

Hair | Hult | Ringle | Sarstedt | Richter | Hauff

Partial Least Squares Strukturgleichungs- modellierung (PLS-SEM)

Eine anwendungsorientierte Einführung



Vahlen

Zum Inhalt:

Die Partial Least Squares Strukturgleichungsmodellierung (PLS-SEM) hat sich in der wirtschafts- und sozialwissenschaftlichen Forschung als geeignetes Verfahren zur Schätzung von Kausalmodellen behauptet. Dank der Anwenderfreundlichkeit des Verfahrens und der vorhandenen Software ist es inzwischen auch in der Praxis etabliert.

Dieses Buch liefert eine anwendungsorientierte Einführung in die PLS-SEM. Der Fokus liegt auf den Grundlagen des Verfahrens und deren praktischer Umsetzung mit Hilfe der SmartPLS-Software. Das Konzept des Buches setzt dabei auf einfache Erläuterungen statistischer Ansätze und die anschauliche Darstellung zahlreicher Anwendungsbeispiele anhand einer einheitlichen Fallstudie. Viele Grafiken, Tabellen und Illustrationen erleichtern das Verständnis der PLS-SEM. Zudem werden dem Leser herunterladbare Datensätze, Videos, Aufgaben und weitere Fachartikel zur Vertiefung angeboten. Damit eignet sich das Buch hervorragend für Studierende, Forscher und Praktiker, die die PLS-SEM zur Gewinnung von Ergebnissen mit den eigenen Daten und Modellen nutzen möchten.

SmartPLS ist das führende Softwareprogramm zur Schätzung von PLS-basierten Strukturgleichungsmodellen. Die Erläuterungen und die im Buch vorgeschlagenen Vorgehensweisen spiegeln den aktuellen Stand der Forschung wider.

Zu den Autoren:

Joseph F. Hair, Jr. ist Professor für Marketing an der University of South Alabama und mit mehr als 50 veröffentlichten Büchern, darunter das mit über 140.000 Zitationen als weltweites Standardwerk zu bezeichnende Buch „Multivariate Data Analysis“, einer der führenden Experten auf dem Gebiet der anwendungs-orientierten Statistik.

G. Tomas M. Hult ist Professor für Marketing und International Business am Eli Broad College of Business an der Michigan State University und mit mehr als 31.000 Zitationen bei Google Scholar einer der meist zitierten Forscher in den Wirtschaftswissenschaften, der sich in seiner Forschung intensiv mit verschiedenen Verfahren der SEM auseinandersetzt.

Christian M. Ringle ist Professor für Betriebswirtschaftslehre und Leiter des Instituts für Personalwirtschaft und Arbeitsorganisation an der Technischen Universität Hamburg (und assoziierter Professor an der University of Newcastle in Australien), Mitentwickler von SmartPLS und einer der prominentesten Vertreter der PLS-SEM in der weltweiten Forschungslandschaft.

Marko Sarstedt ist Professor für Marketing an der Otto-von-Guericke Universität Magdeburg (und assoziierter Professor an der University of Newcastle in Australien), laut Handelsblatt-Ranking einer der führenden Junior-Marketingforscher und einer der prominentesten Vertreter der PLS-SEM in der weltweiten Forschungslandschaft.

Nicole F. Richter ist Professorin für International Business an der University of Southern Denmark und beschäftigt sich seit ihrer Habilitation am Institut von Prof. Ringle in ihren Publikationen kritisch mit dem Einsatz statistischer Verfahren in der internationalen Managementforschung.

Sven Hauff vertritt aktuell die Professur für Arbeit, Personal und Organisation an der Helmut-Schmidt-Universität in Hamburg und wendet seit seiner Dissertation die PLS-SEM in verschiedenen Forschungs- und Publikationsprojekten an.

Partial Least Squares Strukturgleichungs- modellierung

Eine anwendungsorientierte Einführung

von

Prof. Dr. Joseph F. Hair, Jr.

Prof. Dr. G. Tomas M. Hult

Prof. Dr. Christian M. Ringle

Prof. Dr. Marko Sarstedt

Prof. Dr. Nicole F. Richter

PD Dr. Sven Hauff

Verlag Franz Vahlen München

Vorwort

Welche Rolle spielen das Grundeinkommen, Boni und andere Faktoren, um die Zufriedenheit von Mitarbeitern positiv zu beeinflussen? Wie wirkt sich die Zufriedenheit am Arbeitsplatz auf die Loyalität und die Leistung der Mitarbeiter aus? Was sind die wichtigsten Stellschrauben, um aus einem zufriedenen Kunden auch einen loyalen Kunden zu machen? Wie sollte eine internationale Einkaufsorganisation strukturiert sein, um den Erfolg internationaler Einkaufsaktivitäten und schließlich den Unternehmenserfolg positiv zu beeinflussen?

Die Untersuchung solcher und anderer Ursache-Wirkungs-Beziehungen steht häufig im Zentrum der sozialwissenschaftlichen Forschung und ist für die Entscheidungsträger in Unternehmen von hoher Relevanz. Dabei können Forscher und Entscheidungsträger heutzutage oft auf eine Fülle bereits vorhandener Daten zurückgreifen. Zudem gibt es inzwischen zahlreiche Standardsoftwareanwendungen, die für die Auswertung von Daten genutzt werden können. Dies ist für den Forscher und Entscheidungsträger aber nicht nur ein Segen, sondern teilweise auch eine Bürde, da die richtige Interpretation der Daten eine solide Basis an analytischen Fähigkeiten voraussetzt. Diese umfassen neben dem notwendigen theoretischen Wissen, auch solide Kenntnisse in den für die Datenauswertung notwendigen statistischen Verfahren.

Ein für die Untersuchung von Ursache-Wirkungs-Beziehungen sehr geeignetes Analyseverfahren ist die Strukturgleichungsmodellierung. Mit Hilfe der Strukturgleichungsmodellierung können abstrakte Konstrukte (z. B. die Loyalität von Kunden) über mehrere beobachtbare Variablen (z. B. die Weiterempfehlungsbereitschaft oder den Wiederkauf) gemessen werden. Gleichzeitig können Beziehungen zwischen Konstrukten auf verschiedenen Ebenen überprüft werden (z. B. zwischen Kundenzufriedenheit, Image und Kundenloyalität).

Es lassen sich im Wesentlichen zwei Verfahren der Strukturgleichungsmodellierung unterscheiden: die kovarianzbasierte Strukturgleichungsmodellierung und die Partial Least Squares Strukturgleichungsmodellierung. Die kovarianzbasierte Strukturgleichungsmodellierung war über viele Jahre das etablierte und damit dominante Verfahren. In den letzten Jahren ist aber der Einsatz der Partial Least Squares Strukturgleichungsmodellierung immer beliebter geworden, was zahlreiche Studien zu dem Einsatz des Verfahrens in verschiedenen Disziplinen zeigen. Dies liegt vor allem daran, dass die Partial Least Squares Strukturgleichungsmodellierung in vielen Forschungs- und Entscheidungssituationen vorteilhaft ist, z. B. wenn es um die Abbildung komplexer Entscheidungsprobleme in statistischen Modellen geht. Insofern sind Forscher und Entscheidungsträger gut beraten, sich mit der Strukturgleichungsmodellierung und insbesondere mit den Grundlagen der Partial Least Squares Strukturgleichungsmodellierung auseinanderzusetzen.

Hierzu bietet dieses Buch die ideale Grundlage: Es liefert eine anwendungsorientierte Einführung in die Partial Least Squares Strukturgleichungsmodel-

lierung, indem statistische Ansätze so einfach wie möglich erläutert werden. Wo immer möglich verzichten wir auf Gleichungen, Formeln oder griechische Symbole und konzentrieren uns auf die anschauliche Darstellung des Verfahrens. Wir liefern Entscheidungsbäume und Faustregeln für den richtigen Einsatz des Verfahrens. Zudem kommt eine Fallstudie zum Einsatz, die in jedem Kapitel die Anwendung des Verfahrens auch aus praktischer Sicht verdeutlicht. Die Erläuterungen basieren auf den neuesten Erkenntnissen der Forschung. So greift das Buch aktuelle Diskussionen rund um die Partial Least Squares Strukturgleichungsmodellierung auf und stellt neueste Weiterentwicklungen, beispielsweise in den Bereichen Modellevaluation und Mediationsanalysen, vor. Der Leser hat damit den gesamten Werkzeugkoffer an State-of-the-Art Instrumenten zur Lösung des eigenen Forschungs- oder Entscheidungsproblems an der Hand.

Diesem Buch ist eine inzwischen in zweiter Auflage erschienene englischsprachige Fassung vorausgegangen. Diese ist die Grundlage für diese deutsche Fassung des Buches. Der Erfolg des englischsprachigen Titels und die vielen positiven Rückmeldungen haben uns in der Ausrichtung des Buches bekräftigt. Aus vielen Vorlesungen wissen wir, dass obwohl Englisch sich als Wissenschaftssprache inzwischen durchgesetzt haben mag, die Erfassung insbesondere komplexerer Thematiken in der eigenen Muttersprache leichter fällt. Wir sind daher davon überzeugt, dass diese deutsche Fassung den Lernerfolg deutschsprachiger Anwender verbessert und das Verständnis des Verfahrens und die Umsetzung deutlich vereinfacht. Das Glossar referenziert schließlich auf die englischsprachigen Begriffe, so dass der Leser bei Bedarf einen guten Einstieg auch in die englischsprachige Forschung zu diesem Thema findet.

Dieses Buch arbeitet mit einem Datensatz zur Unternehmensreputation (der auch die Grundlage der Anwendungsfälle in der englischsprachigen Fassung ist). Dieser Datensatz steht zusammen mit weiteren nutzbaren Modellen auf der Website des Verlages Vahlen zum freien Download zur Verfügung. Wir arbeiten mit der deutschen Sprachversion des inzwischen als Standard etablierten Softwareprogramms SmartPLS 3, das nicht nur durch sein umfassendes Funktionsangebot sondern auch durch seine Anwenderfreundlichkeit überzeugt. Schließlich finden Sie zu jedem Kapitel dieses Buches ein deutschsprachiges Video auf der Website des Verlages, welches anhand des Anwendungsbeispiels die Umsetzung der statistischen Grundladien sowie die Interpretation der Ergebnisse medial unterstützt. Am Ende dieses Buches finden Sie einen technischen Anhang, der die Optionen zum Herunterladen der Daten und der Software sowie zum Aufruf der Videos genauer beschreibt.

