

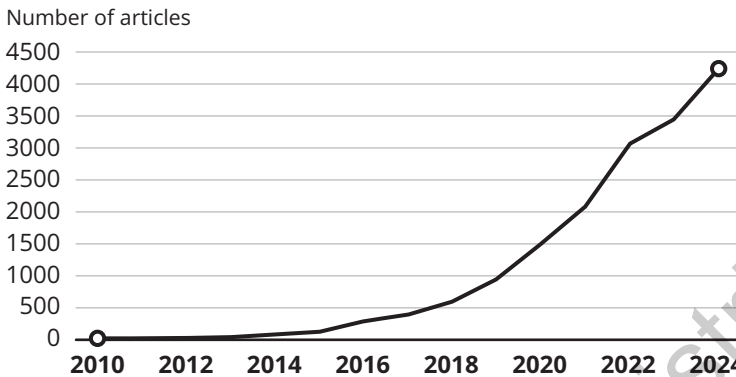
PREFACE

It has been more than a decade since the first edition of our book was published in 2014. In that period of time, the field of structural equation modeling (SEM), and particularly partial least squares structural equation modeling (PLS-SEM), has changed considerably. Although some traditional statistical methods have continued to evolve and extend their capabilities, PLS-SEM has expanded rapidly to include numerous additional analytical options. Much of the focus has been on the development of methods for confirming the quality of composite measures as representations of theoretical concepts (using procedures similar to the traditional confirmatory factor analysis in common factor models) and for assessing a model's out-of-sample predictive power. But many somewhat smaller analytical improvements have emerged as well.

When we wrote the first three editions, we were confident interest in PLS-SEM would increase, and this has been confirmed with each new edition. But even our wildest expectations of the growing interest in the PLS-SEM method have been exceeded. We never anticipated the interest in the PLS-SEM method would expand so rapidly in the last decade! The three previous editions of our book have been cited more than 75,000 times according to Google Scholar, and the books have been translated into more than a dozen other languages, including German (Hair, Hult, et al., 2017, 2024), Italian (Hair, Hult, et al., 2020), Spanish (Hair, Hult, et al., 2019), and French (Hair, Hult, Ringle, Sarstedt, Lux, & Troiville, 2022). Furthermore, the book has been published in an R software edition (Hair, Hult, et al., 2026). *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)* has been the premier text in the field of PLS-SEM for more than a decade, and based on the advances included in this new edition, we are confident it will remain the leading text in the future (Hair, Hult, Ringle, & Sarstedt, 2014, 2017, 2022).

A review of major social sciences journals clearly demonstrates applications of PLS-SEM have grown exponentially in the past decade, as evidenced in the popularity of the terms “partial least squares structural equation modeling,” “PLS-SEM,” and “PLS path modeling” in the Web of Science database (Figure 0.1). Since 2010, PLS-SEM has begun to gain traction in information systems research and marketing. Today, it is established as a standard multivariate analysis method, as reflected in the dedicated chapter on PLS-SEM in Hair, Black, Babin, & Anderson's (2019) *Multivariate Data Analysis* textbook. The method is now widely applied across disciplines, including management, marketing, information systems, medicine, engineering, psychology, political science, and environmental science as well as in specialized fields such as hospitality, accounting, and quality management (see Table 1.2 in Chapter 1, which provides an overview of

FIGURE 0.1 ■ Number of PLS-SEM-related Articles Per Year



Note: Number of articles returned from the Web of Science database for the search terms “partial least squares structural equation modeling,” “PLS-SEM,” and “PLS path modeling.”

PLS-SEM review studies across disciplines). The continuously increasing popularity of PLS-SEM stems from its ability to estimate complex models with latent variables, its relatively relaxed data requirements, and its dual usefulness for both prediction and theory development, making it a versatile tool for researchers seeking to explore complex relationships and predict outcomes.

PLS-SEM has also enjoyed increasingly widespread interest among methods researchers. A growing number of scholars have advanced the method, and they complement the initial core group of authors that have shaped the method (Angelelli, Ciavolino, Ringle, Sarstedt, & Aria, 2025; Ciavolino, Aria, Cheah, & Roldán, 2022; Khan et al., 2019). Their research papers offer novel perspectives on the method, sometimes sparking significant debates. Prominent early examples include the rejoinders to Rigdon’s (2012) *Long Range Planning* article by Bentler and Huang (2014), Dijkstra (2014), Sarstedt, Ringle, Henseler, and Hair (2014), and Rigdon (2014b) himself. Under the general theme “rethinking partial least squares path modeling,” this exchange of thoughts offered the point of departure for some of the most important PLS-SEM developments in the last few years. Other articles have further clarified the similarities and differences between PLS-SEM and covariance-based SEM, which has long been viewed as the default method for analyzing causal models. For example, Bauer, Mayer, Fuchs, and Schamberger (2025), Cho, Sarstedt, and Hwang (2022), Hair, Hult, Ringle, Sarstedt, and Thiele (2017), Sarstedt, Hair, Ringle, Thiele, and Gudergan (2016) discussed the measurement philosophies underlying the two SEM methods and demonstrated the biases that occur when using PLS-SEM and covariance-based SEM (CB-SEM) on models, which are inconsistent with what the methods assume. Related to this debate, Rigdon, Sarstedt, and Ringle (2017)

argued how differences in philosophy of science and different expectations about the research situation tend to induce a preference for one method over the other (see also Hair & Sarstedt, 2019). Similar discussions have emerged in psychology, where researchers acknowledge that reducing measurement to only the philosophy assumed by covariance-based SEM is a restrictive view, which does not apply to nearly all constructs (Rhemtulla, van Bork, & Borsboom, 2020).

Shmueli, Ray, Velasquez Estrada, and Chatla (2016) and Liengaard et al. (2021) have made substantial contributions to the field by shifting researchers' focus to the assessment of PLS path models' predictive power. To transition the emphasis from explanatory model assessment in applications of PLS-SEM, the PLS_{predict} procedure and the cross-validated predictive ability test (CVAPT) were introduced, which enable evaluation of a model's out-of-sample predictive power and predictive model comparisons. Their research has sparked a series of follow-up studies, offering guidelines on how to conduct predictive power assessments in a PLS-SEM framework (e.g., Chin et al., 2020; Sharma et al., 2024; Shmueli et al., 2019).

The growing prominence of PLS-SEM has also attracted scholars who have critically examined aspects of the method. The criticism ranged from the suitability of PLS-SEM for small sample sizes (e.g., Rönkkö & Evermann, 2013) to the performance of model evaluation metrics (e.g., Rönkkö, Lee, Evermann, McIntosh, & Antonakis, 2023) and more recently to its ability to estimate reflective measurement models (e.g., Henseler, Schuberth, Lee, & Kemény, 2025)—see Evermann and Rönkkö (2023). Many of these points have long been debunked in the methodological literature as representing borderline conditions with no relevance for practical research or as grounded in different assumptions (e.g., Cook & Forzani, 2023; Hair, Sarstedt, Ringle, Sharma, Liengaard, 2024; Henseler et al., 2014; Sharma, Liengaard, Sarstedt, Hair, & Ringle, 2023; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). For example, the notion that PLS-SEM cannot estimate reflective measurement models originates from the mistaken conflation of the data generation model with the measurement model specification. Although critics equate the two (i.e., reflective measurement models are common factor models), it is important to differentiate between the model estimation perspective (i.e., the way the method computes proxies of the conceptual variables using either common factor- or composite-based SEM) and the measurement-theoretic perspective (i.e., deciding whether to specify a construct reflectively or formatively; Guenther Guenther, Ringle, Zaefarian, & Cartwright, 2025). Similarly, bemoaning that inference testing in PLS-SEM violates test assumptions in a model with only uncorrelated constructs is certainly tenable (Rönkkö & Evermann, 2013; Rönkkö, Lee, Evermann, McIntosh, & Antonakis, 2023). However, concluding that PLS-SEM should not be used for inference testing, even though this phenomenon occurs only in a specific model setup and does not lead to type I errors, presents at best an incomplete picture of the method's actual performance (Hair, Sarstedt, Ringle, Sharma, Liengaard, 2024; Henseler et al., 2014;

Sharma, Liengaard, Sarstedt, Hair, & Ringle, 2023). We encourage authors to keep a cool head when confronted with such “arguments” and to always consider both sides. We believe it is important to reemphasize our previous call: “Any extreme position that (often systematically) neglects the beneficial features of the other technique and may result in prejudiced boycott calls [. . .] is not good research practice and does not help to truly advance our understanding of methods (or any other research subject)” (Hair, Ringle, & Sarstedt, 2012, p. 313; see also Petter, 2018; Sarstedt, Ringle, Henseler, & Hair, 2014; Russo & Stol, 2023).

Enhancement of the methodological foundations of the PLS-SEM method has been accompanied by the release of multiple new versions of SmartPLS 4 (Ringle, Wende, & Becker, 2024), which implement most of the latest extensions of the method (see <https://www.smartpls.com>). These updates incorporate a broad range of new algorithms and major new features that previously were not available or had to be executed manually (Cheah Magno, & Cassia, 2024). Considering the developments, a new edition of the book is clearly timely and warranted.

Although there are numerous published articles on PLS-SEM, until our first three editions and even today, there are no other widely adopted comprehensive books that explain the fundamental aspects of the method, particularly in a way that can be understood by individuals with limited statistical and mathematical training. This fourth edition of our book updates and extends the coverage of PLS-SEM for social sciences researchers and creates awareness of the most recent developments in an analytical tool that will enable scholars and practitioners to pursue research opportunities in many new and different ways.

The approach of this book is based on the authors’ many years of conducting and teaching research as well as the desire to communicate the fundamentals of the PLS-SEM method to a much broader audience. To accomplish this goal, we have limited the emphasis on equations, formulas, Greek symbols, and so forth that are typical of most books and articles. Instead, we explain in detail the basic fundamentals of PLS-SEM and provide rules of thumb that can be used as general guidelines for understanding and evaluating the results of applying the method. We also rely on a single software package (SmartPLS 4; <https://www.smartpls.com>) that can be used not only to complete the exercises in this book but also in the reader’s own research.

As a further effort to facilitate learning, we focus on a single case study throughout the book. The case is drawn from a published study on corporate reputation, and we believe it is general enough to be understood by many different areas of social sciences research, thus further facilitating comprehension of the method. Review and critical thinking questions are posed at the end of the chapters, and key terms are defined to better understand the concepts. Finally, suggested readings and extensive references are provided to enhance more advanced coverage of the topic.

We are excited to share with you the many new topics we have included in this edition. The table that follows highlights the key changes in the fourth edition of our book.

