwei unterschiedlichen Schätzungen der Varianz des interessierenden Merkmals in der Population möglich (womit sich der etwas eigenartige Name "Varianzanalyse" erklärt). Die eine Schätzung

$$\hat{\sigma}_{in}^2$$

stützt sich ausschliesslich auf die Varianz der Daten innerhalb der Stichproben und abstrahiert damit vollständig von den unterschiedlichen Stichprobenmittelwerten, die andere Schätzung

$$\hat{\sigma}_{zw}^2$$

hingegen basiert nur auf der Unterschiedlichkeit der Stichprobenmittelwerte, konkret auf der Varianz der Stichprobenmittelwerte bezüglich des Gesamtmittelwertes der Daten.

Anhand eines Vergleichs dieser beiden Schätzwerte für die Populationsvarianz

$$\hat{\sigma}^2$$

sind dann die folgenden Schlüsse möglich:

- Wenn sich

$$\hat{\sigma}_{in}^2$$

und

$$\hat{\sigma}_{zw}^2$$

bei einer Irrtumswahrscheinlichkeit

$$\leq$$

p5% signifikant unterscheiden: Mit der ermittelten Irrtumswahrscheinlichkeit darf angenommen werden, dass sich mindestens 2 Stichprobenmittelwerte signifikant unterscheiden.

- Wenn sich

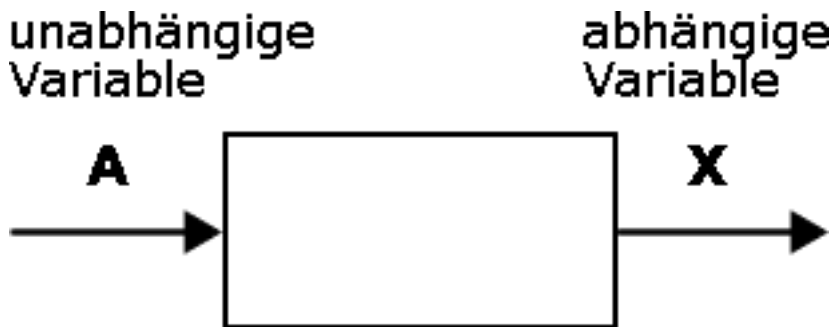
$$\hat{\sigma}_{in}^2$$

und

$$\hat{\sigma}_{zw}^2$$

nicht signifikant unterscheiden ($p > 5\%$): Kein Paarvergleich von Stichprobenmittelwerten wird einen signifikanten Unterschied aufdecken.

Unter dem Gesichtspunkt varianzanalytischer "Wirkungsmodelle" können die beiden denkbaren Resultate einer Varianzanalyse auch anders formuliert werden. Im einfachsten Fall einer univariaten, einfaktoriellen Varianzanalyse bezeichnen wir dabei das die Stichproben definierende Merkmal als unabhängige Variable und das in den Stichproben erhobene Merkmal als abhängige Variable.



Überlegen Sie, welche Aussagen auf Grund des Resultates einer Varianzanalyse unter diesem Gesichtspunkt denkbar sind, und blättern Sie erst dann weiter.

Unter dem Blickwinkel varianzanalytischer "Wirkungsmodelle" sind auf Grund des Resultates einer univariaten, einfaktoriellen Varianzanalyse folgende Schlüsse möglich:

- Wenn sich

$$\hat{\sigma}_{in}^2$$

und

$$\hat{\sigma}_{zw}^2$$

mit einer Irrtumswahrscheinlichkeit p

$$\leq$$

5% signifikant unterschieden: Mit der ermittelten Irrtumswahrscheinlichkeit darf angenommen werden, dass die unabhängige Variable A mit der abhängigen Variablen X in einem stochastischen Zusammenhang steht.

- Wenn sich

$$\hat{\sigma}_{in}^2$$

und

$$\hat{\sigma}_{zw}^2$$

nicht signifikant unterscheiden ($p > 5\%$): Die unabhängige Variable A steht mit der abhängigen Variablen X in keinem statistisch nachweisbaren stochastischen Zusammenhang.

Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen

Eine Varianzanalyse ermöglicht uns damit sowohl eine Aussage bezüglich eines simultanen Vergleichs der Stichprobenmittelwerte als auch eine Aussage zum stochastischen Zusammenhang zwischen der unabhängigen und der abhängigen Variablen.

Varianzanalysen sind mächtige Verfahren, leider stellen die bisher angesprochenen parametrischen Varianzanalysen relativ strenge Anforderungen an die auszuwertenden Daten.

Welche Voraussetzungen müssen im einfachsten Fall, z.B. für eine univariate, einfaktorielle Varianzanalyse ohne Messwiederholungen, erfüllt sein?

- ☐ Die unabhängige Variable muss nominal skaliert sein
- ☐ Die abhängige Variable muss intervall oder proportional skaliert sein
- ☐ Die Stichprobenvarianzen müssen homogen sein
- ☐ Die Stichprobendaten müssen normalverteilt sein
- ☐ Die Daten müssen in der Population normalverteilt sein

Wann sind nichtparametrische (verteilungsfreie) Varianzanalysen angezeigt?

Sind eine oder mehrere Voraussetzungen verletzt, die eine parametrische Varianzanalyse an die Daten stellt, so ist eine parametrische Datenanalyse nicht möglich, wir wählen ein nichtparametrisches (verteilungsfreies) Verfahren, eine sog. Rangvarianzanalyse.

Ganz generell sind drei Fälle denkbar:

1. Wir wählen eine Rangvarianzanalyse, wenn die abhängige Variable ordinal skaliert ist.
2. Wir wählen eine Rangvarianzanalyse, wenn das erhobene Merkmal wohl intervall- oder proportional skaliert ist, für das fokussierte Konstrukt aber unklar ist, inwieweit es den strengen Voraussetzungen einer Intervallskala zu genügen vermag.

Beispiel

Über ein intervall- oder proportional skaliertes Indikatormerkmal wird ein ordinal skaliertes Merkmal eingeschätzt.

Denken Sie z.B. an einen Fragebogen zum Thema "Überdauernde Besorgnis".

Das Konstrukt "Überdauernde Besorgnis" wird im Rahmen eines solchen Fragebogens über die Zahl der vom Probanden positiv beantworteten Fragen eingeschätzt. Der Ausprägungsgrad des konkret erhobenen Merkmals "Anzahl positiv beantworteter Fragen" wird mit ganzen Zahlen ausgedrückt, dieses Merkmal ist proportional skaliert.

Leider ist nun aber nicht anzunehmen, dass auch das eigentlich interessierende Konstrukt "überdauernde Besorgnis" dieser strengen Metrik folgt, d.h. einem Unterschied von einem Punkt im Merkmal "Anzahl positiv beantworteter Fragen" entspricht nicht immer derselbe Unterschied im fokussierten Konstrukt "Überdauernde Besorgnis". Wollen wir Aussagen machen zum Konstrukt "Überdauernde Besorgnis", so müssen wir von ordinal skalierten Daten ausgehen, obwohl das konkret erhobene Merkmal proportional skaliert ist. Konsequenterweise wählen wir auch in diesem Fall ein nichtparametrisches Verfahren.

3. Wir wählen eine Rangvarianzanalyse, wenn eine oder mehrere Verteilungs-Voraussetzungen verletzt sind. Zeigt sich im Rahmen eines varianzanalytischen Untersuchungsdesigns, dass die abhängige Variable in der Population nicht normalverteilt ist oder sind die Stichprobenvarianzen inhomogen, so ist nur noch eine Rangvarianzanalyse möglich.