Wir danken Hermann Schenk und dem Team des Verlages Vahlen für die wunderbare Zusammenarbeit bei der Erstellung dieses Buches. Zudem gilt unser Dank Frau Alexandra Sarstedt für die exzellenten sprachlichen Korrekturen. Weiterhin möchten wir uns bei unseren deutschsprachigen Kollegen für ihre konstruktiven Hinweise der letzten Jahre zu unseren Arbeiten zur Partial Least Squares Strukturgleichungsmodellierung bedanken, die in dieses Buch eingeflossen sind:

Sönke Albers (Kühne Logistics University), Dorothea Alewell (Universität Hamburg), Dennis Ahrholdt (HSBA), Ingo Balderjahn (Universität Potsdam), Jan-Michael Becker (Universität zu Köln), Silke Boenigk (Universität Hamburg), Michel Clement (Universität Hamburg), Adamantios Diamantopoulos (Universität Wien), Andreas Eggert (Universität Paderborn), Margit Enke (TU Bergakademie Freiberg), Bernd Erichson (Otto-von-Guericke-Universität Magdeburg), Georg Fassott (TU Kaiserslautern), Martin Fritze (Universität Rostock), Oliver Götz (ESB Business School), Siggj Gudergan (University of Newcastle, Australia), Michael Haenlein (ESCP Europe), Karl-Werner Hansmann (Universität Hamburg), Jörg Henseler (Universität Twente), Claudia Höck (Universität Hamburg), Michael Höck (TU Bergakademie Freiberg), Roland Holten (Universität Frankfurt), Diana Ingenhoff (Université de Fribourg), Martin Klarmann (Karlsruher Institut für Technologie), Marcel Lichters (Otto-von-Guericke-Universität Magdeburg), Matthias Meyer (TU Hamburg), Christian Nitzl (Universität der Bundeswehr München), Sascha Raithel (Freie Universität Berlin), Ellen Roemer (Hochschule Ruhr West), Alexander Rossmann (Hochschule Reutlingen), Henrik Sattler (Universität Hamburg), Holger Schiele (Universität Twente), Rainer Schlittgen (Universität Hamburg), Jan Hendrik Schreier (Otto-Friedrich-Universität Bamberg), Florian Schuberth (Universität Würzburg), Tobias Schütz (ESB Business School), Manfred Schwaiger (Ludwig-Maximilians-Universität München), Rudolf Sinkovics (University of Manchester), Bernhard Swoboda (Universität Trier), Franziska Völckner (Universität zu Köln), Bodo Vogt (Otto-von-Guericke-Universität Magdeburg), Ralf Wagner (Universität Kassel), Sven Wende (SmartPLS GmbH) und Martin Wetzels (Universität Maastricht).

Wir sind davon überzeugt, dass dieses Buch sich hervorragend für Studierende, Forscher und Praktiker eignet, die die Partial Least Squares Strukturgleichungsmodellierung zur Gewinnung von Ergebnissen mit den eigenen Daten und Modellen nutzen möchten. Wir wünschen Ihnen viel Freude bei der Umsetzung der eigenen Projekte!

Natürlich freuen wir uns stets über Hinweise, Kritik und Verbesserungsvorschläge zum Buch. Diese können gerne über den Verlag oder direkt an uns gerichtet werden.

Für die deutsche Version im Februar 2017

Sven Hauff
Nicole Richter
Christian M. Ringle
Marko Sarstedt

Inhaltsverzeichnis

Vorwort	V
Kapitel 1: Einführung in die Strukturgleichungsmodellierung	1
Kapitelüberblick	2
Was ist Strukturgleichungsmodellierung?	2
Grundlegendes zur Verwendung von Strukturgleichungsmodellen ..	4
Composite-Variablen	5
Messung	5
Skalenniveau	7
Kodierung	8
Verteilung der Daten	9
Strukturgleichungsmodellierung mit Partial Least Squares	
Pfadmodellen	10
Pfadmodelle mit latenten Variablen	10
Messtheorie	12
Strukturtheorie	12
PLS-SEM, CB-SEM und Regressionen auf Basis von Summenwerten ..	13
Dateneigenschaften	20
Stichprobengröße	20
Fehlende Werte	23
Verteilung	23
Skalenniveau	23
Modelleigenschaften	24
Organisation der folgenden Kapitel	26
Zusammenfassung	27
Wiederholungsfragen	29
Weiterführende Fragen	29
Empfohlene Literatur	29
Kapitel 2: Spezifikation des Pfadmodells und Prüfung der Daten ..	31
Kapitelüberblick	32
Schritt 1: Spezifikation des Strukturmodells	32
Mediation	34
Moderation	36
Modelle höherer Ordnung und hierarchische Komponenten-	
modelle	37
Schritt 2: Spezifikation der Messmodelle	38
Reflektiv und formativ spezifizierte Messmodelle	40
Single-Item-Messungen und Summenwerte	45
Schritt 3: Erhebung und Prüfung der Daten	48
Fehlende Werte	48

Antwortmuster	50
Inkonsistente Antworten	50
Ausreißer	51
Verteilung der Daten	52
Anwendungsbeispiel: Spezifikation des PLS-Pfadmodells	53
Schritt 1: Spezifikation des Strukturmodells	53
Schritt 2: Spezifikation der Messmodelle	55
Schritt 3: Erhebung und Prüfung der Daten	57
Erstellung eines Pfadmodells mit der Software SmartPLS	59
Zusammenfassung	65
Wiederholungsfragen	67
Weiterführende Fragen	67
Empfohlene Literatur	67
Kapitel 3: Schätzung des PLS-Pfadmodells	69
Kapitelüberblick	70
Schritt 4: Modellschätzung und der PLS-SEM-Algorithmus	70
Funktionsweise des Algorithmus	70
Statistische Eigenschaften	74
Einstellungen zur Ausführung des Algorithmus	76
Ergebnisse	78
Anwendungsbeispiel: PLS-Pfadmodellschätzung	79
Modellschätzung	79
Ergebnisse der Modellschätzung	82
Zusammenfassung	85
Wiederholungsfragen	86
Weiterführende Fragen	87
Empfohlene Literatur	87
Kapitel 4: Gütebeurteilung von PLS-SEM-Ergebnissen (Teil I)	89
Kapitelüberblick	90
Schritt 5: Evaluation der Messmodelle	90
Schritt 5a: Evaluation reflektiv spezifizierter Messmodelle	96
Interne-Konsistenz-Reliabilität	96
Konvergenzvalidität	97
Diskriminanzvalidität	99
Anwendungsbeispiel: Evaluation reflektiv spezifizierter Messmodelle	106
Ausführen des PLS-SEM-Algorithmus	106
Evaluation der reflektiv-spezifizierten Messmodelle	107
Zusammenfassung	115
Wiederholungsfragen	116
Weiterführende Fragen	116
Empfohlene Literatur	116
Kapitel 5: Gütebeurteilung von PLS-SEM-Ergebnissen (Teil II)	119
Kapitelüberblick	120
Schritt 5b: Evaluation formativ spezifizierter Messmodelle	120

Schritt 1: Prüfung der Konvergenzvalidität	122
Schritt 2: Prüfung der Kollinearität der formativ spezifizierten Messmodelle	123
Schritt 3: Prüfung der Signifikanz und Relevanz der formativen Indikatoren	127
Auswirkungen der Anzahl verwendeter Indikatoren auf die Indikatorgewichte	128
Behandlung von nicht signifikanten Indikatorgewichten	129
Bootstrapping-Verfahren	131
Konzept	131
Bootstrap-Konfidenzintervalle	135
Anwendungsbeispiel: Evaluation formativ spezifizierter Messmodelle	139
Erweiterung des einfachen Pfadmodells	139
Evaluation der reflektiv spezifizierten Messmodelle	147
Evaluation der formativ spezifizierten Messmodelle	149
Zusammenfassung	159
Wiederholungsfragen	160
Weiterführende Fragen	161
Empfohlene Literatur	161
Kapitel 6: Gütebeurteilung von PLS-SEM-Ergebnissen (Teil III)	163
Kapitelüberblick	164
Schritt 6: Evaluation der Ergebnisse des Strukturmodells	164
Schritt 1: Prüfung der Kollinearität	167
Schritt 2: Prüfung der Pfadkoeffizienten im Strukturmodell	168
Schritt 3: Prüfung des Bestimmtheitsmaßes (R^2 -Wert)	170
Schritt 4: Prüfung der f^2 -Effektstärken	173
Schritt 5: Blindfolding und Prüfung der Prognoserelevanz (Q^2 - Wert)	174
Schritt 6: Prüfung der q^2 -Effektstärken	177
Anwendungsbeispiel: Evaluation des Strukturmodells und Ergebnisauswertung	180
Zusammenfassung	189
Wiederholungsfragen	190
Weiterführende Fragen	191
Empfohlene Literatur	191
Kapitel 7: Mediator- und Moderatoranalysen	193
Kapitelüberblick	194
Mediation	195
Einführung	195
Arten von Mediatoreffekten	198
Prüfung mediiender Effekte	200
Messmodellevaluation in Mediatoranalysen	201
Multiple Mediation	201
Anwendungsbeispiel	203

Moderation	206
Einführung	206
Arten von Moderatorvariablen	208
Modellierung von Moderatoreffekten	209
Erstellung eines Interaktionsterms	211
Produktindikatoransatz	211
Orthogonalisierungsansatz	212
Zwei-Stufen-Ansatz	214
Richtlinien zur Erstellung von Interaktionstermen	215
Modellevaluation	216
Ergebnisinterpretation	217
Moderierte Mediation und mediierte Moderation	220
Anwendungsbeispiel	223
Zusammenfassung	230
Wiederholungsfragen	231
Weiterführende Fragen	231
Empfohlene Literatur	231
Kapitel 8: Ausblick auf weiterführende Verfahren	233
Kapitelüberblick	234
Importance-Performance-Analyse	235
Hierarchische Komponentenmodelle	238
Konfirmatorische Tetrad Analyse	242
Umgang mit beobachteter und unbeobachteter Heterogenität	246
Multigruppenanalyse	247
Ermittlung unbeobachteter Heterogenität	250
Messmodellinvarianz	253
Konsistentes PLS-Verfahren	255
Zusammenfassung	259
Wiederholungsfragen	262
Weiterführende Fragen	262
Empfohlene Literatur	262
Anhang	265
Glossar	269
Literatur	293
Stichwortverzeichnis	309

Kapitel 1

Einführung in die Strukturgleichungsmodellierung

Lernziele

1. Die Leser können die Bedeutung der Strukturgleichungsmodellierung (Structural Equation Modeling, SEM) und ihre Beziehung zur multivariaten Datenanalyse diskutieren.
2. Die Leser können prinzipielle Zusammenhänge zur Anwendung multivariater Analyseverfahren beschreiben.
3. Die Leser können grundlegende Konzepte der Partial Least Squares Strukturgleichungsmodellierung (PLS-SEM) erläutern.
4. Die Leser können die Unterschiede zwischen kovarianzbasierter Strukturgleichungsmodellierung (Covariance-Based-SEM, CB-SEM) und PLS-SEM diskutieren und einschätzen, wann welches Verfahren eingesetzt werden sollte.