Chapter 1—An Introduction to Structural Equation Modeling

- Updated compilation of PLS-SEM review studies, textbooks, and edited volumes
- Extended discussion of the differences between measurement-theoretic and model estimation perspectives when running SEM analyses
- More details regarding metrological uncertainty and its sources
- Introduction of the EP-mixed framework
- Extended coverage of recent research criticizing PLS-SEM use for common factor models
- Overview of recent research on the nature of composite-based modeling
- Recent research on missing data treatment options in PLS-SEM

Chapter 2—Specifying the Path Model and Examining Data

- Details on model parsimony and model overfit
- Discussion of nonlinear effects
- Recent research on the interpretation of control variables
- Extended discussion of composite versus causal indicators
- Newest research on the limitations of sum scores
- Updated guidance on how to treat categorical control variables
- Extended discussion of the conceptual foundations of binary moderators

Chapter 3—Path Model Estimation

- Recent research on Mode A/Mode B estimation
- Inclusion of the principal component analysis (PCA) weighting scheme
- Use of unstandardized data in model estimation
- Updated overview of model evaluation steps
- Extended discussion of model fit in PLS-SEM

Chapter 4—Evaluation of Reflective Measurement Models

- Detailed coverage of content validity
- Emphasis of a multimetric approach to measurement validation
- Revised procedure for indicator deletion based on indicator reliabilities
- Removal of outdated discriminant validity procedures, extended coverage of the HTMT metric, and recent discussions regarding its use
- Further coverage of internal consistency reliability using ρ_A

Chapter 5—Evaluation of Formative Measurement Models

- More details on the relationships among Mode A/Mode B and outer loading estimates
- Description of fixed versus random seed initialization in bootstrapping
- Latest research on bootstrapping settings and assessment

(Continued)

(Continued)

Chapter 6—Evaluation of the Structural Model
• Discussion of standardized versus unstandardized path coefficient estimates • Updated guidelines on $PLS_{predict}$ • Inclusion of the cross-validated predictive ability test (CVPAT) • Distinctions among earliest and direct antecedents in predictive power assessment • Guidelines for the tandem use of $PLS_{predict}$ and CVPAT • Extended coverage of prediction-oriented model selection using CVPAT
Chapter 7—Mediator and Moderator Analysis
• Extended case study to showcase parallel mediation • Introduction of the mediation effect size omega • Discussion of binary moderators • More focused discussion of ways to create the interaction term
Chapter 8—Outlook on Advanced Methods
• Description of the combined IPMA (cIPMA) • Coverage of recent research on the Gaussian copula approach for endogeneity assessment • More concise discussion of multigroup analysis • Coverage of two new latent class approaches: PLS-SEM kernel k-means clustering (PLS-SEM KKC) and PLS-SEM-based sparse clustering
Overall: All case studies use the newest version of SmartPLS 4 (the previous edition used SmartPLS 3).

All examples in the edition are updated using the newest version of the most widely applied PLS-SEM software—SmartPLS 4 (Ringle, Wende, & Becker, 2024), whereas previous versions of our book used SmartPLS 3 (Ringle, Wende, & Becker, 2015) and SmartPLS 2 (Ringle, Wende, & Will, 2005).

The book chapters and learning support supplements are organized around the learning outcomes shown at the beginning of each chapter. Moreover, instead of a single summary at the end of each chapter, we present a separate concise summary for each learning outcome. This approach makes the book more understandable and usable for both students and teachers. The SAGE website for the book also includes other support materials to facilitate learning and applying the PLS-SEM method. Additionally, the PLS-SEM Academy (<https://www.pls-sem-academy.com>) offers video-based online courses based on this book and its earlier editions and also on advanced PLS-SEM topics following the explanations of Hair, Sarstedt, Ringle, and Gudergan (2024).

We would like to acknowledge the many insights and suggestions provided by the reviewers as well as a number of our colleagues and students. Most notably, we thank Julen Castillo Apraiz (University of the Basque Country), Vikas Arya (Paris School of Business), Norzalita Aziz (National University of Malaysia), Dipayan Biswas (University of South Florida), Jan-Michael Becker (BI Norwegian Business School), Zakariya Belkhamza (Ahmed Bin Mohammed Military College), Charla Brown (Troy University), Severina Cartwright (University of Liverpool), Fabio Cassia (University of Verona), Wynne W. Chin (University of Houston), Enrico Ciavolino (University of Salento), Gabriel Cepeda-Carrión (University of Seville), Gyeongcheol Cho (Ohio State University), Siegfried P. Gudergan (James Cook University), Jun-Hwa (Jacky) Cheah (University of East Anglia), Svenja Damberg (Aalborg University), Nicholas Danks (Trinity College Dublin), Adamantios Diamantopoulos (University of Vienna), Xuefei N. Deng (California State University, Dominguez Hills), Markus Eberl (Kantar), Janna Ehrlich (University of Applied Sciences Stralsund), George Franke (University of Alabama), Pooja Goel (University of Delhi), Jerónimo García-Fernández (University of Seville), Miriam Guenther (University of Liverpool), Peter Guenther (University of Liverpool), Anne Gottfried (University of Texas, Arlington), Siegfried P. Gudergan (James Cook University), Karl-Werner Hansmann (University of Hamburg), Dana Harrison (East Tennessee State University), Sven Hauff (Helmut Schmidt University), Cornelius Herstatt (Leuphana University), Mike Hollingsworth (Old Dominion University), Philip Holmes (Pensacola Christian College), Chris Hopkins (Auburn University), Lucas Hopkins (Florida State University), Heungsun Hwang (McGill University), John Hulland (University of Georgia), Ida R. Ismail (National University of Malaysia), K.D. Joshi (University of North Carolina Wilmington), Mirella Kleijnen (Vrije Universiteit Amsterdam), Anjala S. Krishen (University of Nevada), Puja Khatri (Guru Gobind Singh Indraprastha University), David Ketchen (Auburn University), Wolfgang Kersten (Hamburg University of Technology), Marcel Lichters (Otto von Guericke University Magdeburg), Benjamin Liengaard (Aarhus University), Yide Liu (Macau University of Science and Technology), Francesca Magno (University of Bergamo), Vasilica M. Margalina (University of Alba Iulia), José A. T. Machado (Polytechnic Institute of Porto), Lucy Matthews (Middle Tennessee State University), Mumtaz Ali Memon (SDU University Kazakhstan), Adam Merkle (University of South Alabama), Ovidiu I. Moisescu (Babeş-Bolyai University), Arthur Money (Henley Business School), Christian Nitzl (University of the Bundeswehr Munich), Michael-Jörg Oesterle (University of Stuttgart), Eleonora Pantano (University of Bristol Business School), Jung-Kun Park (Hanyang University), Torsten Pieper (University of North Carolina), Maria Petrescu (Nova Southeastern University), Catherine Prentice (University of Southern Queensland), Lacramioara Radomir (Babeş-Bolyai University), Arun Rai (Georgia State University), Sascha Raithel (Freie Universität Berlin), S. Mostafa Rasoolimanesh (Taylor's University), Shashank Rao (Auburn University), Soora Rasouli (Eindhoven University of Technology), Soumya Ray (National Tsing Hua University), Nicole F. Richter (University of Southern

Denmark), Ivan Russo (University of Verona), Edward E. Rigdon (Georgia State University), Jeff Risher (Southeastern Oklahoma University), José Luis Roldán (University of Seville), Amit Saini (University of Nebraska-Licolln), Francesco Scafarto (University of Rome “Tor Vergata”), Rainer Schlittgen (University of Hamburg), Manfred Schwaiger (Ludwig-Maximilians-University Munich), Pratyush N. Sharma (University of Alabama), Wen-Lung Shiau (Zhejiang University of Technology), Atul Shiva (Panjab University), Galit Shmueli (National Tsing Hua University), Noemi Sinkovics (Newcastle University), Rudolf R. Sinkovics (Durham University), Donna Smith (Ryerson University), Detmar W. Straub (Temple University), Hiram Ting (i-CATS University College), Ramayah Thurasamy (Universiti Sains Malaysia), Maomi Ueno (The University of Electro-Communications, Tokyo), Carl M. Wallenburg (WHU – Otto Beisheim School of Management), Sven Wende (SmartPLS GmbH), Anita Whiting (Clayton State University), Joachim Wolf (Kiel University), and Ghasem Zaefarian (University of Leeds) for their helpful remarks.

Also, we thank the team of doctoral students and research fellows at Hamburg University of Technology and Ludwig-Maximilians-University Munich—namely, Susanne Adler, Monika Benning, Charlotte Kreienbaum, Johanna Lorenz, Lea Rau, and Tobias Regensburger—for their kind support. In addition, at SAGE we thank Leah Fargotstein for her support and great work. We hope this book will expand knowledge of the capabilities and benefits of PLS-SEM to a much broader group of researchers and practitioners. Last, if you have any remarks, suggestions, or ideas to improve this book, please get in touch with us. We appreciate any feedback on the book’s concept and contents!

Joseph F. Hair, Jr.
University of South Alabama

G. Tomas M. Hult
Michigan State University

Christian M. Ringle
Hamburg University of Technology, Germany and
James Cook University, Australia

Marko Sarstedt
Ludwig-Maximilians-University Munich, Germany and
Babeş-Bolyai University, Romania

Visit the companion site for this book at <https://www.pls-sem.net/>

1

AN INTRODUCTION TO STRUCTURAL EQUATION MODELING

LEARNING OUTCOMES

1. Understand the meaning of structural equation modeling (SEM) and its relationship to multivariate data analysis.
2. Be able to describe the basic elements of a structural equation model.
3. Comprehend the basic concepts of partial least squares structural equation modeling (PLS-SEM).
4. Gain the knowledge to explain the differences between covariance-based structural equation modeling (CB-SEM) and PLS-SEM and when to use each.

CHAPTER PREVIEW

Social science researchers have been using statistical analysis tools for many years to extend their ability to develop, explore, and confirm research findings (e.g., Manley, Williams, & Hair, 2024). Whereas first-generation statistical methods, such as factor analysis and regression analysis, are frequently used in the social sciences and beyond, many problems researchers typically encounter in

their research require the application of second-generation methods. These methods allow estimating more complex models with unobservable variables while accounting for measurement error. In this chapter, we explain the fundamentals of second-generation statistical methods and establish a foundation that will enable you to understand and apply one of the emerging second-generation tools, referred to as **partial least squares structural equation modeling (PLS-SEM)**.

WHAT IS STRUCTURAL EQUATION MODELING?

Statistical analysis has been an essential tool for social science researchers for more than a century. Applications of statistical methods have expanded dramatically, particularly in recent years, with widespread access to many more methods due to user-friendly software and the proliferation of data-driven education and thinking. Researchers initially relied on univariate and bivariate analysis to understand data and relationships. To comprehend more complex relationships associated with current research directions in the social science disciplines, it is increasingly necessary to apply more sophisticated multivariate data analysis methods—including within the context of artificial intelligence and machine learning (e.g., Richter & Tudoran, 2024).

Multivariate analysis involves the application of statistical methods that simultaneously analyze multiple variables. The variables typically represent measurements associated with individuals, companies, events, activities, situations, and so forth. The measurements are often obtained from surveys or observations used to collect primary data, but they are increasingly available from databases consisting of secondary data. Table 1.1 displays some of the major types of statistical methods associated with multivariate data analysis.

TABLE 1.1 ■ Organization of Multivariate Methods		
	Primarily Exploratory	Primarily Confirmatory
First-generation techniques	• Cluster analysis • Exploratory factor analysis • Multidimensional scaling	• Analysis of variance • Logistic regression • Multiple regression • Confirmatory factor analysis (CFA)
Second-generation techniques	• Partial least squares structural equation modeling (PLS-SEM)	• Covariance-based structural equation modeling (CB-SEM)

The statistical methods often used by social scientists are typically called **first-generation techniques** (Fornell, 1982, 1987). These techniques, shown in the upper part of Table 1.1, include regression-based approaches, such as multiple regression, logistic regression, and analysis of variance, but also techniques, such as exploratory and confirmatory factor analysis, cluster analysis, and multidimensional scaling. When applied to a research question, these methods can be used to either confirm a priori established theories or identify data patterns and relationships. Specifically, they are **confirmatory** when testing the hypotheses of existing theories and concepts and **exploratory** when they search for patterns in the data in case there is either no or only little prior knowledge on how the variables are related.