Beispiel

Verletzung einer Verteilungsvoraussetzung für eine parametrische Varianzanalyse

Im Rahmen einer einfachsten Untersuchung sollen Probandinnen und Probanden unterschiedlichen Alters bezüglich ihrer Reaktionszeitgeschwindigkeit untersucht werden. Wir denken an ein varianzanalytisches Design und wählen als unabhängige Variable das Merkmal "Lebensalter", als abhängige Variable das proportionalskalierte Merkmal "mittlere Reaktionszeit in 50 normierten Reaktionsexperimenten".

Wir untersuchen eine Stichprobe lebenswürdiger Personen, die bereit sind, an unserem Experiment teilzunehmen. Dabei registrieren wir für jede Versuchsperson das Lebensalter (unabhängige Variable) und die von ihr erzielte mittlere Reaktionszeit in 50 normierten Reaktionsversuchen (abhängige Variable). Nach Abschluss der Versuche gruppieren wir die Probandinnen und Probanden nach ihrem Lebensalter in 5 Stichproben.

Damit ist eine varianzanalytische Datenauswertung vorbereitet, und wir müssen der Frage nachgehen, ob die Voraussetzungen einer parametrischen Varianzanalyse erfüllt sind. Dabei kommen wir zu folgendem Befund:

Das interessierende Merkmal, unsere abhängige Variable "mittlere Reaktionszeit", ist proportional skaliert. Leider zeigen aber die Ergebnisse älterer Untersuchungen, die sich mit Reaktionszeiten beschäftigen, dass "Reaktionszeiten" in der Population nie normalverteilt sind. Dies leuchtet uns ein, sind doch der menschlichen "Reaktionszeit" gegen unten Grenzen gesetzt. Damit ist eine zentrale Voraussetzung parametrischer Varianzanalysen verletzt, und wir müssen - trotz einer proportional skalierten abhängigen Variablen - ein nichtparametrisches oder, wie man in solchen Fällen auch sagt, ein verteilungsfreies Verfahren wählen. Angezeigt ist im gegebenen Fall eine Rangvarianzanalyse.

Theorie (1 - 3)

Rangvarianzanalysen ermöglichen uns — analog zu den parametrischen Varianzanalysen — sowohl einen simultanen Vergleich der zentralen Tendenz des interessierenden Merkmals in mehreren Stichproben als auch die Prüfung eines stochastischen Zusammenhangs zwischen der unabhängigen und der abhängigen Variablen. Rangvarianzanalysen stellen nur minimale Anforderungen an die auszuwertenden Daten: Einzige Voraussetzung ist, dass das interessierende Merkmal (die abhängige Variable) zumindest ordinal skaliert ist. Explizit ausgedrückt heisst das, dass Rangvarianzanalysen für ordinal, intervall und proportional skalierte Merkmale möglich sind; sie stützen sich aber in jedem Fall nur auf die Ranginformation der Daten.

Eine Zwischenfrage: Warum sprechen wir im Zusammenhang mit Rangvarianzanalysen von der zentralen Tendenz in den Daten und nicht wie üblich vom Mittelwert?

- ☐ Weil die Mittelwerte noch nicht bekannt sind.
- ☐ Zentrale Tendenz und Mittelwert sind bei ordinal skalierten Daten Synonyme.
- ☐ Weil Mittelwerte für ordinal skalierte Daten nicht definiert sind.
- ☐ Weil nichtparametrische Verfahren grundsätzlich eine andere Nomenklatur verwenden.

Als erstes wollen wir den Überlegungen von Kruskal und Wallis folgen, die sich bei der Entwicklung des H-Tests am Grundprinzip orientierten, nach dem alle prüfstatistischen Verfahren aufgebaut sind.