Kapitelüberblick

Sozialwissenschaftler setzen seit vielen Jahren statistische Analyseverfahren ein, um Datenstrukturen zu erforschen, Ergebnisse zu generieren und Theorien zu testen. Dabei dominierten die Analyseverfahren der ersten Generation (z. B. die Faktor- und die Regressionsanalyse) die Forschungslandschaft bis in die 1980er Jahre. Seit den frühen 1990er Jahren finden zunehmend Verfahren der zweiten Generation (z. B. die Varianz- und die Kovarianzbasierte Strukturgleichungsmodellierung) Anwendung, welche in einigen Disziplinen mittlerweile fast 50 % der empirischen Analyseverfahren ausmachen. In diesem Kapitel erläutern wir die Grundlagen dieser zweiten Generation statistischer Analyseverfahren. Hiermit legen wir zugleich die Basis, um eines der zentralen Verfahren der zweiten Generation, die Partial Least Squares Strukturgleichungsmodellierung (PLS-SEM), besser zu verstehen und anwenden zu können.

Was ist Strukturgleichungsmodellierung?

Seit mehr als einem Jahrhundert stellen statistische Analysen ein zentrales Werkzeug von Sozialwissenschaftlern dar. Durch das Aufkommen geeigneter Hard- und Software hat sich die Anwendung statistischer Analyseverfahren rasant verbreitet. Insbesondere die Entwicklung von benutzerfreundlichen Anwendungsoberflächen, die kein umfassendes technisches Vorwissen erfordern, hat in den letzten Jahren den Zugang zu vielfältigen Verfahren ermöglicht. Wissenschaftler haben dabei lange Zeit auf uni- und bivariate Analyseverfahren gebaut, um ihre Daten und die darin enthaltenen Beziehungen zu analysieren. Die aktuelle Forschung in den Sozialwissenschaften beinhaltet aber immer komplexere Beziehungen und Modelle, für die anspruchsvollere multivariate Analyseverfahren nötig sind.

Multivariate Analysen erfordern die Anwendung statistischer Verfahren, mit deren Hilfe mehrere Variablen gleichzeitig analysiert werden können. Diese Variablen beinhalten typischerweise Informationen in Bezug auf Individuen, Unternehmen, Ereignisse, Aktivitäten, Situationen usw. Die Informationen werden dabei direkt durch Umfragen oder Beobachtungen gewonnen oder beruhen auf Sekundärdaten. Abbildung 1.1 verdeutlicht die zentralen Verfahren der multivariaten Datenanalyse.

	Vorwiegend explorativ	Vorwiegend konfirmatorisch
Verfahren der ersten Generation	Clusteranalyse Explorative Faktorenanalyse Multidimensionale Skalierung	Varianzanalyse Logistische Regression Multiple Regression Konfirmatorische Faktorenanalyse
Verfahren der zweiten Generation	Partial Least Squares Strukturgleichungsmodellierung (PLS-SEM)	Kovarianzbasierte Strukturgleichungsmodellierung (CB-SEM)

Abbildung 1.1 Kategorisierung multivariater Analyseverfahren

Die in der oberen Hälfte von Abbildung 1.1 genannten **Verfahren der ersten Generation** (Fornell, 1982, 1987) beinhalten regressionsbasierte Ansätze wie die multiple Regression, die logistische Regression und die Varianzanalyse, aber auch explorative und konfirmatorische Faktorenanalysen, Clusteranalysen sowie die multidimensionale Skalierung. Werden diese Verfahren auf bestimmte Forschungsfragen angewendet, so können sie entweder zur Bestätigung von zuvor aufgestellten Theorien oder zur Identifikation von Mustern oder Beziehungen zwischen Variablen eingesetzt werden. Sie werden **konfirmatorisch** genannt, wenn sie zum Test von Hypothesen benutzt werden, und **explorativ**, wenn sie zur Erkennung von Mustern und Zusammenhängen in Daten dienen, über die vorher kein oder wenig Wissen besteht.

Die Unterscheidung zwischen einem konfirmatorischen und einem explorativen Ansatz ist nicht immer eindeutig. Beispielsweise werden im Rahmen einer Regressionsanalyse die abhängigen und die unabhängigen Variablen normalerweise auf der Basis theoretischer Überlegungen definiert. Ziel der Regressionsanalyse ist es, diese theoretischen Überlegungen zu testen. Allerdings kann dieses Verfahren auch eingesetzt werden, um zu prüfen, inwieweit andere unabhängige Variablen ebenfalls einen Einfluss auf die abhängige Variable haben. Hierdurch kann eine vorhandene Theorie gegebenenfalls erweitert werden. Typischerweise fokussiert die Interpretation der Ergebnisse zuerst auf die unabhängigen Variablen, welche die abhängige Variable statistisch signifikant erklären (eher konfirmatorisch) und dann darauf, welche unabhängigen Variablen, die abhängige Variable relativ besser erklären (eher explorativ). Faktorenanalysen hingegen werden normalerweise eingesetzt, um Beziehungen zwischen Variablen zu entdecken, wodurch sich die Zahl der Variablen auf wenige zusammengesetzte Faktoren (d.h. Kombinationen von Variablen) reduzieren lässt. Die finale Zusammenstellung der Faktoren ergibt sich aus der explorativen Analyse und Bestimmung der (sofern vorhanden) Beziehungen in den Daten. Auch wenn dieses Verfahren vom Ansatz her explorativ ist (daher der Name), haben Forscher meist schon ein gewisses Vorwissen, wodurch zum Beispiel die Entscheidung über die Zahl der zu extrahierenden Faktoren beeinflusst werden kann (Sarstedt & Mooi, 2014). Im Vergleich dazu zielt die konfirmatorische Faktorenanalyse darauf ab, einen vorher definierten Faktor und seine zugehörigen Indikatoren zu testen und zu untermauern.

Die Verfahren der ersten Generation haben in den Sozialwissenschaften eine breite Anwendung gefunden. In den letzten 20 Jahren wenden sich Forscher allerdings zunehmend den **Verfahren der zweiten Generation** zu, um die Defizite der ersten Verfahrensgeneration zu überwinden (Abbildung 1.1). Diese sind unter dem Begriff der **Strukturgleichungsmodellierung (SEM)** bekannt und ermöglichen es Forschern nicht direkt beobachtbare Variablen in ihre Analysen zu integrieren, welche indirekt über Indikatorvariablen gemessen werden. Zudem erleichtern sie die Berücksichtigung von Messfehlern in den beobachteten Variablen (Chin, 1998).

Es lassen sich im Wesentlichen zwei Verfahren der SEM unterscheiden: die **kovarianzbasierte SEM** (Covariance-Based-SEM, CB-SEM) und die **Partial Least Squares SEM** (PLS-SEM, auch **PLS-Pfadmodellierung** genannt). Die CB-SEM wird hauptsächlich zur Bestätigung (oder Widerlegung) von Theorien (d. h. einem anhand wissenschaftlicher Methoden aufgestellten Set an systematisch verbundenen Hypothesen, die zur Erklärung und Prognose von bestimmten Zielgrößen dienen) eingesetzt. Hierzu wird geprüft, wie gut ein theoretisches Modell durch empirische Daten abgebildet werden kann. Im Gegensatz dazu wird die PLS-SEM hauptsächlich zur Entwicklung von Theorien im Rahmen der explorativen Forschung angewendet. Hierbei fokussiert die Analyse auf die Erklärung der abhängigen Variable(n). Wir werden diesen Unterschied im weiteren Verlauf des Kapitels genauer erläutern.

Die PLS-SEM findet zunehmend Verbreitung. Neben zahlreichen Artikeln zur Einführung in das Verfahren (z. B. Chin, 1998; 2010; Haenlein & Kaplan, 2004; Hair et al., 2011; Henseler et al., 2012; Henseler et al., 2009; Mateos-Aparicio, 2011; Rigdon, 2013; Roldán & Sánchez-Franco, 2012; Tenenhaus et al., 2005; Wold, 1985) und zu dessen Verbreitung in unterschiedlichsten Disziplinen (do Valle & Assaker, 2015; Hair et al., 2012a; Hair et al., 2012b; Richter et al., 2016; Ringle et al., 2012; Sarstedt et al., 2014b; Nitzl, 2016) gibt es mittlerweile auch umfassende englischsprachige Lehrbücher zur PLS-SEM (z. B. Hair et al., 2017a; Hair et al., 2018). Das hier vorliegende deutsche Buch soll Forschern im deutschsprachigen Raum den Einstieg in die PLS-SEM erleichtern und ihnen damit neue Forschungsmöglichkeiten eröffnen.

Grundlegendes zur Verwendung von Strukturgleichungsmodellen

Forscher müssen sich auf Basis der zugrundeliegenden Forschungsfrage sowie der zur Verfügung stehenden Daten für ein geeignetes Analyseverfahren entscheiden. Unabhängig davon, ob Verfahren der ersten oder der zweiten Generation angewendet werden, sollten stets einige grundsätzliche Überlegungen bei der Wahl des geeigneten Verfahrens angestellt werden. Zu den wichtigsten Überlegungen zählen folgende Aspekte: (1) Aggregation oder Zusammenführung von Variablen zu sogenannten Composite-Variablen, (2) Messung, (3) Skalenniveau, (4) Kodierung und (5) Verteilung der Daten.