It is important to note the distinction between confirmatory and exploratory is not always as clear-cut as it may seem. For example, when running a regression analysis, researchers usually select the dependent and independent variables based on established theories and concepts. The goal of the regression analysis is then hypothesized relationships. However, the technique can also be used to explore whether additional independent variables prove valuable for extending the concept being tested. The findings typically focus first on which independent variables are statistically significant predictors of the single dependent variable (more confirmatory) and then on which independent variables are, relatively speaking, better predictors of the dependent variable (more exploratory). In a similar fashion, when exploratory factor analysis is applied to a data set, the method searches for relationships between the variables in an effort to reduce a large number of variables to a smaller set of factors. The final set of factors is a result of exploring relationships in the data and reporting the meaningful relationships that are found (if any). Nevertheless, although the technique is exploratory in nature (as the name already suggests), researchers often have theoretical knowledge that may, for example, guide their decision on how many factors to extract from the data (Sarstedt & Mooi, 2019; Chapter 8). In contrast, confirmatory factor analysis is specifically designed for testing and substantiating an a priori determined set of factor(s) and its assigned indicators.

First-generation techniques have been widely applied by social science researchers and have substantially shaped the way we view the world today. In particular, methods such as multiple regression, logistic regression, and analysis of variance have often been used to test relationships among variables of interest empirically. However, what is common to these techniques is that they share three limitations (Haenlein & Kaplan, 2004):

1. *Postulation of a simple model structure:* Concerning this first limitation, multiple regression analysis and its extensions postulate a simple model structure involving one layer of dependent and independent variables. Causal chains such as “A leads to B leads to C” or more complex nomological networks involving a large number of intervening variables can only be estimated piecewise rather than simultaneously, which can

have severe consequences for the quality of the results (Sarstedt, Hair, Nitzl, Ringle, & Howard, 2020).

2. *Assumption that all variables can be considered observable:* With regard to this second limitation, standard regression methods are restricted to processing observable variables, such as age or sales (in units or dollars). Theoretical concepts, which are “abstract, unobservable properties or attributes of a social unit or entity” (Bagozzi & Philipps, 1982, p. 465), can be considered only after prior stand-alone validation by means of, for example, a confirmatory factor analysis. This ex-post inclusion of measures of theoretical concepts, however, comes with limitations.
3. *Conjecture that all variables are measured without error:* In accordance with the previous point, it is important to bear in mind that each observation of the real world is subject to a certain amount of measurement error, which can be either systematic or random (see Chapter 4). First-generation techniques are, strictly speaking, applicable only when there is neither systematic nor random error. This situation is, however, rarely encountered in reality, particularly when the aim is to estimate relationships among measures of theoretical concepts. Because the social sciences and many other fields of scientific inquiry routinely deal with theoretical concepts such as perceptions, attitudes, and intentions, these limitations of first-generation techniques are fundamental.

To overcome these limitations, researchers have increasingly been turning to **second-generation techniques**. These methods, referred to as **structural equation modeling (SEM)**, enable researchers to simultaneously model and estimate complex relationships among multiple dependent and independent variables. The concepts under consideration are typically unobservable and measured indirectly by multiple indicator variables. In estimating the relationships, SEM accounts for measurement error in observed variables. As a result, the method obtains a more precise measurement of the theoretical concepts of interest (Cole & Preacher, 2014). We will discuss these aspects in the following sections and chapters in greater detail.

There are two types of SEM methods: **covariance-based structural equation modeling (CB-SEM)** and partial least squares structural equation modeling (PLS-SEM; also called **PLS path modeling**). CB-SEM is primarily used to confirm (or reject) theories (i.e., a set of systematic relationships among multiple variables that can be tested empirically). It does this by determining how well a proposed theoretical model can estimate the covariance matrix for a sample data set. In contrast, PLS has been introduced as a “causal-predictive” approach to SEM (Jöreskog & Wold, 1982, p. 270; see also Wold, 1982), which focuses on explaining the variance in the model’s dependent variables (Chin et al., 2020;

Hair, Sarstedt, Ringle, Sharma, & Liengaard, 2024). We explain these concepts and further differences in more detail later in the chapter.

PLS-SEM has evolved rapidly as a statistical modeling technique. Over recent decades, there have been numerous introductory articles on the method (e.g., Chin, 1998; Haenlein & Kaplan, 2004; Hair, Risher, Sarstedt, & Ringle, 2019; Legate, Hair, Chretien, & Risher, 2023; Legate, Ringle, & Hair, 2024; Nitzi & Chin, 2017; Rigdon, 2013; Ringle, Sarstedt, Sinkovics, & Sinkovics, 2023; Roldán & Sánchez-Franco, 2012; Sarstedt, Hair, & Ringle, 2023; Tenenhaus Esposito Vinzi, Chatelin, & Lauro, 2005; Wold, 1985) as well as review articles examining how researchers across different disciplines have used PLS-SEM (Table 1.2). In light of the increasing maturation of the field, researchers have also started exploring the knowledge infrastructure of methodological research on PLS-SEM by analyzing the structures of authors, countries, and co-citation networks (Angelelli, Ciavolino, Ringle, Sarstedt, & Aria, 2025; Ciavolino, Aria, Cheah, & Roldán, 2022; Hwang, Sarstedt, Cheah, & Ringle, 2020; Khan et al., 2019).

TABLE 1.2 ■ Review Articles on PLS-SEM Use

Discipline	References
Accounting	Lee, Petter, Fayard, and Robinson (2011), <i>International Journal of Accounting Information Systems</i>
	Nitzi (2016), <i>Journal of Accounting Literature</i>
Business research (general)	Petter and Hadavi (2023), <i>Partial Least Squares Path Modeling: Basic Concepts, Methodological Issues and Applications</i>
	Vaithilingam, Ong, Moisescu, and Nair (2024), <i>Journal of Business Research</i>
	Widaryanti, Abdullah, Sitawati, and Luhglatno. (2025), <i>Quantity & Quality</i>
	Batra (2025), <i>Engineering, Construction and Architectural Management</i>
Engineering and construction management	de Oliveira Moraes and Watanuki (2025), <i>Gestão & Produção</i>
	Purwanto and Sudargini (2021), <i>Journal of Industrial Engineering & Management Research</i>
	Zeng, Liu, Gong, Hertogh, and König (2021), <i>Frontiers of Engineering Management</i>

(Continued)

TABLE 1.2 ■ (Continued)

Entrepreneurship	Manley, Hair, Williams, and McDowell (2021), <i>The TQM Journal</i>
Family business	Sarstedt, Ringle, Smith, Reams, and Hair (2014), <i>Journal of Family Business Strategy</i>
Higher education and e-learning	Ghasemy, Teeroovengadum, Becker, and Ringle (2020), <i>Higher Education</i>
	Lin et al. (2020), <i>British Journal of Educational Technology</i>
Hospitality and tourism	Ali, Rasoolimanesh, Sarstedt, Ringle, and Ryu (2018), <i>International Journal of Contemporary Hospitality Management</i>
	do Valle and Assaker (2016), <i>Journal of Travel Research</i>
	Mostafiz, Islam, and Pahlevan Sharif (2019), <i>Quantitative Tourism Research in Asia: Current Status and Future Directions</i>
	Usakli and Kucukergin (2018), <i>International Journal of Contemporary Hospitality Management</i>
Human resource management	Legate, Hair, Chretien, and Risher (2023), <i>Human Resource Development Quarterly</i>
	Ringle, Sarstedt, Mitchell, and Gudergan (2020), <i>International Journal of Human Resource Management</i>
International business research	Henseler, Ringle, and Sinkovics (2009), <i>Advances in International Marketing</i>
	Richter, Hauff, Ringle, and Gudergan (2022), <i>Management International Review</i>
	Richter, Sinkovics, Ringle, and Schlägel (2016), <i>International Marketing Review</i>
Knowledge management	Cepeda-Carrión, Cegarra-Navarro, and Cillo (2019), <i>Journal of Knowledge Management</i>
Management information systems	Hair, Hollingsworth, Randolph, and Chong (2017), <i>Industrial Management & Data Systems</i>
	Kante and Michel (2023), <i>Computers in Human Behavior Reports</i>
	Ringle, Sarstedt, and Straub (2012), <i>MIS Quarterly</i>
	Sabol, Hair, Cepeda-Carrión, Roldán, and Chong (2023), <i>Industrial Management & Data Systems</i>

(Continued)

TABLE 1.2 ■ (Continued)

Marketing	Hair, Sarstedt, Ringle, and Mena (2012), <i>Journal of the Academy of Marketing Science</i>
	Guenther, Guenther, Ringle, Zaefarian, and Cartwright (2023), <i>Industrial Marketing Management</i>
	Sarstedt, Hair, et al. (2022), <i>Psychology & Marketing</i>
Operations research	Bayonne, Marin-Garcia, & Alfalla-Luque, (2020), <i>Journal of Industrial Engineering and Management</i>
	Peng and Lai (2012), <i>Journal of Operations Management</i>
Psychology	Willaby, Costa, Burns, MacCann, and Roberts (2015), <i>Personality and Individual Differences</i>
Quality management	Barcia, Garcia-Castro, and Abad-Mora (2022), <i>Sustainability</i>
	Magno, Cassia, and Ringle (2024a), <i>The TQM Journal</i>
Software engineering	Russo and Stol (2021), <i>ACM Computing Surveys</i>
Strategic management	Hair, Sarstedt, Pieper, and Ringle (2012), <i>Long Range Planning</i>
	Hulland (1999), <i>Strategic Management Journal</i>
	Shela, Ramayah, Aravindan, Ahmad, and Alzahrani (2023), <i>Heliyon</i>
Supply chain management	Abbas, Salem, Akram, and Abbas (2025), <i>Operations Research Forum</i>
	Kaufmann and Gaeckler (2015), <i>Journal of Purchasing and Supply Management</i>
	Wang, Cheah, Wong, and Ramayah (2024), <i>International Journal of Logistics Management</i>

Until the first edition of this book, published in 2014, there was no comprehensive textbook that explained the fundamental aspects of the PLS-SEM method, particularly in a way that can be comprehended by the nonstatistician. In recent years, a growing number of follow-up textbooks (e.g., Garson, 2016; Hair, Sarstedt, Ringle, & Gudergan, 2024; Henseler, 2020; Piaw, 2024; Ramayah, Cheah, Chuah, Ting, & Memon, 2018; Wong, 2019), several of which have been translated into various languages, and also edited books on the method (e.g., Avkiran & Ringle, 2018; Esposito Vinzi, Chin, Henseler, &

Wang, 2010; Latan & Noonan, 2017; Latan, Hair, & Noonan, 2023; Radomir et al., 2023) have been published, which helped further popularize PLS-SEM. Special issues dedicated to PLS-SEM in diverse disciplines—such as behavioral research methods (Sarstedt & Hwang, 2020), business research (Gudergan, Moisescu, Radomir, Ringle, & Sarstedt, 2025; Hair, Cheah, Ringle, Sarstedt, & Ting, 2021), internet research (Shiau, Sarstedt, & Hair, 2019), information systems research (Hwang, Ringle, & Sarstedt, 2021; Shiau, Pu, Ray, & Chen, 2020), logistics management (Cheah, Kersten, Ringle, & Wallenburg, 2023), marketing (Cheah & Hair, 2025; Sarstedt & Liu, 2024), and total quality management (Magno, Cassia & Ringle, 2024b)—have contributed to its establishment as a standard method for multivariate data analysis in scientific research.

CONSIDERATIONS IN USING STRUCTURAL EQUATION MODELING

Depending on the underlying research question and the empirical data available, researchers must select an appropriate multivariate analysis method. Regardless of whether a researcher is using first- or second-generation multivariate analysis methods, several considerations are necessary in deciding to use multivariate analysis, particularly SEM. Among the most important are the following five elements: (1) composite variables, (2) measurement, (3) measurement scales, (4) coding, and (5) data distributions.