Lassen Sie uns zusammentragen, was aus formaler Sicht zu jedem prüfstatistischen Verfahren gehört. Aus rein formaler Sicht, d.h. ohne Berücksichtigung der Interpretation der Resultate, brauchen wir:

- ☐ eine Prüfgrösse
- ☐ Kenntnis, was eine Ablehnung von H_0 inhaltlich bedeutet
- ☐ eine Arbeitshypothese H_0 und eine Alternativhypothese H_1
- ☐ Kenntnis, wie die Prüfgrösse bei Gültigkeit von H_0 verteilt ist
- ☐ Kenntnis, wie die Prüfgrösse verteilt ist, wenn H_0 falsch ist
- ☐ ein Signifikanzniveau

Auch der von Kruskal und Wallis entwickelte H-Test basiert auf

1. einer Prüfgrösse,
2. einer Arbeitshypothese H_0 und eine Alternativhypothese H_1 ,
3. einer Prüfverteilung, d.h. einer Verteilung, die beschreibt, wie die Prüfgrösse bei Gültigkeit der Arbeitshypothese H_0 verteilt ist,
4. einem Signifikanzniveau, das die Beurteilung eines konkret vorliegenden Ausprägungsgrades der Prüfgrösse im Rahmen der Prüfverteilung erlaubt.

Die grundlegende Idee von Kruskal und Wallis, auf der die von ihnen entwickelte Prüfgrösse aufbaut, ist sehr einfach zu verstehen. Wir wollen sie graphisch veranschaulichen.

Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [\[link\]](#)

(0)		Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
(1)	Als erstes bilden wir eine Gesamtstichprobe, d.h. wir führen die Daten aus allen Stichproben zusammen.	Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
(2)	Dann rangieren wir die Daten in dieser Gesamtstichprobe. Bei den 7 unterschiedlichen Ausprägungen vergeben wir die Ränge 1 bis 7.	Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
(3)	Wie Sie erkennen können, werden die Daten aus den einzelnen Stichproben durchmischt. Für die Gesamtstichprobe ermitteln wir den mittleren Rang M , indem wir die Summe aller vergebenen Ränge durch die Zahl der Daten dividieren.	Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
(4)	Nun ordnen wir die Daten wieder in die Stichproben ein, aus denen sie ursprünglich stammten und geben ihnen die Rangplätze mit, die sie in der Gesamtstichprobe erzielten.	Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
http://mesosworld.ch - Stand vom: 20.1.2010		Dieses Element (Animation, Video etc.) kann in der PDF version nicht dargestellt werden und ist nur in der online Version sichtbar. [link]
(5)	Wir bestimmen für jede Stichprobe den mittleren Rang M_i , indem wir	

Und nun zum wahrscheinlichkeitstheoretischen Kern der Idee von Kruskal und Wallis, auf den Sie leicht selber kommen könnten.

Stammen die Stichproben bezüglich ihrer zentralen Tendenz aus derselben Population, unterscheiden sie sich also nur zufällig, so müssen sich die in der Gesamtstichprobe erzielten Ränge auch zufällig auf die einzelnen Stichproben verteilen. Oder anders ausgedrückt: In allen Stichproben müssen in vergleichbarer Weise kleine und grosse Rangplätze zu finden sein.

Dies führt zu mittleren Rängen M_1 , M_2 und M_3 in den Stichproben, die

- ☐ mit dem mittleren Rang M der Gesamtstichprobe identisch sind.
- ☐ sich nur zufällig vom mittleren Rang M der Gesamtstichprobe unterscheiden.
- ☐ sich systematisch vom mittleren Rang M der Gesamtstichprobe unterscheiden.
- ☐ absolut identisch sind, aber vom mittleren Rang M der Gesamtstichprobe abweichen.

Wenn die Stichproben bezüglich der zentralen Tendenz aus derselben Population stammen, so dürfen die mittleren Ränge $M_1, M_2, \dots, M_p$ der Stichproben nur zufällig vom mittleren Rang M der Gesamtstichprobe abweichen.

Auf der Grundlage dieser Überlegungen haben Kruskal und Wallis eine Prüfgrösse H hergeleitet und bewiesen, dass diese Prüfgrösse H -verteilt ist, wenn die Arbeitshypothese H_0 gültig ist. Dass die Definitionsformel dieser Prüfgrösse nicht ganz einfach ist, insbesondere dann, wenn verbundene Ränge vorkommen, haben Sie der Vorbereitungs-Literatur zu diesem Lernschritt entnehmen können.