Composite-Variablen

Eine **Composite-Variable** (im Englischen auch als *Variate* bezeichnet) ist eine Linearkombination von verschiedenen Variablen, die anhand der Forschungsfrage ausgewählt werden (Hair et al., 2010). Die Kombination von Variablen beinhaltet die Bestimmung von Gewichten (z.B. w_1 und w_2), deren anschließende Multiplikation mit den jeweiligen Beobachtungen (z.B. x_1 und x_2) und die Zusammenfassung aller Werte. Die mathematische Formel für eine Linearkombination mit fünf Variablen sieht wie folgt aus (wobei Composite-Variablen aus einer beliebigen Anzahl Variablen gebildet werden können):

$$\text{Composite-Variable} = w_1 \cdot x_1 + w_2 \cdot x_2 + \dots + w_5 \cdot x_5,$$

wobei x die individuellen Variablen und w deren Gewicht bezeichnen. Alle x Variablen (z.B. Fragen in einem Fragebogen) beinhalten die Antworten von mehreren Befragten und können in einer Datenmatrix angeordnet werden. Abbildung 1.2 zeigt eine solche Datenmatrix, wobei i einen Index darstellt, welcher die Nummer der Befragten (d.h. die jeweiligen Fälle) bezeichnet. Für jeden der i Befragten kann ein aggregierter Wert für die Composite-Variable errechnet werden.

Fall	x_1	x_2	...	x_5	Composite-Variable
1	x_{11}	x_{21}	...	x_{51}	v_1
...	...	...	...	...	...
i	x_{1i}	x_{2i}	...	x_{5i}	v_i

Abbildung 1.2 Datenmatrix

Messung

Die **Messung** ist ein grundlegendes Konzept bei der Durchführung sozialwissenschaftlicher Analysen. Wenn wir an Messung denken, erscheint uns sofort das Bild eines Lineals, mit dem wir die Größe eines Menschen oder die Länge eines Möbelstücks messen können. In unserem Leben gibt es aber noch viele andere Arten der Messung. Im Auto etwa finden wir Anzeigen zu Geschwindigkeit, Motortemperatur und Tankfüllung. Wenn wir krank sind, nutzen wir ein Thermometer, um die Höhe des Fiebers zu bestimmen und wenn wir eine Diät machen, stellen wir uns auf die Waage.

Die Messung beschreibt den Prozess der Zuordnung von Zahlen zu einer Variablen basierend auf eindeutigen und konstanten Regeln (Hair et al., 2016a). Diese Regeln werden angewendet, damit die Zahlen der Variable so zugeordnet werden, dass die Variable eine treffende Messung ermöglicht. Bei einigen Variablen sind diese Regeln sehr einfach zu befolgen, bei anderen Variablen kann dies aber auch schwierig sein. Ist die Variable zum Beispiel Geschlecht,

dann ist es leicht eine 1 für Frauen und eine 0 für Männer zuzuordnen. Auch bei Alter oder Größe ist die Bestimmung einer Zahl einfach. Was aber wenn die Variablen Zufriedenheit oder Vertrauen sind? Die Messung ist hier viel schwieriger, da derartige Phänomene abstrakt und komplex sind und sich nicht direkt beobachten lassen. Wir sprechen daher von der Messung von **latenten** (d. h. nicht beobachtbaren) **Variablen** oder **Konstrukten**.

Abstrakte Konstrukte wie Zufriedenheit oder Vertrauen können nicht direkt gemessen werden. Allerdings lassen sich charakteristische Elemente und Indikatoren von Zufriedenheit mit oder Vertrauen in beispielsweise Marken, Produkte oder Unternehmen messen. Nicht direkt beobachtbare Konstrukte lassen sich somit indirekt mit Hilfe von verschiedenen Indikatoren (auch Items oder manifeste Variablen genannt) messen. Jeder Indikator repräsentiert dabei einen eigenen Aspekt des betrachteten abstrakten Konstrukts. Ist das Konstrukt zum Beispiel die Zufriedenheit mit einem Restaurant, können die folgenden Indikatorvariablen (Fragen im Fragebogen) verwendet werden:

1. Der Geschmack des Essens war ausgezeichnet.
2. Die Schnelligkeit der Bedienung entsprach meinen Erwartungen.
3. Die Bedienung kannte sich gut mit der Speisekarte aus.
4. Die Hintergrundmusik war angenehm.
5. Das Preis-Leistungs-Verhältnis war gut.

Durch die Kombination verschiedener Indikatorvariablen bzw. **Items** lässt sich das übergeordnete Konstrukt Zufriedenheit mit einem Restaurant indirekt messen. Forscher verwenden häufig mehrere Items und bilden damit eine Multi-Item-Skala zur indirekten Messung abstrakter Konstrukte wie in dem Restaurantbeispiel. Die unterschiedlichen Items werden dabei zu einer Variablen zusammengefasst. In einigen Fällen geschieht dies durch einfache Addition der verschiedenen Items. In anderen Fällen erfolgt die Zusammenfassung unter Verwendung einer linearen Gewichtung der jeweiligen Items (siehe die Ausführungen zu Composite-Variablen). Die Logik hinter diesem Vorgehen besteht darin, dass abstrakte Konstrukte wie die Zufriedenheit mit einem Restaurant präziser gemessen werden, wenn mehrere Variablen verwendet werden. Die erwartete Verbesserung der Messgenauigkeit beruht auf der Annahme, dass verschiedene Items auch verschiedene Aspekte eines Konstrukts widerspiegeln, wodurch das Konstrukt selbst präziser gemessen wird. Dies beinhaltet auch eine Reduktion des **Messfehlers**, welcher die Differenz zwischen dem wahren Wert einer Variablen und dem empirisch gemessenen Wert beschreibt. Messfehler können unterschiedliche Ursachen haben, wie etwa ungenau formulierte Fragen in einem Fragebogen, Missverständnisse bei der Skalierung oder die nicht korrekte Anwendung eines statistischen Verfahrens. In multivariaten Analysen sind Messfehler somit sehr wahrscheinlich. Ziel ist es daher, diese Messfehler so weit wie möglich zu reduzieren.

Teilweise entscheiden Forscher sich für die Verwendung von **Single-Items** anstelle multipler Items, um Konstrukte wie Zufriedenheit oder Kaufbereitschaft zu messen. Zum Beispiel könnten wir nur die Frage „Alles in allem bin ich mit dem Restaurant sehr zufrieden“ anstelle der oben beschriebenen fünf Items verwenden. Hierdurch wird natürlich der Fragebogen deutlich kürzer,

gleichzeitig sinkt aber auch die Präzision unserer Messung. Wir werden die Grundlagen zur Messung und Beurteilung von Messmodellen in den folgenden Kapiteln genauer diskutieren.

Skalenniveau

Das **Skalenniveau** beschreibt eine vorbestimmte Anzahl geschlossener Antworten, die zur Beantwortung einer Frage verwendet werden können. Es lassen sich vier Typen von Messskalen unterscheiden, die jeweils ein unterschiedliches Skalenniveau repräsentieren: die Nominalskala, die Ordinalskala, die Intervallskala und die Ratioskala. **Nominalskalen** haben das geringste Skalenniveau, da sie nur in einem sehr beschränkten Maße den Einsatz statistischer Analyseverfahren erlauben. Eine Nominalskala beinhaltet Zahlen, anhand derer Objekte (z. B. Menschen, Unternehmen, Produkte etc.) identifiziert und klassifiziert werden können. Daher wird diese Messskala manchmal auch als **Kategorialskala** bezeichnet. Wenn beispielsweise in einem Interview nach dem Beruf gefragt wird und als Antwortkategorien Arzt/Ärztin, Anwalt/Anwältin, Lehrer/Lehrerin, Ingenieur/Ingenieurin usw. zur Verfügung stehen, dann liegt eine Nominalskala vor. Nominalskalen haben zwei oder mehrere Kategorien, wobei sich die Kategorien wechselseitig ausschließen. Jeder Kategorie kann eine Zahl zugewiesen werden, welche dann dazu verwendet werden kann, die Anzahl der Befragten in den einzelnen Kategorien oder deren prozentuale Verteilung zu bestimmen.

Die Messskala mit dem nächsthöheren Skalenniveau wird **Ordinalskala** genannt. Liegt eine Messskala mit einem ordinalen Skalenniveau vor, dann hat die Zu- bzw. Abnahme des Zahlenwertes eine inhaltliche Bedeutung. Werden beispielsweise Kunden bezüglich der Nutzung eines Produktes als Nichtnutzer = 0, Gelegenheitsnutzer = 1 und häufige Nutzer = 2 kodiert, dann wissen wir, dass höhere Werte mit einer gesteigerten Nutzungshäufigkeit einhergehen. Wird ein Sachverhalt mit Hilfe einer Ordinalskala gemessen, dann liegt eine Information über die Rangfolge zwischen den Beobachtungen vor. Eine Ordinalskala erlaubt allerdings keine Aussagen über die Abstände zwischen den Kategorien. Wir wissen beispielsweise nicht, ob die Differenz zwischen „Nichtnutzer“ und „Gelegenheitsnutzer“ genauso groß wie die Differenz zwischen „Gelegenheitsnutzer“ und „häufige Nutzer“ ist, obwohl die Differenz in den Werten (also $0 - 1$ und $1 - 2$) gleich ist. Daher können arithmetische Mittel oder Varianzen für Ordinalskalen nicht berechnet werden.

Lässt sich eine **Intervallskala** für die Messung verwenden, so existieren präzise Informationen über die Rangfolge der Merkmalsausprägungen und deren Abstand. Ist die Temperatur zum Beispiel 25°C , wissen wir, dass bei einem Rückgang auf 20°C die Differenz genau 5°C beträgt. Die Differenz liegt ebenso bei 5°C , wenn die Temperatur von 25°C auf 30°C steigt. Der einheitliche Abstand zwischen Skalenpunkten wird Äquidistanz genannt. Solche äquidistanten Messskalen sind für bestimmte Analyseverfahren, wie etwa die SEM, notwendig. Intervallskalen beinhalten allerdings keinen absoluten Nullpunkt. Liegt die Temperatur bei 0°C , dann ist es zwar kalt, die Temperatur kann aber noch weiter zurückgehen. Der Wert von 0 bedeutet somit nicht, dass es keine

Temperatur mehr gibt (Sarstedt & Mooi, 2014). Der Vorteil von Intervallskalen ist, dass fast jede mathematische Operation, inklusive der Berechnung von Mittelwerten und Standardabweichungen, durchgeführt werden kann. Zudem können Intervallskalen in alternative Intervallskalen umgewandelt werden. So wird etwa in einigen Ländern Grad Fahrenheit (°F) statt Grad Celsius (°C) zur Messung der Temperatur verwendet. Die Temperatur lässt sich anhand der folgenden Formel von Celsius in Fahrenheit umwandeln: Grad Fahrenheit = Grad Celsius $\cdot 9 / 5 + 32$. Auf die gleiche Art und Weise (mittels **Reskalierung**) lassen sich Daten, die auf einer Skala von 1 bis 5 gemessen wurden, in eine Skalierung von 0 bis 100 umwandeln: $[(\text{Datenpunkt auf der Skala von 1 bis 5}) - 1] / 4 \cdot 100$.