Composite Variables

A **composite variable** (also referred to as a **variate**) is a linear combination of several variables that are chosen based on the research problem at hand (Hair, Black, Babin, & Anderson, 2019). The process for combining the variables involves calculating a set of weights, multiplying the weights (e.g., w_1 and w_2) with the associated data observations for the variables (e.g., x_1 and x_2), and summing them. The mathematical formula for this linear combination with five variables is shown as follows (note that the composite value can be calculated for any number of variables):

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

where x stands for the individual variables and w represents the weights. All x variables (e.g., questions in a questionnaire) have responses from many respondents that can be arranged in a data matrix. Table 1.3 shows such a data matrix, where i is an index that stands for the number of responses (i.e., cases). A composite value is calculated for each of the i respondents in the sample.

TABLE 1.3 ■ Data Matrix

Case	x_1	x_2	...	x_5	Composite Value
1	x_{11}	x_{21}	...	x_{51}	v_1
...	...	...	...	...	...
i	x_{1i}	x_{2i}	...	x_{5i}	v_i

Measurement

Measurement is a fundamental concept in conducting research. When we think of **measurement**, the first thing that comes to mind is often a ruler, which could be used to measure someone's height or the length of a piece of furniture. But there are many other examples of measurement in life. When you drive, you use a speedometer to measure the speed of your vehicle, a heat gauge to measure the temperature of the engine, and a gauge to determine how much fuel remains in your tank. If you are sick, you use a thermometer to measure your temperature, and when you go on a diet, you measure your weight on a bathroom scale.

Measurement is the process of assigning numbers to a variable based on a set of rules (Hair, Page, Brunsveld, Merkle, & Cleton, 2023). The rules are used to assign the numbers to the variable in a way that accurately represents the variable. With some variables, the rules are easy to follow, whereas with other variables, the rules are much more difficult to apply. For example, if the variable is purchase decision, then it is easy to assign a 1 for yes and a 0 for no. Similarly, if the variable is price, it is again easy to assign a number. But what if the variable is satisfaction or trust? Measurement in these situations is much more difficult because the phenomenon that is supposed to be measured is abstract, complex, and not directly observable. To clarify their measurement, we discuss the concept of **latent variables** or **constructs**.

We cannot directly measure abstract concepts such as satisfaction or trust. However, we can measure indicators of what we have agreed to call satisfaction or trust, for example, as they relate to a brand, product, or company. Specifically, when concepts are difficult to measure, one approach is to measure them indirectly by using a set of directly observable and measurable **indicators** (also called **items** or **manifest variables**). Each indicator represents a single separate aspect of a larger abstract concept. For example, if the concept is restaurant satisfaction, then several indicators that could be used to measure this might be the following:

1. The taste of the food was excellent.
2. The speed of service met my expectations.

3. The waitstaff was knowledgeable about the menu items.
4. The background music in the restaurant was pleasant.
5. The meal was a good value compared with the price.

By combining several indicators to form a scale (or index; Chapter 2), we can indirectly measure the overall concept of restaurant satisfaction. Typically, researchers use several items to form a multi-item scale that indirectly measures a concept, as in the preceding restaurant satisfaction example. The measures may then be combined to form a single **composite score** (i.e., the score of the variate). In some instances, the composite score is a simple summation of several measures (i.e., sum score). In other instances, the scores of the individual measures are combined to form a composite score by using a linear weighting process. The logic of using several individual variables to measure an abstract concept such as restaurant satisfaction is that the measure will be more accurate. The anticipated improved accuracy is based on the assumption that using several items to measure a single concept is more likely to represent all the different aspects of that concept. The accuracy is improved if the items are individually weighted instead of equally weighted. By using multiple variables, researchers seek to reduce **measurement error**, which is the difference between the true value and the value obtained by a measurement. There are many sources of measurement error, including poorly worded questions in a survey, misunderstanding of the scaling approach, and incorrect application of a statistical method. Indeed, all measurements used in multivariate analysis are likely to contain some degree of measurement error. The objective, therefore, is to reduce the measurement error as much as possible.

Rather than using multiple items, researchers sometimes opt for the use of **single-item constructs** to measure concepts such as satisfaction or purchase intention. For example, we may use only “Overall, I’m satisfied with this restaurant” to measure restaurant satisfaction instead of all five items described. Although this is an efficient way to make the questionnaire shorter, it also reduces the quality of your measurement. We discuss the fundamentals of measurement and measurement evaluation in the following chapters.

Measurement Scales

A **measurement scale** is a tool with a predetermined number of closed-ended responses that can be used to obtain an answer to a question. There are four types of measurement scales, each representing a different level of measurement—nominal, ordinal, interval, and ratio. Nominal scales are the lowest level of scales because they are the most restrictive in terms of the type of analysis that can be carried out. A **nominal scale** assigns numbers that can be used to identify and classify objects (e.g., people, companies, products) and is also referred to as a **categorical scale**. For example, if a survey asked a respondent to identify their

profession and the categories are doctor, lawyer, teacher, engineer, and so forth, the question has a nominal scale. Nominal scales can have two or more categories, but each category must be mutually exclusive, and all possible categories must be included. A number would be assigned to identify each category, and the numbers could be used to count the responses in each category, or the modal response or percentage in each category.

If we have a variable measured on an **ordinal scale**, we know that if the value of that variable increases or decreases, this gives meaningful information. For example, if we code customers' use of a product as nonuser = 0, light user = 1, and heavy user = 2, we know that if the value of the variable increases, the level of use also increases. Therefore, when an attribute or characteristic is measured on an ordinal scale, the values provide information about the order of our observations. However, we cannot assume the differences in the order are equally spaced. That is, we do not know if the difference between "nonuser" and "light user" is the same as between "light user" and "heavy user," even though the differences in the values (i.e., 0–1 and 1–2) are equal. Therefore, it is not appropriate to calculate arithmetic means or variances for ordinal data.

If an attribute or characteristic is measured on an **interval scale**, we have precise information on the rank order at which something is measured, and in addition, we can interpret the magnitude of the differences in values directly. For example, if the temperature is 80°F, we know that if it drops to 75°F, the difference is exactly 5°F. This difference of 5°F is the same as the increase from 80°F to 85°F. This exact "spacing" is called **equidistance**, and equidistant scales are necessary for certain analysis techniques. What the interval scale does not include is an absolute zero point. If the temperature is 0°F, it may feel cold, but the temperature can drop further. The value of 0 therefore does not mean that there is no temperature at all (Sarstedt & Mooi, 2019; Chapter 3). The value of interval scales is that almost any type of mathematical computations can be carried out, including the mean and standard deviation. Moreover, you can convert and extend interval scales to alternative interval scales. For example, instead of degrees Fahrenheit (°F), many countries use degrees Celsius (°C) to measure the temperature. Whereas 0°C marks the freezing point, 100°C depicts the boiling point of water. You can convert temperature from Fahrenheit into Celsius by using the following equation: Degrees Celsius (°C) = (degrees Fahrenheit (°F) – 32) · 5/9. In a similar way, you can convert data (via rescaling) on a scale from 1 to 5 into data on a scale from 0 to 100: $\left(\left[\text{data point on the scale from 1 to 5}\right] - 1\right) / (5 - 1) * 100$.

A **ratio scale** provides the most information and is considered the highest level of measurement. If something is measured on a ratio scale, we know that a value of 0 means that a particular characteristic for a variable is not present. For example, if a customer buys no products (value = 0), then they really buy no products. Or, if we spend no money on advertising a new product (value = 0), we spend no money. Therefore, the zero point or origin of the variable is equal

to 0. The measurement of length, mass, and volume as well as time elapsed uses ratio scales. With ratio scales, all types of mathematical computations are possible.

Coding

The assignment of numbers to categories in a manner that facilitates measurement is referred to as **coding**. In survey research, data are often precoded. Precoding is assigning numbers ahead of time to answers (e.g., scale points) that are specified on a questionnaire. For example, a 10-point agree–disagree scale typically would assign the number 10 to the highest endpoint “agree” and a 1 to the lowest endpoint “disagree,” and the points between would be coded 2–9. Postcoding is assigning numbers to categories of responses after data are collected. The responses might be to an open-ended question in a quantitative survey or to an interview response in a qualitative study.

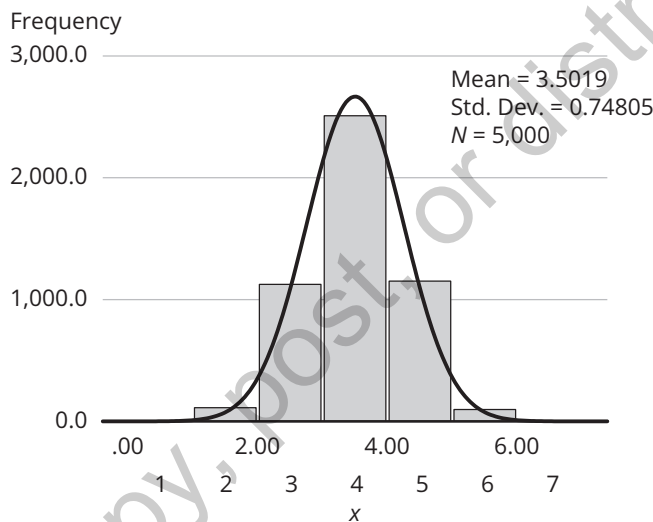
Coding is important in the application of multivariate analysis because it determines when and how types of scales can be used. For example, variables measured on interval and ratio scales can always be used with multivariate analysis. However, when using ordinal scales such as Likert scales (which is common within an SEM context), researchers have to pay special attention to the coding to fulfill the requirement of equidistance. For example, when using a typical 7-point Likert scale with the categories (1) *fully disagree*, (2) *disagree*, (3) *somewhat disagree*, (4) *neither agree nor disagree*, (5) *somewhat agree*, (6) *agree*, and (7) *fully agree*, the inference is that the “distance” between categories 1 and 2 is the same as between categories 3 and 4. In contrast, the same type of Likert scale but using the categories (1) *fully disagree*, (2) *disagree*, (3) *neither agree nor disagree*, (4) *somewhat agree*, (5) *agree*, (6) *strongly agree*, and (7) *fully agree* is unlikely to be equidistant because there are only two categories below the neutral category “neither agree nor disagree,” whereas four categories score above the neutral category. This would clearly bias any result in favor of a better outcome. A suitable Likert scale, as in our first example, will present symmetry of Likert items about a middle category that have clearly defined linguistic qualifiers for each category. In such symmetric scaling, equidistant attributes will typically be more clearly observed or, at least, inferred. When a Likert scale is perceived as symmetric and equidistant, it will behave more like an interval scale. So, whereas a Likert scale is ordinal, if it is well presented, then it is likely the Likert scale can approximate an interval-level measurement, and the corresponding variables can be used in SEM.

Data Distributions

When researchers collect quantitative data, the answers to the questions asked are reported as a distribution across the available (predefined) response categories.

For example, if responses are requested using a 7-point agree–disagree scale, then a distribution of the answers in each of the possible response categories (1, 2, 3, . . . , 7) can be calculated and displayed in a table or chart. Figure 1.1 shows an example of the frequencies of a corresponding variable x . As can be seen, most respondents indicated a 4 on the 7-point scale, followed by 3 and 5, and finally (barely visible), 1 and 7. Overall, the frequency count approximately follows a bell-shaped, symmetric curve around the mean value of 4. This bell-shaped curve is the normal distribution, which many statistical methods require for their analyses.