Da wir für den H -Test in der Regel ein Statistikprogramm einsetzen, verzichten wir an dieser Stelle auf die Wiedergabe der Berechnungsformel.

Nun noch ein Wort zur Prüfverteilung:

Wir wollen p Stichproben mit dem H -Test simultan vergleichen. Sind in mindestens $p = 4$ Stichproben je mindestens 5 Daten erhoben worden, so geht die H -Verteilung über in eine

χ^2

-Verteilung mit dem Freiheitsgrad $df = p-1$. In allen anderen Fällen müssen wir entweder die spezielle H -Verteilung benutzen, die in Lehrbüchern (z.B. (Bortz et al. 1998)) oder Tabellenwerken nachgeschlagen werden kann, oder wir prüfen, ob die Verteilungen des Merkmals in den hinter den Stichproben stehenden Populationen zumindest ähnlich sind. Auch in diesem, meist nicht leicht zu prüfenden Fall darf anstelle der H -Verteilung die

χ^2

-Verteilung ($df = p-1$) als Prüfverteilung gewählt werden.

Wir machen uns die Sache leicht, wenn wir den H -Test, d.h. Rangvarianzanalysen ohne Messwiederholungen, von SPSS oder einem anderen Statistikprogramm rechnen lassen. SPSS wählt selbständig die adäquate Prüfverteilung und gibt Ihnen direkt die Überschreitungswahrscheinlichkeit der Prüfgrösse aus.

Bevor wir uns dies an einem Beispiel veranschaulichen, stellen wir abschliessend die Elemente des H -Tests nochmals zusammen:

Arbeitshypothese H_0 und Alternativhypothese H_1 :	H_0: Die Stichproben stammen aus Populationen, die bezüglich der zentralen Tendenz im interessierenden Merkmal identisch sind. Anders ausgedrückt: Die Stichproben unterscheiden sich bezüglich der zentralen Tendenz des Merkmals nur zufällig. H_1: Die Stichproben unterscheiden sich bezüglich der zentralen Tendenz. Prüfen wir auf einen stochastischen Zusammenhang zwischen der unabhängigen Variablen (d.h. dem Definitionsmerkmal der Stichproben) einerseits und der abhängigen Variablen (dem in den Stichproben erhobenen Merkmal) andererseits, so können die Hypothesen auch wie folgt formuliert werden: H_0: Zwischen der abhängigen und der unabhängigen Variablen besteht kein statistisch nachweisbarer stochastischer Zusammenhang. H_1: Zwischen der abhängigen und der unabhängigen Variablen besteht ein stochastischer Zusammenhang.
Prüfgrösse:	Die Prüfgrösse H bestimmt sich nach der von Kruskal und Wallis hergeleiteten Formel (in diesem Lernschritt nicht wiedergegeben).
Prüfverteilung:	Prüfverteilung ist die H-Verteilung. Die H-Verteilung geht über in eine χ^2-Verteilung mit dem Freiheitsgrad $df = p-1$ (p steht für die Zahl der Stichproben), wenn entweder $p \geq 4$ und in jeder Stichprobe mindestens 5 Daten erhoben wurden, oder wenn die Verteilungen des Merkmals in den hinter den Stichproben stehenden Populationen ähnlich sind.
Signifikanzniveau:	Wir wählen das Signifikanzniveau, d.h. das α-Fehler-Risiko, den Gepflogenheiten entsprechend mit 5%, 1% oder 0,1%.

Fallbeispiel

Im Rahmen eines einfachsten Experimentes soll untersucht werden, inwieweit die "Reaktionszeiten" von Versuchspersonen in genormten Reaktionsexperimenten mit dem "Lebensalter" der Probandinnen und Probanden in einem stochastischen Zusammenhang stehen. Zur Prüfung dieser Frage wählen wir ein varianzanalytisches Design mit dem Merkmal "Lebensalter" als unabhängige und dem proportional skalierten Merkmal "mittlere Reaktionszeit in 50 genormten Reaktionsexperimenten" als abhängige Variable.