Die **Ratioskala** hat den größten Informationsgehalt. Wird etwas auf einer Ratioskala gemessen, wissen wir, dass bei einem Wert von 0 eine bestimmte Ausprägung einer Variablen nicht vorhanden ist. Wenn beispielsweise ein Kunde laut Skala keine Produkte kauft (Wert = 0), dann wissen wir, dass er tatsächlich nichts gekauft hat. Ein anderes Beispiel sind Unternehmen, die kein Geld für die Werbung neuer Produkte ausgeben (Wert = 0). Der Nullpunkt oder Ursprung der Variablen entspricht somit 0. Ratioskalen werden auch zur Messung von Längen, Mengen, Volumina oder zur Zeitmessung (z.B. verstrichene Minuten) eingesetzt. Ratioskalen ermöglichen sämtliche mathematischen Operationen.

Kodierung

Kodierung bezeichnet die Zuordnung von Nummern zu Skalen, um die Messung von Attributen zu ermöglichen. Im Rahmen von Umfrageforschungen werden Daten oft vorkodiert, d.h. schon vor der eigentlichen Befragung werden bestimmten Antworten (z.B. Skalenpunkte) Zahlen zugeordnet. Beispielsweise wird bei einer 9-Punkte-Skala zur Zustimmung die Zahl 9 der höchsten Zustimmung „stimme voll und ganz zu“ und die 1 der niedrigsten Zustimmung „stimme überhaupt nicht zu“ zugeordnet. Alle dazwischen abgestuften Zustimmungsurteile werden mit 2 bis 8 kodiert. Alternativ kann die Kodierung auch im Anschluss an eine Befragung erfolgen. Dies ist insbesondere dann sinnvoll, wenn im Rahmen von quantitativen Umfragen offene Fragen gestellt werden oder vollständig qualitative Befragungen durchgeführt werden.

Bei der Anwendung multivariater Analysen ist die Frage der Kodierung sehr wichtig, da diese festlegt, wann und wie verschiedene Skalentypen verwendet werden können. So lassen sich Variablen, die mit einer Intervall- oder Ratioskala gemessen wurden, jederzeit in multivariaten Analysen verwenden. Werden dagegen Ordinalskalen genutzt (was im Rahmen der SEM häufig der Fall ist), dann sollten Forscher darauf achten, dass die Bedingung äquidistanter Ausprägungen eingehalten wird. Wird beispielsweise eine 5-Punkte Likert-Skala mit den Kategorien (1) *stimme überhaupt nicht zu*, (2) *stimme nicht zu*, (3) *weder noch*, (4) *stimme zu* und (5) *stimme voll und ganz zu* verwendet, dann können wir annehmen, dass die „Distanz“ zwischen 1 und 2 die Gleiche wie zwischen 3 und 4 ist. Im Gegensatz dazu weist der gleiche Typ von Likert-Skala, bei dem jedoch (1) *stimme nicht zu*, (2) *weder noch*, (3) *stimme etwas zu*, (4)

stimme zu und (5) *stimme voll und ganz zu* als Kategorien verwendet werden, keine Äquidistanz auf, da nur eine einzige Bewertung unterhalb der neutralen Kategorie „weder noch“ möglich ist. Zudem wird eine derartige Skala die Ergebnisse sehr wahrscheinlich in Richtung eines deutlich besseren Ergebnisses verzerren. Eine gute Likert-Skala weist eine Symmetrie der Likert-Items um eine mittlere Kategorie sowie eine klare sprachliche Unterscheidung aller Kategorien auf. In einer derartigen symmetrischen Skala kann die Eigenschaft der Äquidistanz normalerweise angenommen werden. Wird eine Likert-Skala als symmetrisch und äquidistant wahrgenommen, funktioniert sie ähnlich einer Intervallskala. Somit kann eine gut entwickelte Likert-Skala, obwohl sie ordinal ist, einer Intervallmessung sehr nahe kommen und die betreffende Variable kann im Rahmen der SEM verwendet werden.

Verteilung der Daten

Wenn quantitative Daten erhoben werden, dann können die Antworten zu den gestellten Fragen als eine Verteilung über die vorhandenen (vordefinierten) Antwortkategorien ausgegeben werden. Wird beispielsweise eine 9-Punkte Zustimmungsskala verwendet, dann kann eine Verteilung der Antworten in den jeweiligen Antwortkategorien (1, 2, 3, ..., 9) berechnet und in einer Tabelle oder einem Diagramm dargestellt werden. Abbildung 1.3 zeigt ein

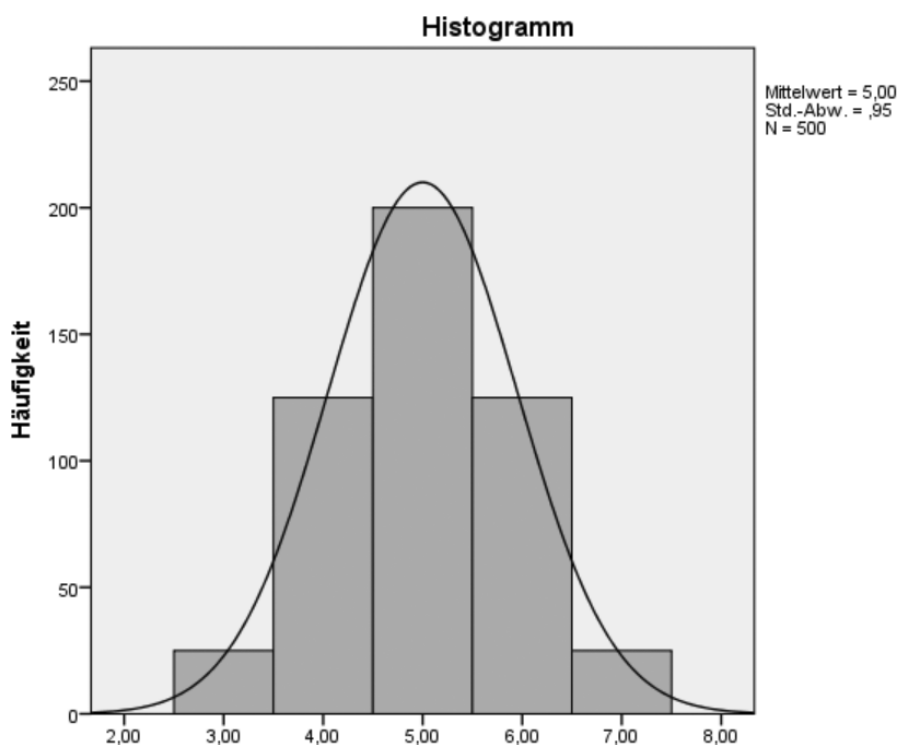


Abbildung 1.3 Häufigkeitsverteilung

Beispiel für die Häufigkeitsverteilung einer Variablen x . Es wird deutlich, dass die meisten Befragten eine 5 auf der 9-Punkte Skala, gefolgt von 4 und 6 und schließlich 3 und 7 angegeben haben. Die Häufigkeitsverteilung entspricht annähernd einer glockenförmigen, symmetrischen Kurve um den Mittelwert von 5. Eine solche Glockenkurve ist die Normalverteilung. Ihr Vorliegen ist für die Anwendung zahlreicher multivariater Analyseverfahren Voraussetzung, um einwandfrei Ergebnisse zu generieren.

Auch wenn viele verschiedene Formen von Verteilungen existieren (z.B. die Normal-, Binomial- oder die Poissonverteilung), brauchen Forscher, die mit der SEM arbeiten, nur zwischen normalen und nicht-normalen Verteilungen zu unterscheiden. Normalverteilungen sind für gewöhnlich wünschenswert, insbesondere im Rahmen der CB-SEM. Im Vergleich dazu macht die PLS-SEM keine Annahmen über die Datenverteilung. Wir werden später aber noch zeigen, dass es dennoch sinnvoll ist, die Verteilung zu berücksichtigen, wenn mit der PLS-SEM gearbeitet wird. Um herauszufinden, ob Daten normalverteilt sind, können statistische Tests wie der Kolmogorov-Smirnov-Test und der Shapiro-Wilk-Test eingesetzt werden (Sarstedt & Mooi, 2014). Zusätzlich können Forscher zwei Verteilungsmaße (Schiefe und Kurtosis; Kapitel 2) verwenden, mit deren Hilfe das Ausmaß der Abweichung von der Normalverteilung näher bestimmt werden kann (Hair et al., 2010).

Strukturgleichungsmodellierung mit Partial Least Squares Pfadmodellen

Pfadmodelle mit latenten Variablen

Pfadmodelle sind grafische Darstellungen, welche im Rahmen der SEM zur Veranschaulichung von Hypothesen und den Beziehungen zwischen Variablen erstellt werden (Hair et al., 2011; Hair et al., 2016a). Abbildung 1.4 zeigt ein Beispiel eines Pfadmodells.

Konstrukte (d.h. Variablen, die nicht direkt gemessen werden) werden in Pfadmodellen als Kreise oder Ellipsen (Y_1 bis Y_4) dargestellt. Die **Indikatoren**, auch Items oder manifeste Variablen genannt, sind die direkt gemessenen Variablen, in denen die Rohdaten enthalten sind. Sie werden in Pfadmodellen in der Form von Rechtecken dargestellt (x_1 bis x_{10}). Die Beziehungen zwischen den Konstrukten sowie zwischen den Konstrukten und den ihnen zugewiesenen Indikatoren werden durch Pfeile dargestellt. Die Pfeile haben bei der PLS-SEM immer nur eine Spitze, d.h. es werden immer direkte Beziehungen abgebildet. Hierdurch wird stets eine gerichtete Beziehung beschrieben, die auch als kausale Beziehung interpretiert werden kann, sofern dies durch eine Theorie gestützt wird.