FIGURE 1.1 ■ Distribution of Responses



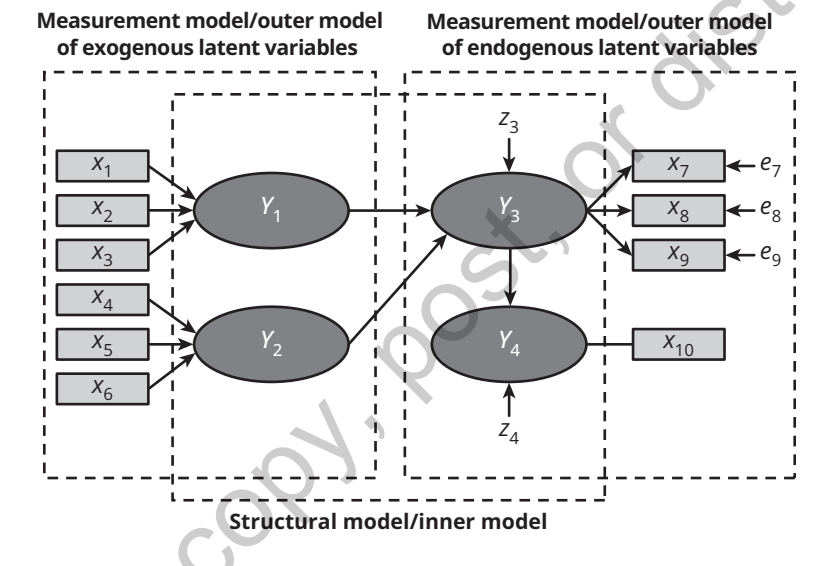
Although many different types of distributions exist (e.g., normal, binomial, Poisson), researchers working with SEM generally only need to distinguish normal from nonnormal distributions. Normal distributions are usually desirable, especially when working with CB-SEM. In contrast, PLS-SEM generally makes no assumptions about the data distributions. However, for reasons discussed in later chapters, it is worthwhile to consider the distribution when working with PLS-SEM. To assess whether the data follow a normal distribution, researchers can apply statistical analyses such as the Anderson–Darling test and the Crámer-von Mises test (e.g., Linnet, 1988; Yap & Sim, 2011). In addition, researchers can examine two measures of distributions—skewness and kurtosis (Chapter 2)—which facilitate assessing to what extent the data deviate from normality (Hair, Black, Babin, & Anderson, 2019).

PRINCIPLES OF STRUCTURAL EQUATION MODELING

Path Models With Latent Variables

Path models are diagrams used to visually display the hypotheses and variable relationships that are examined when SEM is applied (Hair, Page, Brunsveld, Merkle & Cleton, 2023; Hair, Ringle, & Sarstedt, 2011; Sarstedt, Hair, & Ringle, 2023). An example of a **path model** is shown in Figure 1.2.

FIGURE 1.2 ■ A Simple Path Model



Constructs (i.e., variables that are not directly measured) are represented in path models as circles or ovals (Y_1 – Y_4). The indicators, also called items or manifest variables, are the directly measured variables that contain the raw data. They are represented in path models as rectangles (x_1 – x_{10}). Relationships among multiple constructs as well as between constructs and their assigned indicators are shown as arrows. In PLS-SEM, the arrows are always single-headed, thus representing assumed directional relationships.

A PLS path model consists of two elements. First, there is a **structural model** (also called the **inner model** in the context of PLS-SEM) that comprises the constructs (circles or ovals). The structural model also displays the relationships (paths) between the constructs. Depending on their position in the structural

model, constructs can be exogenous or endogenous. **Exogenous constructs** are those that explain other constructs in the model, whereas **endogenous constructs** are those that are being explained in the model. Second, each construct has a **measurement model** (also referred to as the **outer model** in PLS-SEM) that displays the relationships between the construct and its indicator variables (rectangles). For example, x_1 – x_3 are the indicators used in the measurement model of Y_1 , whereas Y_4 has only the x_{10} indicator in the measurement model.

The **error terms** (e.g., e_7 or e_8 ; Figure 1.2) are connected to the (endogenous) constructs and (reflectively) measured variables by single-headed arrows. Error terms represent the unexplained variance when path models are estimated (i.e., the difference between the model's in-sample prediction of a value and an observed value of a manifest or latent variable). In Figure 1.2, error terms e_7 – e_9 are on those indicators whose relationships point from the construct (Y_3) to the indicator (i.e., reflectively measured constructs).

In contrast, the formative indicators x_1 – x_6 , where the relationships go from the indicators to the constructs (Y_1 and Y_2), do not have error terms (Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). Finally, for the single-item construct Y_4 , the direction of the relationships between the construct and the indicator is not relevant because the construct and item are equivalent. For the same reason, there is no error term connected to x_{10} . The structural model also contains error terms. In Figure 1.2, z_3 and z_4 are associated with the endogenous latent variables Y_3 and Y_4 (note that error terms on constructs and measured variables are labeled differently). In contrast, the exogenous latent variables (Y_1 and Y_2) that explain only other latent variables in the structural model do not have an error term, regardless of whether they are specified reflectively or formatively.

Testing Theoretical Relationships

Path models are developed based on theory and are often used to test theoretical relationships. A **theory** is a set of systematically related hypotheses developed following the scientific method that can be used to explain and predict outcomes. Thus, hypotheses are individual conjectures, whereas theories are multiple hypotheses that are logically linked together and can be tested empirically. Two types of theory are required to develop path models: measurement theory and structural theory. **Measurement theory** specifies which indicators are used to measure a certain theoretical concept and how the indicators are used. In contrast, **structural theory** specifies how the constructs are related to each other in the structural model.

Testing theory using PLS-SEM follows a two-step process (Hair, Black, Babin, & Anderson, 2019). We first test the measurement theory to confirm the reliability and validity of the measurement models. After the measurement models' quality is confirmed, we move on to testing the structural theory. The logic is we must first confirm the measurement theory before testing the structural theory because structural theory cannot be confirmed if the measures are unreliable or invalid.

Measurement Theory

Measurement theory specifies how the latent variables (constructs) are measured. Generally, there are two different ways to measure unobservable variables. One approach is referred to as reflective measurement, and the other is formative measurement (e.g., Bollen & Diamantopoulos 2017; Diamantopoulos & Winklhofer, 2001; Rigdon et al., 2014; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). Constructs Y_1 and Y_2 in Figure 1.2 are modeled based on a **formative measurement model**. In this measurement model type, the directional arrows are pointing from the indicator variables (x_1 – x_3 for Y_1 and x_4 – x_6 for Y_2) to the construct, suggesting they jointly form the construct. In contrast, Y_3 in the figure is modeled based on a **reflective measurement model**. With reflective indicators, the direction of the arrows is from the construct to the indicator variables, suggesting that the construct causes the measurement or, more precisely, the covariation of the indicator variables. As indicated in Figure 1.2, reflective measures have an error term associated with each indicator, which is not the case with formative measures. The latter are assumed to be error free (Diamantopoulos, 2006). Finally, note that Y_4 is measured using a single item rather than multi-item measures. Therefore, the relationship between construct and indicator is undirected.

Deciding whether to measure the constructs reflectively versus formatively and whether to use multiple items or a single-item measure are fundamental when developing path models. We therefore explain these two approaches to modeling constructs as well as their variations in more detail in Chapter 2.

Structural Theory

Structural theory defines how the latent variables are related to each other (i.e., it shows the constructs and their path relationships in the structural model). The location and sequence of the constructs are based on theory, or the researcher's experience and accumulated knowledge, or both. When path models are developed, the sequence is typically from left to right. The variables on the left side of the path model are independent variables, and any variable on the right is the dependent variable. Moreover, variables on the left are shown as sequentially preceding and predicting the variables on the right. However, when variables are in the middle of the path model (between the variables that serve only as independent or dependent variables— Y_3) they may also serve as both independent and dependent variables in the structural model.

When latent variables serve only as independent variables, they are called exogenous latent variables (Y_1 and Y_2). When latent variables serve only as dependent variables (Y_4) or as both independent and dependent variables (Y_3), they are called endogenous latent variables. Any latent variable that has only a single-headed arrow going out of it is an exogenous latent variable. In contrast, endogenous latent variables can have either single-headed arrows going both into and out of them (Y_3) or only going into them (Y_4). Note that the exogenous

latent variables Y_1 and Y_2 do not have error terms because these constructs are the entities (independent variables) that are explaining the dependent variables in the path model.

PLS-SEM, CB-SEM, AND REGRESSIONS BASED ON SUM SCORES

There are two main approaches to estimating the parameters of a structural equation model. One approach is CB-SEM (Jöreskog, 1978, 1982), as explained, for instance, in the textbooks by Hair, Black, Babin, and Anderson (2019; see also Hair, Babin, Ringle, Sarstedt, & Becker, 2025, on the use of CB-SEM in SmartPLS) and Kline (2023). The other approach is PLS-SEM (Hair, Ringle, & Sarstedt, 2011; Sarstedt, Ringle, Hair, 2023), on which this book focuses. Each approach is appropriate for a different research context, and researchers need to understand the differences to apply the correct method (Guenther, Guenther, Ringle, Zaefarian, & Cartwright, 2023; Marcoulides & Chin, 2013; Rigdon, Sarstedt, & Ringle, 2017). Note that some researchers have argued for using regressions based on sum scores instead of some type of indicator weighting as executed by PLS-SEM. The sum scores approach offers practically no added value over the PLS-SEM weighted approach (e.g., Hair, Sharma, Sarstedt, Ringle, & Liengaard (2024). For this reason, in the following, we only briefly discuss sum scores in Box 1.1 and instead focus on the PLS-SEM and CB-SEM methods.

A crucial conceptual difference between PLS-SEM and CB-SEM relates to the way each method treats the latent variables included in the model. CB-SEM represents a **common factor-based SEM** method, which assumes the data stem from a common factor model population where the indicator covariances define the nature of the data (e.g., Cook & Foran, 2023; Marcoulides, Chin, & Saunders, 2012; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). PLS-SEM, on the other hand, assumes the data stem from a composite model population, where linear combinations of the indicators define the data's nature (e.g., Cho & Choi, 2020; Hair, Hult, Ringle, Sarstedt, & Thiele, 2017; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). PLS is therefore regarded as a **composite-based SEM** method (Hwang, Sarstedt, Cheah, & Ringle, 2020) because the model estimation involves combining the indicators based on a linear method to form composite variables (Chapter 3).

In this context, it is important to differentiate between the model estimation perspective (i.e., the way the method computes proxies of the conceptual variables using either common factor- or composite-based SEM) and the measurement-theoretic perspective (i.e., deciding whether to specify a construct reflectively or formatively). Deciding on the epistemic (measurement) nature between a construct and its indicators is first and foremost a conceptual decision grounded in measurement theory (Chapter 2). This decision is, however,

independent from the way researchers estimate a model (Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). The choice of an estimator rests, therefore, on assumptions regarding the data generation process. This data generation process can be composite based but yet produce highly correlated variables (e.g., Hair, Hult, Ringle, Sarstedt, & Thiele, 2017). As Guenther, Guenther, Ringle, Zaefarian, and Cartwright (2025, p. 2) noted, “Indicators in a measurement model might be conceptually (and empirically) highly correlated, giving rise to a reflective specification, even if the mental model underlying the respondents’ answering behavior might follow a composite logic.” This is because attitudes—the primary concept of interest in construct measurement—are the result of constructive processes (Zaller & Feldman, 1992), supporting the notion of a composite data generation logic. Considering this, a reflex-like equating of common factor models with reflective measurement and composite models with formative measurement is inappropriate (Guenther, Guenther, Ringle, Zaefarian, & Cartwright, 2025).