Wir untersuchen eine grössere Zahl von Personen, die bereit sind, an unserem Experiment teilzunehmen. Dabei registrieren wir für jede Versuchsperson das Lebensalter und die von ihr erzielte mittlere Reaktionszeit in 50 normierten Reaktionsversuchen.

Nach Abschluss der Versuche gruppieren wir die Probandinnen und Probanden nach ihrem Lebensalter in 5 Stichproben:

- Stichprobe a1: 21-30-Jährige
- Stichprobe a2: 31-40-Jährige
- Stichprobe a3: 41-50-Jährige
- Stichprobe a4: 51-60-Jährige
- Stichprobe a5: 61-70-Jährige

Damit ist eine varianzanalytische Datenauswertung vorbereitet; wir wollen sie in den üblichen Schritten durcharbeiten.

1. Zusammenstellung und Überprüfung der von den Daten erfüllten Voraussetzungen:

Das interessierende Merkmal, unsere abhängige Variable "mittlere Reaktionszeit", ist proportional skaliert. Aus zahlreichen älteren Untersuchungen wissen wir aber, dass "Reaktionszeiten" nie normalverteilt sind, weil menschlichen "Reaktionszeiten" gegen unten Grenzen gesetzt sind. Damit ist eine zentrale Voraussetzung parametrischer Varianzanalysen verletzt, und wir müssen - trotz einer proportional skalierten abhängigen Variablen - ein nichtparametrisches (verteilungsfreies) Verfahren wählen. Angezeigt ist also eine Rangvarianzanalyse für unabhängige Stichproben, die bezüglich der Verteilung des Merkmals in der Population nichts voraussetzt. Die Rangvarianzanalyse wird aber nur die Ranginformation der Daten auswerten, dies obwohl die Daten auf einer Proportionalskala erhoben wurden.

2. Formulierung der Arbeitshypothese H0 und Alternativhypothese H1:

Da wir uns für den stochastischen Zusammenhang zwischen dem "Lebensalter" (unabhängige Variable) und der "mittleren Reaktionszeit" (abhängige Variable) interessieren, formulieren wir unsere Arbeitshypothese H0 wie folgt: Zwischen der unabhängigen und der abhängigen Variablen besteht kein statistisch nachweisbarer stochastischer Zusammenhang.

Als Alternativhypothese H1 postulieren wir: Die unabhängige und die abhängige Variable stehen in einem stochastischen Zusammenhang.

3. Prüfung der Arbeitshypothese H0:

Wir werten unsere Daten anhand einer Rangvarianzanalyse für unabhängige Stichproben (H-Test nach Kruskal & Wallis) aus und geben hierfür die folgenden SPSS-Befehle ein:

1. GET FILE 'X:\\SPSS_DAT\\reaktion.sav'.
LIST alter reakt.

Wir lesen unser Datenfile mit dem Namen **reaktion.sav** und lassen uns eine Liste der Rohdaten ausgeben. Die Variable **alter** beschreibt das Lebensalter, die Variable **reakt** die mittlere Reaktionszeit einer Person.

Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen

SPSS Output 1 zeigen

ALTER	REAKT
21	166
22	146
31	134
24	201
25	205
26	198
27	165
22	199
23	188
23	176
33	199
34	179
35	199
31	201
39	225
34	265
.	.
.	.
58	289
56	299
61	266
67	309
66	267
65	276
68	299
66	311
65	300
67	276
61	188
60	206
45	207

Number of cases read: 58 Number of cases listed: 58

2. RECODE alter (21 THRU 30 = 1) (31 THRU 40 = 2)
(41 THRU 50 = 3) (51 THRU 60 = 4)
(61 THRU 70 = 5) INTO altkat.
List alter **altkat** reakt.