Ein PLS-Pfadmodell besteht aus zwei Elementen. Dies ist zum einen das **Strukturmodell** (bei PLS-Pfadmodellen auch inneres Modell genannt), welches die Konstrukte (Kreise oder Ellipsen) repräsentiert. Zum anderen gibt es das **Messmodell** der Konstrukte (bei PLS-Pfadmodellen auch als äußeres Modell bezeichnet), welches die Beziehungen zwischen den Konstrukten und den

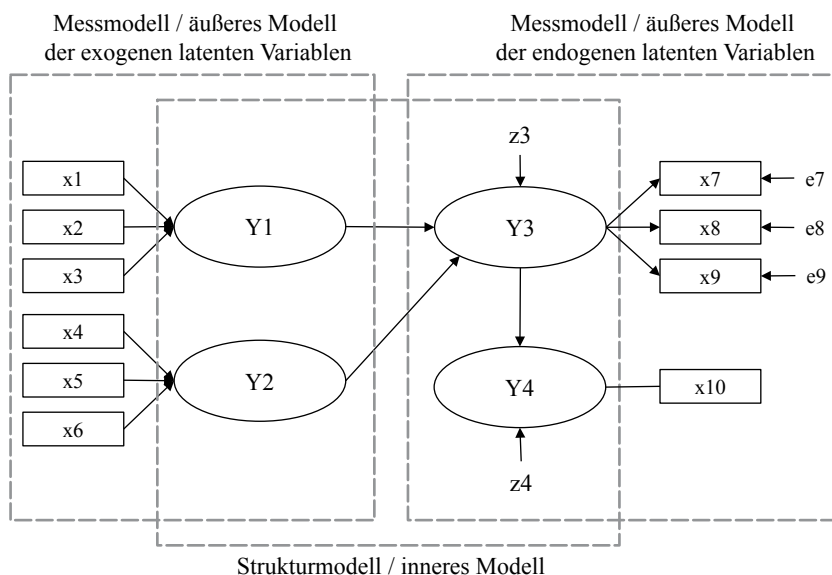


Abbildung 1.4 Ein einfaches Pfadmodell

Indikatoren (Rechtecke) darstellt. In Abbildung 1.4 gibt es zwei verschiedene Messmodelle: ein Messmodell für die exogenen latenten Variablen (d. h. diejenigen Konstrukte, welche die anderen Konstrukte im Modell erklären) und ein weiteres für die endogenen latenten Variablen (d. h. die Konstrukte, die im Modell erklärt werden). Statt zwischen den verschiedenen Messmodellen für exogene und endogene Variablen zu unterscheiden, beziehen Forscher sich meist auf das Messmodell einer konkreten latenten Variable. Beispielsweise sind x_1 bis x_3 die Indikatoren im Messmodell für Y_1 , während Y_4 nur durch x_{10} gemessen wird.

Die **Fehlerterme** (z. B. e_7 oder e_8 ; Abbildung 1.4) sind mit den (endogenen) Konstrukten und (reflektiv) gemessenen Variablen durch einfache Pfeile verbunden. Fehlerterme repräsentieren die nicht erklärte Varianz des geschätzten Pfadmodells. In Abbildung 1.4 finden sich die Fehlerterme e_7 bis e_{10} nur bei Indikatoren, bei denen die Beziehung vom Konstrukt zu den Indikatoren geht (d. h. es handelt sich hier um reflektive Indikatoren). Im Vergleich dazu haben die (formativen) Indikatoren x_1 bis x_6 , bei denen die Beziehung vom Indikator zum Konstrukt geht, keine Fehlerterme. Bei dem Single-Item-Konstrukt Y_4 ist die Richtung der Beziehung irrelevant, da hier das Konstrukt und der Indikator äquivalent sind. Aus dem gleichen Grund gibt es auch für x_{10} keinen Fehlerterm. Das Strukturmodell beinhaltet ebenfalls Fehlerterme. In Abbildung 1.4 sind z_3 und z_4 mit den endogenen latenten Variablen Y_3 und Y_4 verbunden (wobei zu beachten ist, dass die Fehlerterme von Konstrukten und gemessenen Variablen unterschiedlich bezeichnet werden). Die exogenen latenten Variablen, welche lediglich andere latente Variablen erklären, haben dagegen keinen Fehlerterm.

Pfadmodelle werden anhand einer Theorie entwickelt. Eine **Theorie** ist ein anhand wissenschaftlicher Methoden aufgestelltes Set an systematisch verbundenen Hypothesen, die zur Erklärung und Prognose von bestimmten Zielgrößen dienen. Hypothesen sind somit einzelne Annahmen, während Theorien mehrere logisch verknüpfte Hypothesen enthalten, die empirisch prüfbar sind. Zur Entwicklung von Pfadmodellen werden zwei Arten von Theorien benötigt: die Messtheorie und die Strukturtheorie. Letztere beschreibt, wie die Konstrukte innerhalb des Strukturmodells miteinander verbunden sind. Die Messtheorie macht dagegen Aussagen dazu, wie jedes einzelne Konstrukt gemessen wird.

Messtheorie

Die **Messtheorie** spezifiziert, wie die latenten Variablen (Konstrukte) gemessen werden. Grundsätzlich lassen sich zwei verschiedene Formen der Messung latenter Variablen unterscheiden: die formative und die reflektive Messung. Die Konstrukte Y_1 und Y_2 in Abbildung 1.4 basieren auf einem **formativen Messmodell**. Hierbei ist zu beachten, dass die Pfeile von den Indikatorvariablen (x_1 bis x_3 für Y_1 und x_4 bis x_6 für Y_2) zu den Konstrukten zeigen, was auf eine kausale (vorhersagende) Beziehung in diese Richtung hinweist.

Im Gegensatz dazu basieren Y_3 und Y_4 auf einem **reflektiven Messmodell**. Bei reflektiven Indikatoren zeigen die Pfeile vom Konstrukt zu den Indikatoren, was darauf hinweist, dass das Konstrukt die Messung der Indikatorvariablen (genauer gesagt deren Kovarianz) verursacht. Abbildung 1.4 verdeutlicht zudem, dass bei reflektiven Messungen jedem Indikator ein Fehlerterm zugeordnet ist, was bei formativen Messungen nicht der Fall ist. Bei letzteren wird angenommen, dass diese fehlerfrei sind (Diamantopoulos, 2011). Schließlich sei noch darauf hingewiesen, dass Y_4 mit einem einzigen Item anstelle multipler Items gemessen wird.

Die Frage der Messung von Konstrukten (d. h. formativ vs. reflektiv und Multi-Item vs. Single-Item) ist von grundlegender Bedeutung bei der Entwicklung von Pfadmodellen. In Kapitel 2 werden wir noch ausführlicher auf die jeweiligen Messmodelle und deren Unterschiede eingehen.

Strukturtheorie

Die **Strukturtheorie** beschreibt die Beziehungen zwischen den latenten Variablen (d. h. sie begründet die Konstrukte und die Pfadbeziehungen innerhalb des Strukturmodells). Die Anordnung und Abfolge der Konstrukte basiert auf einer Theorie oder der Erfahrung und dem Vorwissen eines Forschers. Werden Pfadmodelle entwickelt, werden sie von links nach rechts aufgestellt. Die Variablen auf der linken Seite stellen somit die unabhängigen Variablen dar, die Variablen auf der rechten Seite sind dagegen die abhängigen Variablen. Entsprechend sind die Variablen auf der linken Seite sequentiell vorgelagert und beeinflussen die Variablen auf der rechten Seite. Variablen können allerdings auch gleichzeitig als abhängige und unabhängige Variablen dienen.

Wenn latente Variablen nur als unabhängige Variablen dienen, werden sie als **exogene latente Variablen** bezeichnet (Y_1 und Y_2). Dienen latente Variablen dagegen nur als abhängige Variablen (Y_4) oder als unabhängige und abhängige Variablen (Y_3), werden sie **endogene latente Variablen** genannt. Jede Variable, die innerhalb des Strukturmodells nur von ihr ausgehende Pfeile hat, stellt eine exogene latente Variable dar. Endogene latente Variablen können dagegen Pfeile haben, die entweder ein- und ausgehen (Y_3) oder nur eingehen (Y_4). Zudem können wir festhalten, dass exogene latente Variablen keine Fehlerterme haben, da diese Konstrukte die unabhängigen Variablen darstellen, welche die abhängigen Variablen im Pfadmodell erklären.

PLS-SEM, CB-SEM und Regressionen auf Basis von Summenwerten

Es gibt zwei wesentliche Verfahren, um die Beziehungen innerhalb eines Strukturgleichungsmodells zu schätzen (Hair et al., 2010; Hair et al., 2011). Das eine ist die weiter verbreitete CB-SEM; das andere ist die PLS-SEM, welche im Fokus des vorliegenden Buches liegt. Beide Verfahren sind für unterschiedliche Forschungskontexte geeignet, daher müssen Forscher die Unterschiede kennen, um das jeweils korrekte Verfahren anzuwenden. Zudem haben sich einige Forscher für die Verwendung von Regressionen auf Basis von Summenwerten ausgesprochen anstelle einer Gewichtung von Indikatoren, wie sie in der PLS-SEM vorgenommen wird. Da dieser Ansatz allerdings keine erkennbaren Vorteile gegenüber der PLS-SEM enthält, werden wir hier nur kurz darauf eingehen und uns stattdessen auf die PLS-SEM und die CB-SEM fokussieren.

Um die Frage zu beantworten, wann die Verwendung der PLS-SEM oder der CB-SEM angebracht ist, sollten Forscher die Charakteristika und Ziele der beiden Verfahren vergleichen (Hair et al., 2012b). Falls eine Theorie noch nicht umfassend etabliert ist, sollten Forscher die Anwendung der PLS-SEM als Alternative zur CB-SEM in Betracht ziehen. Gleiches gilt, wenn das primäre Ziel der Verwendung der SEM die Erklärung und Prognose von Zielkonstrukten ist (Rigdon, 2012).