Researchers have long noted that the composites produced by the PLS-SEM method are approximations of the conceptual variables they purport to measure

Box 1.1 ■ Sum Scores Approach

In a **sum scores** approach to measurement, each indicator contributes equally to forming the composite (e.g., Hair & Sarstedt, 2019; Henseler et al., 2014). Referring to our descriptions of composite variables at the beginning of this chapter, this would imply that all indicator weights w are set to 1. As noted earlier, the resulting mathematical formula for a linear combination with five variables would be as follows:

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

For example, if a respondent has the scores 4, 5, 4, 6, and 7 on the five variables, the corresponding composite value would be 26. Although easy to apply, regressions using sum scores equalize any differences in the individual item weights. Such differences are, however, common in research reality, and ignoring them conceals measurement model issues (e.g., indicators with low reliability). Simulation studies have shown that using sum scores may entail substantial biases in the parameter estimates and reduced levels of statistical power (e.g., Hair, Hult, Ringle, Sarstedt, & Thiele, 2017) and has adverse consequences for the model’s predictive power (Hair, Sharma, Sarstedt, Ringle, & Liengaard, 2024).

Apart from these statistical concerns, learning about individual item weights offers important insights because the researcher learns about each item’s importance for forming the composite in a certain context (i.e., its relationships with other composites in the structural model). When measuring customer satisfaction, for example, the researcher learns which aspects covered by the individual items are of particular importance for the shaping of satisfaction (Hair, Sharma, Sarstedt, Ringle, & Liengaard, 2024).

(e.g., Rigdon, 2012). As a consequence, some scholars incorrectly view CB-SEM as a more precise method to empirically measure theoretical concepts (e.g., Rönkkö, McIntosh, & Antonakis, 2015), whereas PLS-SEM provides approximations. Other scholars contend, however, that such a view is shortsighted because common factors derived in CB-SEM are also not necessarily equivalent to the theoretical concepts that are the focus of research (e.g., Rigdon, 2012; Rigdon & Sarstedt, 2022; Rossiter, 2011). Indeed, Rigdon, Becker, and Sarstedt (2019a) confirmed that common factor models can be subject to considerable degrees of metrological uncertainty. **Metrological uncertainty** refers to the dispersion of the measurement values that can be attributed to the object or concept being measured (JCGM/WG1, 2008). Rigdon, Becker, and Sarstedt (2019a) have shown that researchers' common practice to restrict the number of indicators per construct to improve model fit in CB-SEM analyses (Hair, Matthews, Matthews, & Sarstedt, 2017) leads to a substantial increase in uncertainty—with adverse consequences for the validity of the results. However, metrological uncertainty can come in many other forms, including ambiguous conceptual definitions, coverage error in sampling, and measurement scale design issues (Rigdon & Sarstedt, 2022; Rigdon, Sarstedt, & Becker, 2020; Sarstedt, 2026). Similarly, researchers' degrees of freedom in research projects—even when estimating the same model using the same data—induce considerable uncertainty and, as a consequence, can result in variability in applications of SEM (Sarstedt, Adler, et al., 2024).

These issues do not imply that composite models are superior, but they cast considerable doubt on the assumption of some researchers that CB-SEM constitutes the gold standard when measuring theoretical concepts. In fact, researchers in various fields of science show increasing appreciation that common factors (relying only on common variance and not including specific variance) may not always be the right approach to measure theoretical concepts (e.g., Rhemtulla, van Bork, & Borsboom, 2020; Rigdon, 2016). Similarly, Rigdon, Becker, and Sarstedt (2019b) demonstrated that using sum scores can significantly increase the degree of metrological uncertainty, which casts additional doubt on this measurement practice (see Box 1.1).

Apart from differences in the philosophy of measurement, the differing treatment of latent variables and, more specifically, the availability of construct scores also has consequences for the methods' areas of application. Specifically, although it is possible to estimate construct scores within a CB-SEM framework, these estimated scores are not unique. That is, an infinite number of different sets of latent variable scores that will fit the model equally well are possible (Guttman, 1955). A crucial consequence of this **factor (score) indeterminacy** is that the correlations between a common factor and any variable outside the factor model—including the theoretical concept, which the constructs purports to represent—are themselves indeterminate (Rigdon, Becker, & Sarstedt, 2019a). That is, they may be high or low, depending on which set of factor scores one chooses.

This limitation has adverse consequences for using CB-SEM for predictive purposes (e.g., Hair & Sarstedt, 2021a; Dijkstra, 2014). In contrast, a major advantage of PLS-SEM is that it always produces a single specific (i.e., determinate) composite score for each case once the weights are established. These determinate scores are proxies of the theoretical concepts being measured, just as factors are proxies for the concepts in CB-SEM (Rigdon, Sarstedt, & Ringle, 2017; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). Using these proxies (determinate scores) as input, PLS-SEM applies ordinary least squares regression with the objective of minimizing the error terms (i.e., the residual variance) of the endogenous constructs in the structural model and indicators in measurement models. In other words, PLS-SEM estimates coefficients such that the amount of explained variance in the endogenous constructs (i.e., the **R^2 values**) and indicators is being maximized (see also Hair, Sarstedt, Ringle, Sharma, & Liengard, 2024).

This feature of maximizing the amount of explained variance achieves the (in-sample) prediction objective of PLS-SEM, which is therefore the preferred method when the research objective is testing a theoretical model while maintaining a predictive focus. In sum, because of the way it estimates model parameters, PLS-SEM is regarded a **variance-based SEM** approach. Specifically, the logic of the PLS-SEM approach is that all of the indicators' variance should be used to estimate the model relationships with particular focus on prediction of the dependent variables (e.g., McDonald, 1996). In contrast, CB-SEM divides the total variance into three types—common, unique, and error variance—but utilizes only common variance (i.e., the variance shared with other indicators in the same measurement model) in the model estimation (Hair, Black, Babin, & Anderson, 2019). That is, CB-SEM explains only the covariance between the indicators (Jöreskog, 1973) rather than considering the indicator's entire variance.

Note that PLS-SEM is similar but not equivalent to **PLS regression**, another popular multivariate data analysis technique (Abdi, 2010; Wold, Sjöström, & Eriksson, 2001). PLS regression is a regression-based approach that explores the linear relationships among multiple independent variables and a single or multiple dependent variable(s). PLS regression differs from regular regression, however, because in developing the regression model, it derives composites from the multiple independent variables by means of principal component analysis. PLS-SEM, in contrast, relies on prespecified networks of relationships among constructs as well as between constructs and their measures (see Mateos-Aparicio, 2011, for a more detailed comparison between PLS-SEM and PLS regression).

In addition, it is important to distinguish PLS-SEM from **path analysis** (e.g., Lleras, 2005; Pedhazur, 1997), which was developed by Wright (1921). The path analysis does not use constructs but examines the relationships between observed variables. It serves as an extension of multiple regression analysis in that it considers more than one dependent variables. The path analysis thereby facilitates evaluating both direct and indirect effects, including mediation, moderation, and moderated-mediation (Hayes, 2022; Sarstedt, Hair, Nitzl, Ringle, & Howard,

2020). In light of its nature, the path analysis can be considered a special case of the PLS-SEM (i.e., a structural model that consists of only observed variables and their relationships).

CONSIDERATIONS WHEN APPLYING PLS-SEM

Key Characteristics of the PLS-SEM Method

Several considerations are important when deciding whether to apply PLS-SEM (Ringle, Sarstedt, Sinkovics, & Sinkovics, 2023). These considerations have their roots in the method's characteristics, which have a bearing on the data and model used. Moreover, the properties of the PLS-SEM method affect the evaluation of the results. In light of these considerations, there are four critical issues relevant to the application of PLS-SEM (Hair, Ringle, & Sarstedt, 2011; Hair, Risher, Sarstedt, & Ringle, 2019): (1) data characteristics, (2) model characteristics, (3) model estimation, and (4) model evaluation. Table 1.4 summarizes the key characteristics of the PLS-SEM method. An initial overview of these issues is provided in this chapter, and a more detailed explanations are provided in later chapters of the book, particularly as they relate to the PLS-SEM algorithm and evaluation of results.

TABLE 1.4 ■ Key Characteristics of PLS-SEM	
Data Characteristics	
Sample size	• Neglectable identification issues with small sample sizes • Achieves high levels of statistical power with small sample sizes
Distribution	• No distributional assumptions—PLS-SEM is a nonparametric method • Influential outliers and collinearity possible influencing the results
Scale of measurement	• Works with metric data and quasi-metric (ordinal) scaled variables • The standard PLS-SEM algorithm accommodates binary coded variables but with additional considerations required when they are used as control variables, moderators, and in the analysis of data from discrete choice experiments

(Continued)

TABLE 1.4 ■ (Continued)

Model Characteristics	
Number of items in each construct's measurement model	Handles constructs measured with single- and multi-item measures
Relationships between constructs and their indicators	Easily incorporates reflective and formative measurement models
Model complexity	Handles complex models with many structural model relationships
Model setup	No causal loops (no circular relationships) allowed in the structural model
Model Estimation	
Objective	Aims at maximizing the amount of explained variance in the endogenous constructs in the structural model and indicators in the measurement models
Efficiency	Converges after a few iterations (even in situations with complex models and/or large sets of data) to the optimum solution (i.e., the algorithm is efficient)
Nature of constructs	Viewed as proxies of the latent concepts under investigation, represented by composites of indicators
Construct scores	- Estimated as linear combinations of their indicators (i.e., they are determinate) - Used for predictive purposes - Can be used as input for subsequent analyses
Parameter estimates	- Structural model relationships generally underestimated and measurement model relationships generally overestimated when solutions are obtained using data from common factor models - Unbiased and consistent when estimating data from composite models - High levels of statistical power compared to alternative methods such as CB-SEM

(Continued)

TABLE 1.4 ■ (Continued)

Model Evaluation	
Evaluation of the overall model	• The concept of fit—as defined in CB-SEM—not fully applicable to PLS-SEM. Efforts to introduce model fit measures have proven difficult
Evaluation of the measurement models	• Consideration of fundamental psychometric properties prior to the data analysis (content and expert validity, unidimensionality) • Reflective measurement models assessed on the grounds of indicator reliability, internal consistency reliability, convergent validity, and discriminant validity • Formative measurement models assessed on the grounds of convergent validity, indicator collinearity, and the significance and relevance of indicator weights
Evaluation of the structural model	• Collinearity among sets of predictor constructs • Significance and relevance of path coefficients • Criteria to assess the model's in-sample (i.e., explanatory) power and out-of-sample predictive power
Additional analyses	• Methodological research substantially extending the original PLS-SEM method by introducing advanced modeling, assessment, and analysis procedures—some examples include the following: ◦ Confirmatory tetrad analysis ◦ Discrete choice modeling ◦ Endogeneity assessment ◦ Higher-order constructs ◦ Importance-performance map analysis ◦ Latent class analysis ◦ Measurement model invariance ◦ Mediation analysis ◦ Moderating effects ◦ Multigroup analysis ◦ Necessary condition analysis ◦ Nonlinear effects

Source: Adapted and extended from Hair, J. F., Ringle, C. M., & Sarstedt, M. (2011). PLS-SEM: Indeed a silver bullet. *Journal of Marketing Theory and Practice*, 19(2), 139–151. Copyright © 2011 by M. E. Sharpe, Inc. Reprinted with permission of the publisher (Taylor & Francis Ltd., <http://www.tandfonline.com>).