Wir gruppieren die Personendaten in fünf Alterskategorien, die über die Hilfsvariable altkat gespeichert werden und lassen uns nochmals eine Kontrollliste ausgeben.

Einfaktorielle Rangvarianzanalyse ohne Messwiederholungen

SPSS Output 2 zeigen

ALTER	ALTKAT	REAKT
21	1	166
22	1	146
31	2	134
24	2	134
24	1	201
25	1	205
26	1	198
27	1	165
22	1	199
23	1	188
23	1	176
33	2	199
34	2	179
35	2	199
31	2	201
39	2	225
34	2	265
.	.	.
.	.	.
58	4	289
56	4	299
61	5	266
67	5	309
66	5	267
66	5	276
68	5	299
66	5	311
65	5	300
67	5	276
61	5	188
60	4	206
45	3	207

Number of cases read: 58 Number of cases listed: 58

3. VALUE LABELS altkat 1 '21-30 Jahre' 2 '31-40 Jahre'
3 '42-50 Jahre' 4 '51-60 Jahre'
5 '61-70 Jahre.'.

VARIABLE LABELS altkat 'Alters-Kategorie'.

VARIABLE LABELS reakt 'mittlere Reaktionszeiten'.

NPAR TESTS /K-W reakt BY altkat(1,5).

Wir führen für die Ausgabe Bezeichnungen für die Variablen und die einzelnen Alterskategorien ein und lassen den H-Test nach Kruskal und Wallis rechnen.

SPSS Output 3 zeigen

Ranks

	ALTKAT Alters- Kategorie	N	Mean Rank
REAKT mittlere Reaktionszeiten	1 21-30 Jahre	9	14.72
	2 31-40 Jahre	14	23.89
	3 42-50 Jahre	18	27.17
	4 51-60 Jahre	8	40.63
	5 61-70 Jahre.	9	47.78
	Total	58	

Test Statistics(a,b)

	REAKT mittlere Reaktionszeiten
Chi-Square	22.830
df	4
Asymp. Sig.	.000
a Kruskal Wallis Test	
b Grouping Variable: ALTKAT Alters-Kategorie	

Der SPSS-Ausgabe entnehmen wir die Grössen der fünf Stichproben, die mittleren Ränge in den Stichproben, den Ausprägungsgrad der Prüfgrösse und die Überschreitungswahrscheinlichkeit dieses Ausprägungsgrades für den Fall, dass H_0 gültig ist.

Bei der ermittelten Überschreitungswahrscheinlichkeit $p < 0,1\%$ lehnen wir H_0 ab.

4. Interpretation

Dass H_0 mit einer Irrtumswahrscheinlichkeit $p < 0,1\%$ abgelehnt werden kann, bedeutet, dass mit einer Irrtumswahrscheinlichkeit kleiner $0,1\%$ angenommen werden darf, dass zwischen der unabhängigen Variablen "Lebensalter" und der abhängigen Variablen "mittlere Reaktionszeit" ein stochastischer Zusammenhang besteht. Über die Art dieses Zusammenhangs geben uns die für die Stichproben bestimmten mittleren Ränge einen groben Hinweis.

Die Reaktionszeiten

- ☐ nehmen mit zunehmendem Alter ab.
- ☐ bleiben in etwa gleich.
- ☐ nehmen mit zunehmendem Alter zu.

5. Mögliche weitere Datenanalysen:

Soll z.B. untersucht werden, ob sich die recht ähnlichen Reaktionszeiten der 51-60-Jährigen signifikant von denjenigen der 61-70-Jährigen unterscheiden, so vergleichen Sie die beiden Stichproben mit dem

- ☐ t-Test für unabhängige Stichproben.
- ☐ t-Test für abhängige Stichproben.
- ☐ U-Test nach Mann-Whitney.
- ☐ Wilcoxon-Test.
- ☐ Scheffé-Test.