Ein wesentlicher konzeptioneller Unterschied zwischen der PLS-SEM und der CB-SEM besteht darin, wie das Verfahren latente Variablen innerhalb des Modells behandelt. Die CB-SEM betrachtet Konstrukte als Faktoren, welche die Kovariation zwischen den dazugehörigen Indikatoren erklären. Die Werte dieser gemeinsamen Faktoren sind weder bekannt noch werden sie benötigt, um die Parameter des Modells zu schätzen. Die PLS-SEM verwendet dagegen Proxies, welche das Konstrukt über die Gewichtung und Zusammenfassung von Indikatorvariablen, also über sogenannte Composite-Variablen, darstellen. Die PLS-SEM stellt daher einen composite-basierten Ansatz der SEM dar, welcher die strenge Annahme der CB-SEM lockert, wonach jegliche Kovariation innerhalb eines Sets an Indikatoren durch einen gemeinsamen Faktor erklärt wird (Henseler et al., 2014; Rigdon, 2012; Rigdon et al., 2014; Sarstedt et al., 2016c). Die Nutzung von gewichteten Composite-Variablen ermöglicht gleichzeitig die Berücksichtigung von Messfehlern, wodurch die PLS-SEM

klassischen multiplen Regressionsanalysen, in denen **Summenwerte** einzelner Variablen verwendet werden, überlegen ist. In letzterem Fall gehen Forscher von einem gleichen Gewicht der Indikatoren aus, so dass jeder Indikator in gleichem Maße zur Composite-Variablen beiträgt (Henseler et al., 2014). Bezugnehmend auf unsere Beschreibung von Composite-Variablen zu Beginn des Kapitels würde dies bedeuten, dass alle Gewichte w auf 1 gesetzt werden. Die entsprechende mathematische Formel für eine Linearkombination mit fünf Variablen wäre dann wie folgt:

$$\text{Composite-Variable} = 1 \cdot x_1 + 1 \cdot x_2 + \dots + 1 \cdot x_5.$$

Liegen beispielsweise bei einem Befragten die Werte 4, 5, 4, 6 und 7 für die fünf Variablen vor, dann ist der entsprechende Summenwert 26. Derartige Aggregationen oder Summenwerte sind einfach anzuwenden, allerdings werden hierdurch jegliche Unterschiede in den Gewichtungen der einzelnen Items nivelliert. Dies ist zwar ein gängiges Vorgehen in der Forschungspraxis, führt aber zu systematischen Fehlern bei der Parameterschätzung (z. B. Hair et al., 2017b). Darüber hinaus liefert die Kenntnis der Gewichte einzelner Indikatoren wichtige Hinweise auf deren Bedeutung für das Konstrukt in einem bestimmten Kontext (d. h. in Relation zu anderen Composite-Variablen innerhalb des Strukturmodells). Wird zum Beispiel die Kundenzufriedenheit gemessen, erfahren Forscher welche der durch die jeweiligen Items erfassten Aspekte den größten Einfluss auf die Ausprägung der Zufriedenheit haben.

Im Rahmen der PLS-SEM wird nicht davon ausgegangen, dass die erstellten Proxies mit den Konstrukten, welche sie ersetzen, identisch sind. Vielmehr werden sie explizit als Approximationen verstanden (Rigdon, 2012). Daher sehen einige Forscher die CB-SEM als die direktere und präzisere Methode zur empirischen Messung theoretischer Konstrukte an. Eine derartige Sichtweise ist allerdings insofern kurzfristig, da die in der CB-SEM abgeleiteten gemeinsamen Faktoren ebenfalls nicht notwendigerweise mit den konzeptionellen Variablen eines theoretischen Modells gleichgesetzt werden können. Tatsächlich gibt es immer eine große (Validitäts-)Lücke zwischen der Messung eines bestimmten Konstruktes und der theoretischen oder intendierten Aussage des Konstruktes (z. B. Rigdon, 2012; Rossiter, 2011).

Messungen als Approximationen zu verstehen scheint im Rahmen sozialwissenschaftlicher Forschung realistischer (z. B. Rigdon, 2014a) und macht die Differenzierung zwischen der PLS-SEM und der CB-SEM mit Blick auf die Behandlung von Konstrukten fragwürdig. Diese Ansicht wird auch durch die Art, wie die CB-SEM in der Forschungspraxis angewendet wird, gestützt. Wird die CB-SEM angewendet, dann führt das ursprünglich aufgestellte hypothetische Modell meist zu einem ungenügenden Fit. In der Konsequenz müssten Forscher das Modell verwerfen und die Studie überdenken (was normalerweise eine neue Datenerhebung erforderlich macht), insbesondere dann, wenn mehrere Variablen entfernt werden müssen, um den entsprechenden Fit zu erreichen (Hair et al., 2010). Alternativ wird häufig das ursprüngliche, theoretisch entwickelte Modell modifiziert, um Fit-Maße über empfohlene Gren-

zwerte zu bringen. Hierdurch gelangen Forscher zu einem Modell mit einem akzeptablen Fit, welches ihrer Ansicht nach durch die Theorie gestützt ist. Dies stellt allerdings ein Best-Case-Szenario dar, welches in der Realität meist nicht zutrifft. Vielmehr suchen Forscher explorativ nach Modellspezifikationen, in denen Teile des ursprünglichen Modells modifiziert werden, um einen zufriedenstellenden Modellfit zu erlangen. Die dabei entstehenden Modelle korrespondieren allerdings meist nicht sehr gut mit dem Ausgangsmodell und sind häufig zu simpel (Sarstedt et al., 2014a).

Neben den Unterschieden in der Philosophie des Messens hat die unterschiedliche Behandlung von latenten Variablen einen Einfluss auf die Anwendungsmöglichkeiten der Verfahren. Die CB-SEM erlaubt es Konstruktwerte zu schätzen, allerdings sind die geschätzten Werte nicht einheitlich. Somit kann über eine unendliche Anzahl verschiedener möglicher Sets an Konstruktwerten jeweils ein ähnlicher Modellfit erzielt werden. Eine zentrale Konsequenz dieser Unbestimmtheit der Konstruktwerte ist, dass die Korrelationen zwischen den latenten Variablen (hier den gemeinsamen Faktoren) und jeder Variablen außerhalb des Messmodells der latenten Variablen (hier des Faktormodells) ebenfalls unbestimmt sind. Die Korrelationen können hoch oder niedrig sein, je nachdem welches Set an Konstruktwerten (hier Faktorkonstrukten) gewählt wird. Diese Einschränkung macht die CB-SEM für Prognosen unbrauchbar (z. B. Dijkstra, 2014). Im Vergleich dazu hat die PLS-SEM den wesentlichen Vorteil, dass immer ein einziger, spezifischer (d. h. determinierter) Wert für jede latente Variable generiert wird, sobald die Indikatorgewichte bestimmt sind. Diese determinierten Werte sind Proxies für das zu messende Konstrukt, genau wie Faktoren die Proxies der konzeptionellen Variablen in der CB-SEM sind (Becker et al., 2013a). Die PLS-SEM verwendet die sogenannte Ordinary Least Squares (OLS) Regression mit diesen Proxies als Input und dem Ziel der Minimierung der Fehlerterme (d. h. der Residualvarianz) der endogenen Konstrukte. Kurz gesagt schätzt die PLS-SEM Koeffizienten (d. h. die Beziehungen im Pfadmodell), welche die R^2 -Werte der endogenen (Ziel-) Konstrukte maximieren. Diese Funktion verwirklicht das Prognoseziel der PLS-SEM. Die PLS-SEM ist somit die bevorzugte Methode, wenn die Theorieentwicklung und die Erklärung der Varianz (oder die Prognose von Konstrukten) die Zielsetzungen sind. Daher wird die PLS-SEM auch als ein **varianzbasierter** Ansatz der SEM angesehen.

Ferner sollte beachtet werden, dass die PLS-SEM zwar Ähnlichkeiten mit der **PLS-Regression**, einem weiteren multivariaten Analyseverfahren, aufweist, damit aber nicht gleichzusetzen ist. Die PLS-Regression ist ein regressionsbasierter Ansatz, welcher lineare Beziehungen zwischen mehreren unabhängigen und einer oder mehreren abhängigen Variablen untersucht. Im Unterschied zur OLS-Regression werden im Rahmen der PLS-Regression aggregierte Variablen oder Faktoren mittels Hauptkomponentenanalysen für die unabhängigen und abhängigen Variablen gebildet. Demgegenüber beruht die PLS-SEM auf vordefinierten Beziehungsnetzwerken zwischen den Konstrukten sowie zwischen den Konstrukten und ihren Indikatoren (Mateos-Aparicio, 2011 liefert eine detailliertere Gegenüberstellung der PLS-SEM und der PLS-Regression).

Bei der Frage der Nutzung oder Nichtnutzung der PLS-SEM sollten mehrere Aspekte berücksichtigt werden, die auch auf den grundlegenden Charakteristika des Verfahrens beruhen. So haben etwa die statistischen Eigenschaften des PLS-SEM-Algorithmus wichtige Funktionen in Bezug auf die verwendeten Daten und das verwendete Modell. Darüber hinaus beeinflussen die Eigenschaften des PLS-SEM-Verfahrens auch die Bewertung der Ergebnisse. Letztendlich hängt die Anwendung der PLS-SEM von vier wesentlichen Entscheidungskriterien ab (Hair et al., 2011; Hair et al., 2012b; Ringle et al., 2012): (1) den Daten, (2) den Modelleigenschaften, (3) dem PLS-SEM-Algorithmus und (4) den Fragen der Modellbewertung. Abbildung 1.5 fasst die Hauptmerkmale der PLS-SEM zusammen. Im weiteren Verlauf dieses Kapitels wird ein erster Einblick in diese Eigenschaften gegeben. Tiefergehende Erklärungen, insbesondere in Bezug auf den PLS-SEM-Algorithmus sowie die Bewertung der Ergebnisse, finden sich in den späteren Kapiteln dieses Buches.