PLS-SEM works efficiently with small sample sizes and complex models (Casel, Hackl, & Westlund, 1999; Chin, 2010; Chin, & Newsted, 1999). In addition, different from maximum likelihood-based CB-SEM, which requires normally distributed data, PLS-SEM makes no distributional assumptions (i.e., it is non-parametric). PLS-SEM can easily handle reflective and formative measurement models, as well as single-item constructs, with no identification problems. It can therefore be applied in a wide variety of research situations. When applying PLS-SEM, researchers also benefit from high efficiency in parameter estimation, which is manifested in the method's greater **statistical power** compared to CB-SEM. Greater statistical power means PLS-SEM is more likely to render a specific relationship significant when it is in fact present in the population. The same holds for the comparison with regression based on sum scores, which compared to PLS-SEM, has lower statistical power (Hair, Hult, Ringle, Sarstedt, & Thiele, 2017).

There are, however, several limitations of PLS-SEM. In its basic form, the technique cannot be applied when structural models contain causal loops or circular relationships among the latent variables (i.e., nonrecursive models). Early extensions of the basic PLS-SEM algorithm that have not yet been implemented in standard PLS-SEM software packages do, however, enable handling of circular relationships (Lohmöller, 1989). Furthermore, model fit testing in a PLS-SEM framework is still in its infancy. Although a few researchers have promoted adopting goodness-of-fit measures as commonly used in a CB-SEM context (e.g., Schubert, Henseler, & Dijkstra, 2018; Schubert, Rademaker, & Henseler, 2023), others have raised meaningful concerns about their applicability and general performance (e.g., Hair, Howard, & Nitzl, 2020; Hair, Sarstedt, & Ringle, 2019)—see Chapter 6 for details.

In contrast, PLS-SEM-based model estimation and assessment follows a causal-predictive paradigm, where the aim is to test the predictive power within the confinements of a model carefully developed on the grounds of theory and logic. The underlying causal-predictive logic follows what Gregor (2006) referred to as **explaining and predicting (EP) theories**. EP theories imply an understanding of the underlying causes and prediction as well as descriptions of theoretical constructs and their relationships. According to Gregor (2006, p. 626), this type of theory “corresponds to commonly held views of theory in both the natural and social sciences.”

Numerous seminal theories and models such as Oliver's (1980) expectation-disconfirmation theory or the technology acceptance models (e.g., Davis, 1989; Venkatesh, Morris, Davis, & Davis, 2003) follow an EP-theoretic approach in that they aim to explain *and* predict. Drawing on this notion, Sharma et al. (2024) proposed the **EP-mixed framework**, which distinguishes among four theorizing modes: confirmatory, exploratory, confirmatory-first mixed, and exploratory-first mixed. Within each of these modes, the EP-mixed framework provides clear prescriptions for analytical methods, guidelines, and metrics to assess both explanatory power and predictive accuracy. Furthermore, the framework suggests how these modes may be utilized to create new research programs or extend existing

ones focusing on replicable EP theories. As an additional important contribution, Sharma et al.'s (2024) EP-mixed framework positions predictive modeling at the forefront of theory building and clarifies how it can be integrated with explanatory modeling (Hofman, Sharma, & Watts, 2017; Hofman et al., 2021).

PLS-SEM is perfectly suited to investigate models derived within the EP-mixed framework. More specifically, the method strikes a balance between machine learning methods, which are fully predictive in nature, and CB-SEM, which focuses on confirmation and model fit (Richter, Cepeda-Carrión, Roldán, & Ringle, 2016). The causal-predictive nature of PLS-SEM, therefore, makes the method particularly appealing for research in fields that focus on deriving recommendations for practice. For example, recommendations in managerial implications sections that populate business research journals are always proposed in the form of predictive statements (“our results suggest managers should . . .”). Making such statements requires a prediction focus in model estimation and evaluation. PLS-SEM perfectly emphasizes this need because the method sheds light on the mechanisms (i.e., the structural model relationships) through which the predictions were generated (Hair, 2020; Hair & Sarstedt, 2019, 2021b; Sarstedt & Danks, 2022).

In early writing, researchers noted that PLS estimation is “deliberately approximate” to factor-based SEM (Hui & Wold 1982, p. 127), a characteristic that has been incorrectly referred to as the **PLS-SEM bias** (e.g., Chin, Marcoulin, & Newsted, 2003). A few studies have used simulations to demonstrate this supposed “bias” (e.g., Goodhue, Lewis, & Thompson, 2012; McDonald, 1996; Rönkkö & Evermann, 2013), which manifests itself in measurement model estimates that are higher, while structural model estimates are lower compared to the prespecified values. The studies have concluded that parameter estimates will approach what has been labeled the “true” parameter values when both the number of indicators per construct and sample size increase (Hui & Wold, 1982). Importantly, all these simulation studies used CB-SEM as the benchmark against which the PLS-SEM estimates were evaluated with the assumption that they should be the same. But because PLS-SEM assumes a composite model population, some bias can be expected in such an assessment (Lohmöller, 1989; Schlittgen, Sarstedt, & Ringle, 2020; Schneeweiß, 1991). Not surprisingly, the same issues apply when CB-SEM is used on data generated from composite model populations. In fact, Sarstedt, Hair, Ringle, Thiele, and Gudergan (2016) have shown that the biases produced by CB-SEM are far more severe than those of PLS-SEM when applying the method to the wrong type of model (i.e., estimating composite models with CB-SEM vs. estimating common factor models with PLS-SEM).

Building on the notion of the PLS-SEM bias, Rönkkö, Lee, Evermann, McIntosh, and Antonakis (2023) set out to show that the method does not reliably detect misspecification issues, such as when indicators are assigned to the wrong measurement model. This is only the case, however, when grounding the measurement model assessment in a selective set of (partially outdated) measures. When using the common set of model evaluation metrics, as discussed in Chapters 4–6,

no such issues emerge (Hair, Sarstedt, 2024c). Moreover, acknowledging the different nature of common factor and composite models, most of the issues voiced by critics of the PLS-SEM method (Rönkkö McIntosh, Antonakis, & Edwards, 2016) are not in fact an issue (Cook & Forzani, 2023; Henseler et al., 2014).

Apart from these conceptual concerns, simulation studies show that the differences between PLS-SEM and CB-SEM estimates when assuming the latter as a standard of comparison are small provided the measurement models meet minimum recommended standards in terms of measurement quality (i.e., reliability and validity). Specifically, when the measurement models have four or more indicators and the indicator loadings meet the common standards (≥ 0.708 ; see Chapter 4), there is practically no difference between the two methods in terms of parameter accuracy (e.g., Reinartz, Haenlein, & Henseler, 2009; Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). Thus, the extensively discussed PLS-SEM bias is of no practical relevance for the vast majority of applications (e.g., Binz Astrachan, Patel, & Wanzenried, 2014).

Finally, methodological research has substantially extended the original PLS-SEM method by introducing advanced modeling, assessment, and analysis procedures (see Table 1 in Gudergan et al., 2025). Examples include different types of robustness checks (Vaithilingam, Ong, Moisescu, & Nair, 2024), discrete choice modeling (Hair, Ringle, et al., 2019), necessary condition analysis (Richter, Schubring, Hauff, Ringle, & Sarstedt, 2020) and its extensions (Hauff, Richter, Sarstedt, & Ringle, 2024), out-of-sample prediction metrics (Hair, 2020), endogeneity assessment (Hult et al., 2018; Liengaard et al., 2025), and higher-order constructs (Sarstedt, Hair, Cheah, Becker, & Ringle, 2019). Chapter 8 of this book as well as Hair, Sarstedt, Ringle, and Gudergan (2024) provide an introduction to several of these advanced issues. In the following sections, we discuss aspects related to data characteristics (e.g., minimum sample size requirements) and model characteristics (e.g., model complexity).

Data Characteristics

Data characteristics, such as **minimum sample size requirements**, nonnormal data, and measurement scales (i.e., the use of different scale types), are among the most often stated reasons for applying PLS-SEM across numerous disciplines (e.g., Ghasemy, Teeroovengadam, Becker, & Ringle, 2020; Ringle, Sarstedt, Mitchell, & Gudergan, 2020; Sarstedt, Hair, et al., 2022). Although some of the arguments are consistent with the method's capabilities, others are not. In the following sections, we discuss these and related data characteristics.

Minimum Sample Size Requirements

Small sample size is probably the most often misused argument for using PLS-SEM, with some researchers considering unacceptably low sample sizes (Goodhue, Lewis, & Thompson, 2012; Marcoulides & Saunders, 2006).

These researchers oftentimes believe there is some “magic” in the PLS-SEM approach that enables them to use a small sample (e.g., less than 100) to obtain results representing the effects that exist in a population of several million elements or individuals. No multivariate analysis technique, including PLS-SEM, has this kind of “magic” capability (Petter, 2018).

PLS-SEM can certainly be used with smaller samples, but the population’s nature determines the situations in which small sample sizes are acceptable (Rigdon, 2016). For example, in business-to-business research, populations are often restricted in size. One example would be examining the relationships among teams in organizations, which are typically small. Assuming other situational characteristics are equal, the more heterogeneous the population, the larger the sample size needed to achieve an acceptable sampling error (Cochran, 1977). If basic sampling theory guidelines are not considered (Sarstedt, Bengart, Shaltoni, & Lehmann, 2018), questionable results are produced.

In addition, when applying multivariate analysis techniques, the technical dimension of the sample size becomes relevant. Adhering to the minimum sample size guidelines ensures the results of a statistical method such as PLS-SEM have adequate statistical power. In these regards, an insufficient sample size may not reveal an effect that exists in the underlying population (which results in committing a type II error). Moreover, executing statistical analyses based on minimum sample size guidelines will ensure the results of the statistical method are robust and the model is generalizable to another sample from that same population. Thus, an insufficient sample size may lead to PLS-SEM results that differ from those of another sample. In the following, we focus on the PLS-SEM method and its technical requirements determining the minimum sample size.

The overall complexity of a structural model has little influence on the sample size requirements for PLS-SEM. The reason is that the PLS-SEM algorithm does not compute all relationships in the structural model at the same time. Instead, it uses ordinary least squares regressions to estimate the model’s partial regression relationships. Two early studies systematically evaluated the performance of PLS-SEM with small sample sizes and concluded the method performed well (e.g., Chin & Newsted, 1999; Hui & Wold, 1982). Subsequent simulation studies by, for example, Hair, Hult, Ringle, Sarstedt, and Thiele (2017) and Reinartz, Haenlein, and Henseler (2009), suggest that PLS-SEM is the method of choice when the sample size is small—regardless of whether used on common factor or composite model data. Moreover, compared with its covariance-based counterpart, PLS-SEM has higher levels of statistical power in situations with complex model structures and smaller sample sizes. Similarly, Henseler et al. (2014) have shown that solutions can be obtained with PLS-SEM when other methods such as CB-SEM do not converge or provide inadmissible solutions. For instance, problems are often encountered when using CB-SEM on complex models, especially when the sample size is limited.