Dateneigenschaften	
Stichprobengröße	• Keine Identifikationsprobleme bei kleinen Stichproben • Auch bei kleinen Stichprobengrößen wird im Allgemeinen eine hohe Teststärke (Power) erzielt • Große Stichproben erhöhen die Präzision (d.h. die Konsistenz) der PLS-SEM Schätzungen
Fehlende Werte	• Hoch robust, solange fehlende Werte unterhalb eines vertretbaren Niveaus bleiben
Verteilung	• Keine Verteilungsannahmen; die PLS-SEM ist ein nicht-parametrisches Verfahren
Skalenniveau	• Arbeitet mit metrischen Daten, quasi-metrischen (ordinalen) Daten und binär-kodierten Variablen (mit Einschränkungen) • Gewisse Einschränkungen, beispielsweise, wenn kategoriale Daten zur Messung endogener latenter Variablen verwendet werden
Modelleigenschaften	
Zahl der Items in den Messmodellen der Konstrukte	• Verarbeitet Konstrukte mit Single- und Multi-Item-Messungen
Beziehung zwischen den Konstrukten und ihren Indikatoren	• Einfache Integration von reflektiven und formativen Messmodellen möglich
Modellkomplexität	• Verarbeitet komplexe Modelle mit vielen Beziehungen im Strukturmodell • Größere Anzahl an Indikatoren hilfreich, um den PLS-SEM Bias zu reduzieren

Modellaufbau	Keine kausalen Schleifen (keine zirkulären Beziehungen) innerhalb des Strukturmodells erlaubt
Eigenschaften des PLS-SEM Algorithmus	
Ziel	Minimiert die Höhe der unerklärten Varianz bzw. maximiert die R^2-Werte
Effizienz	Konvergiert nach wenigen Iterationen zur optimalen Lösung (selbst bei komplexen Modellen und/oder umfangreichen Stichproben); effizienter Algorithmus
Natur der Konstrukte	Werden als Proxies der untersuchten theoretischen Konzepte gesehen; repräsentiert durch Composite-Variablen
Konstruktwerte	- Berechnet als Linearkombination ihrer Indikatoren - Für Prognosezwecke genutzt - Können als Input für anschließende Analysen verwendet werden - Nicht durch Unzulänglichkeiten in den Daten beeinflusst
Parameterschätzungen	- Beziehungen im Strukturmodell werden generell unterschätzt (PLS-SEM Bias) - Beziehungen im Messmodell werden generell überschätzt (PLS-SEM Bias) - Consistency at Large - Hohe Teststärke
Modellbewertung	
Bewertung des Gesamtmodells	Kein etabliertes globales Gütekriterium (Goodness-of-Fit-Maß)
Bewertung der Messmodelle	- Reflektive Messmodelle: Bewertung der Reliabilität und Validität anhand multipler Kriterien - Formative Messmodelle: Bewertung der Validität, Signifikanz und Relevanz der Indikatorgewichte sowie der Kollinearität der Indikatoren
Bewertung des Strukturmodells	Kollinearität zwischen Konstrukten, Signifikanz der Pfadkoeffizienten, Kriterien zur Bewertung der Prognosefähigkeit des Modells
Zusätzliche Analysen	- Importance-Performance-Analysen - Mediierende Effekte - Hierarchische Komponentenmodelle - Multigruppenanalysen - Identifikation und Behandlung unbeobachteter Heterogenität - Messmodellinvarianz - Moderierende Effekte

Abbildung 1.5 Hauptmerkmale der PLS-SEM

Quelle: In Anlehnung an Hair et al., 2011. Copyright © 2011, M. E. Sharpe, Inc. Nachdruck mit freundlicher Genehmigung durch Taylor & Francis Ltd., <http://www.tandfonline.com>.

Die PLS-SEM arbeitet effizient mit kleinen Stichproben, komplexen Modellen und macht keine Annahmen bezüglich der Verteilung der verwendeten Daten (Cassel et al., 1999). Zudem kann die PLS-SEM mit reflektiven und formativen Messmodellen sowie mit Single-Item-Konstrukten umgehen, ohne dass es zu Identifikationsproblemen kommt. Somit kann die PLS-SEM bei einem breiten Spektrum von Forschungssituationen eingesetzt werden. Dabei profitieren Forscher auch von der hohen Effizienz bei der Parameterschätzung, welche sich im Vergleich zu der CB-SEM in einer höheren **Teststärke** manifestiert. Eine höhere Teststärke bedeutet, dass die in der Population tatsächlich signifikanten Beziehungen in der PLS-SEM mit einer höheren Wahrscheinlichkeit auch als signifikant ausgewiesen werden. Das Gleiche gilt im Vergleich zu Regressionen mit Summenwerten, die ebenfalls eine geringere Teststärke als die PLS-SEM haben (Hair et al., 2017b).

Die PLS-SEM hat allerdings auch Grenzen. In ihrer Grundform kann die PLS-SEM nicht angewendet werden, wenn das Strukturmodell kausale Schleifen oder zirkuläre Beziehungen enthält. Frühe Erweiterungen des PLS-Algorithmus, die bislang noch nicht in die gängigen PLS-SEM-Softwareprogramme implementiert wurden, sind allerdings auch in der Lage mit zirkulären Beziehungen umzugehen (Lohmöller, 1989). Darüber hinaus sind die Möglichkeiten zum Testen und Bestätigen von Theorien eingeschränkt, da die PLS-SEM über kein globales Gütekriterium verfügt. In der neueren Forschung finden sich allerdings Bemühungen zur Entwicklung von globalen Gütekriterien, um die Anwendungsmöglichkeiten des Verfahrens zu erweitern (Bentler & Huang, 2014). Zum Beispiel haben Henseler et al. (2014) den Standardized-Root-Mean-Square-Residual (SRMR)-Index (also die standardisierte Wurzel der mittleren quadrierten Residuen) zur Validierung von PLS-Pfadmodellen vorgestellt, welcher die Abweichung zwischen den beobachteten Korrelationen und den durch das Modell implizierten Korrelationen misst. Dieses Maß ist ebenfalls in der SmartPLS 3 Software implementiert und wird im Rahmen der Ausführungen zur Evaluation von Strukturmodellen in Kapitel 6 ausführlicher diskutiert.

Eine weitere Eigenschaft der PLS-SEM ist, dass die Parameterschätzungen nicht optimal hinsichtlich ihrer Konsistenz sind – eine Eigenschaft, die fälschlicherweise oft als der PLS-SEM Bias angesehen wird (Kapitel 3). Obwohl die Befürworter der CB-SEM diesen Unterschied zwischen den beiden Verfahren betonen, zeigen Simulationsstudien, dass die Unterschiede zwischen den PLS-SEM- und den CB-SEM-Schätzungen sehr gering sind, sofern die Messmodelle die Minimalanforderungen hinsichtlich der Zahl der Indikatoren und deren Ladungen erfüllen. Genauer gesagt, sofern die Messmodelle mindestens vier Indikatoren enthalten und deren Ladungen den gängigen Normen (≥ 0.70) entsprechen, gibt es praktisch keinen Unterschied bei der Parametergenauigkeit der beiden Verfahren (z. B. Reinartz et al., 2009; Hair et al., 2017b). Für die meisten Studien hat dies somit keine praktische Relevanz (z. B. Binz Astrachan et al., 2014). Tatsächlich sollte die Abweichung der Parameterschätzungen in der PLS-SEM nicht als Bias betrachtet werden, sondern vielmehr als Ergebnis eines Algorithmus, in dem ein anderes Optimierungsziel verfolgt wird. Zudem liefert Dijkstra und Henselers (2015b, 2015a) konsistenter PLS-Ansatz

(consistent PLS, PLSc) korrekte Modellschätzer ohne auf die Stärken des PLS-Ansatzes (wie die Möglichkeit komplexe Modelle mit geringen Stichproben und/oder formativen Messmodellen zu schätzen) zu verzichten (für einen alternativen Ansatz vgl. Bentler & Huang, 2014). Auf <https://www.smartpls.com/documentation/pls-sem-compared-with-cb-sem> zeigen wir anhand des bekannten Technology Acceptance Models (Davis, 1989) eine vergleichende Analyse, bei der PLS, PLSc und verschiedene CB-SEM-Schätzverfahren wie der Maximum-Likelihood Ansatz verwendet werden. Der Vergleich zeigt, dass die Ergebnisse von PLS, PLSc und des CB-SEM-Maximum-Likelihood-Ansatzes sehr ähnlich sind, während andere CB-SEM-Schätzverfahren zu recht unterschiedlichen Ergebnissen führen.

In bestimmten Fällen, insbesondere dann, wenn wenig Vorwissen über die Beziehungen im Strukturmodell oder die Messung der Konstrukte vorliegt oder wenn der Schwerpunkt eher auf der Exploration statt Konfirmation liegt, ist die PLS-SEM der CB-SEM überlegen. Die PLS-SEM ist auch dann eine gute Alternative zum Testen von Theorien, wenn die Annahmen der CB-SEM bezüglich der Normalverteilung, der Stichprobengröße oder der Modellkomplexität verletzt werden oder wenn Unregelmäßigkeiten im Verlauf der Modellschätzung auftreten.

Abbildung 1.6 enthält einige Faustregeln bezüglich der Frage, ob die CB-SEM oder die PLS-SEM angewendet werden soll. Dabei wird deutlich, dass die PLS-SEM nicht als eine universelle Alternative zur CB-SEM empfohlen werden kann. Beide Verfahren unterscheiden sich aus statistischer Perspektive, sind für unterschiedliche Zwecke ausgelegt und basieren auf unterschiedlichen

Die PLS-SEM sollte angewendet werden, wenn

- das Ziel in der Prognose von Zielkonstrukten oder in der Identifikation wesentlicher „Treiber-Konstrukte“ liegt.
- formativ gemessene Konstrukte Teil des Strukturmodells sind. Hierbei ist zu beachten, dass formative Messungen auch in der CB-SEM verwendet werden können, allerdings erfordert dies Modifikationen der Konstruktspezifikationen (z. B. benötigt das Konstrukt sowohl formative als auch reflektive Indikatoren zur Identifikation).
- das Strukturmodell komplex ist (d. h. viele Konstrukte und Indikatoren enthält).
- die Stichprobe klein ist und/oder die Daten nicht normalverteilt sind.
- Konstruktwerte in anschließenden Analysen verwendet werden sollen.

Die CB-SEM sollte angewendet werden, wenn

- das Ziel im Testen von Theorien, der Theoriebestätigung oder im Vergleich alternativer Theorien besteht.
- die Fehlerterme eine zusätzliche Spezifikation benötigen (z. B. die Kovarianz).
- das Strukturmodell zirkuläre Beziehungen enthält.
- der Forschungsansatz ein globales Gütekriterium erfordert.

Abbildung 1.6 Faustregeln zur Auswahlentscheidung zwischen der PLS-SEM und der CB-SEM

Quelle: In Anlehnung an Hair et al., 2011. Copyright © 2011, M. E. Sharpe, Inc. Nachdruck mit freundlicher Genehmigung durch Taylor & Francis Ltd., <http://www.tandfonline.com>.