Unfortunately, some researchers believe sample size considerations do not play a role in the application of PLS-SEM. This idea has been fostered by the often-cited **10 times rule** (Barclay, Higgins, & Thompson, 1995), which suggests the sample size should be equal to 10 times the number of independent variables in the most complex regression in the PLS path model (i.e., considering both measurement and structural models). This rule of thumb is equivalent to saying the minimum sample size should be 10 times the maximum number of arrowheads pointing at a latent variable anywhere in the PLS path model. Although this rule offers a rough guideline, the minimum sample size requirement should also consider the statistical power of the estimates. To assess statistical power, researchers can interpret power tables (Cohen, 1992) or power analyses using programs such as G*Power (Faul, Erdfelder, Buchner, & Lang, 2009), which is available free of charge at <http://www.gpower.hhu.de/>. These approaches do not explicitly consider the entire model but use the most complex regression in the (formative) measurement models and structural model of a PLS path model as point of reference for assessing the statistical power. In doing so, researchers typically aim at achieving a power level of 80%. However, the minimum sample size resulting from these calculations may still be too small (Kock & Hadaya, 2018).

Addressing these concerns, Kock and Hadaya (2018) proposed the **inverse square root method**, which considers the probability that the ratio of a path coefficient and its standard error will be greater than the critical value of a test statistic for a specific significance level. Therefore, the results for the technically required minimum sample size depend on only one path coefficient and do not depend on the size of the most complex regression in the model. Assuming a common power level of 80% and significance levels of 1%, 5%, and 10%, the minimum sample size ($n_{\min}$) is given by the following equations, where $p_{\min}$ is the value of the path coefficient with the minimum magnitude in the PLS path model, which is expected to be statistically significant:

$$\text{Significance level} = 1\% : n_{\min} > \left(\frac{3.168}{|p_{\min}|} \right)^2$$

$$\text{Significance level} = 5\% : n_{\min} > \left(\frac{2.486}{|p_{\min}|} \right)^2$$

$$\text{Significance level} = 10\% : n_{\min} > \left(\frac{2.123}{|p_{\min}|} \right)^2$$

For example, assuming a significance level of 5% and a minimum path coefficient of 0.2, the minimum sample size is given by

$$n_{\min} > \left(\frac{2.486}{0.2} \right)^2 = 154.505$$

This result needs to be rounded to the next integer, so the minimum sample size is 155.

The inverse square root method is rather conservative in that it slightly overestimates the sample size required to render an effect significant at a given power level. Most importantly, the method stands out due to its ease of use because it can be readily implemented.

Nevertheless, two considerations are important when using the inverse square root method. First, by using the smallest statistical path coefficient as point of reference, the method can be misleading because researchers will not expect marginal effects to be significant. For example, assuming a 5% significance level and a minimum path coefficient of 0.01 would require a sample size of 61,802! Hence, researchers should choose a higher path coefficient as input, depending either on whether the model produces overall weak or strong effects or on the smallest relevant (to be detected) effect.

Second, by relying on model estimates, the inverse square root method follows a retrospective approach. Such an assessment can be used as a basis for additional data collection or adjustments in the model. If possible, however, researchers should follow a prospective approach by trying to derive the minimum expected effect size prior to data analysis. To do so, they can draw on prior research involving a comparable conceptual background or models with similar complexity or, preferably, the results of a pilot study, which tested the hypothesized model using a smaller sample of respondents from the same population. For example, if the pilot study produced a minimum path coefficient of 0.15, this value should be chosen as input for computing the required sample size for the main study. In most cases, however, researchers have only limited information regarding the expected effect sizes even if a pilot study has been conducted. Hence, it is reasonable to consider ranges of effect sizes rather than specific values to determine the sample size required for a specific study.

Table 1.5 shows the minimum sample size requirement for different significance levels and varying ranges of $p_{\min}$. In deriving the minimum sample size, it is reasonable to consider the upper boundary of the effect range as reference because the inverse square root method is rather conservative. For example, when assuming that the minimum path coefficient expected to be significant is between 0.11 and 0.20, one would need approximately 155 observations to render the corresponding effect significant at 5%. Similarly, if the minimum path coefficient expected to be significant is between 0.31 and 0.40, then the recommended sample size would be 39.

TABLE 1.5 ■ Minimum Sample Sizes for Different Levels of Minimum Path Coefficients ($p_{\min}$) and Significance Levels

$p_{\min}$	Significance level		
	1%	5%	10%
0.05–0.1	1,004	619	451
0.11–0.2	251	155	113
0.21–0.3	112	69	51
0.31–0.4	63	39	29
0.41–0.5	41	25	19

Nonnormal Data

The use of PLS-SEM has two other key advantages related to data characteristics (i.e., distribution and scales). In situations where it is difficult or impossible to meet the stricter requirements of more traditional multivariate techniques (e.g., normal data distribution), PLS-SEM is the preferred method. PLS-SEM’s greater flexibility is described by the label “soft modeling,” coined by Wold (1982), who developed the method. It should be noted, however, that “soft” is attributed only to the distributional assumptions and not to the concepts, models, or estimation technique itself (Lohmöller, 1989). PLS-SEM’s statistical properties provide robust model estimations with data that have normal as well as extremely non-normal distributional properties (Hair, Hult, Ringle, Sarstedt, & Thiele, 2017; Reinartz, Haenlein, & Henseler, 2009). It must be remembered, however, that outliers and collinearity do influence the ordinary least squares regressions in PLS-SEM, and researchers should evaluate the data and results for these issues (Hair, Black, Babin, & Anderson, 2019).

Scales of Measurement

The PLS-SEM algorithm generally requires variables to be measured on a **metric scale** (ratio or interval measurement) for the measurement model indicators. But the method also works well with ordinal scales with equidistant data points (i.e., quasi-metric scales; Sarstedt & Mooi, 2019; Chapter 3) and with binary coded data. The use of binary coded data is often a means of including categorical control variables or moderators in path models. In short, binary indicators can be included in models and their PLS-SEM estimation but require special attention (e.g., 0/1 coding). For example, using PLS-SEM in discrete choice experiments where the aim is to predict a binary dependent variable requires specific designs and estimation routines (Hair, Ringle, et al., 2019). Similarly, using binary moderators requires adjustments to adequately account

for standard deviation changes triggered by a change from one category (e.g., 0) to the other (e.g., 1) (Becker, Cheah, Gholamzade, Ringle, & Sarstedt, 2023). In Chapter 7, we briefly discuss the integration, estimation, and interpretation of binary moderators.

Secondary Data

Secondary data are data that have already been gathered, often for a different research purpose and sometimes a long time ago (Sarstedt & Mooi, 2019; Chapter 3). Secondary data are increasingly available to explore real-world phenomena. Research based on secondary data typically focuses on a different objective than in a standard CB-SEM analysis, which is strictly confirmatory in nature. More precisely, secondary data are mainly used in exploratory research to propose causal relationships in situations that have little clearly defined theory (Hair, Risher, Sarstedt, & Ringle, 2019; Hair, Hollingsworth, Randolph, & Chong, 2017). Such settings require researchers to place greater emphasis on examining all possible relationships rather than achieving model fit (Nitzl, 2016). By its nature, this process creates large, complex models that are difficult to analyze using the CB-SEM method. In contrast, due to its less stringent requirements on the data, PLS-SEM offers the flexibility needed for the interplay between theory and data (Nitzl, 2016). Or, as Wold (1982, p. 29) noted, “Soft modeling is primarily designed for research contexts that are simultaneously data-rich and theory-skeletal.” Furthermore, the increasing popularity of secondary data, most notably in the form of user-generated content such as social media posts, videos, and messages (Kübler, Adler, Welke, Sarstedt, & Pauwels, 2025), shifts the research focus from strictly confirmatory to predictive and causal-predictive modeling. Such research settings are an excellent fit for the prediction-oriented PLS-SEM approach (Sharma et al., 2024).

PLS-SEM is also valuable for analyzing secondary data from a measurement theory perspective. First, unlike survey measures, which are usually crafted to confirm a well-developed theory, measures used in secondary data sources are typically not created and refined over time for confirmatory analyses. Thus, achieving model fit is unlikely with secondary data measures in most research situations when using CB-SEM. Second, researchers who use secondary data do not have the opportunity to revise or refine the measurement model to achieve fit. Third, a major advantage of PLS-SEM when using secondary data is that it permits the unrestricted use of single-item and formative measures. This is extremely valuable for research involving secondary data because many measures included in corporate databases are artifacts, such as financial ratios and other firm-fixed factors (Henseler, 2017). Such artifacts are typically reported in the form of formative indices, which can readily be estimated with PLS-SEM. Box 1.2 summarizes key considerations related to data characteristics.

BOX 1.2 ■ Data Considerations When Applying PLS-SEM

- The 10 times rule is not a reliable indication of sample size requirements in PLS-SEM. Statistical power analyses provide a more reliable minimum sample size estimate. Researchers can also draw on the inverse square root method as a more conservative way of assessing minimum sample size requirements.
- When the construct measures meet recommended guidelines in terms of reliability and validity, results from CB-SEM and PLS-SEM are generally similar.
- PLS-SEM can handle extremely nonnormal data (e.g., data with high levels of skewness).
- Due to its flexibility in handling different data and measurement types, PLS-SEM is the method of choice when analyzing secondary data.
- PLS-SEM works with metric, quasi-metric, and categorical (i.e., dummy-coded) scaled data, although there are certain limitations. Processing of data from discrete choice experiments or using binary moderators requires specific designs and estimation routines.

Model Characteristics

PLS-SEM is flexible in its modeling properties. In its basic form, the PLS-SEM algorithm requires all models to be specified without circular relationships or loops of relationships between the latent variables in the structural model. Although causal loops are sometimes specified in business research, this characteristic does not limit the applicability of PLS-SEM as Lohmöller's (1989) extensions of the basic PLS-SEM algorithm allow for handling such model types. Other model specification requirements that constrain the use of CB-SEM, such as when estimating models with formative measures, are generally not relevant with PLS-SEM.

Measurement model issues are one of the major obstacles to obtaining a solution with CB-SEM. For instance, estimation of complex models with many latent variables and/or indicators is often impossible with CB-SEM. In contrast, PLS-SEM can be used in such situations because it is not constrained by identification and other technical issues. Consideration of reflective and formative measurement models is a key issue in the application of SEM (Bollen & Diamantopoulos, 2017). PLS-SEM can easily handle both formative and reflective measurement models and is considered the primary approach when the hypothesized model incorporates formative measures. Note that CB-SEM can accommodate formative indicators, but to ensure model identification, researchers must follow rules that impose specific constraints on the model to ensure model identification (Bollen & Davis, 2009; Diamantopoulos & Riefler, 2011). As Hair, Sarstedt, Ringle, and Mena (2012, p. 420) noted, "These constraints often contradict theoretical considerations, and the question arises whether model design should guide theory or vice versa."

In contrast, PLS-SEM does not have such requirements and handles formative measurement models without any limitation. This also applies to model settings in which endogenous constructs are measured formatively. The applicability of CB-SEM to such model settings has been subject to considerable debate (Cadoğan & Lee, 2013; Rigdon, 2014a), but due to PLS-SEM's multistage estimation process (Chapter 3), which separates measurement from structural model estimation, the inclusion of formatively measured endogenous constructs is not an issue (Rigdon et al., 2014).

Different from CB-SEM, PLS-SEM facilitates easy specification of interaction terms to map moderation effects in a path model. This makes PLS-SEM the method of choice in simple moderation models and more complex conditional process models that combine moderation and mediation effects (Sarstedt, Hair, et al., 2020; Cheah, Nitzl, Roldán, Cepeda-Carrión, & Gudergan, 2021). Similarly, higher-order constructs, which facilitate specifying a construct simultaneously on different levels of abstraction (Sarstedt, Hair, Cheah, Becker, & Ringle, 2019) can readily be implemented in PLS-SEM.

Finally, PLS-SEM is capable of estimating complex models. For example, if theoretical or conceptual assumptions support large models and sufficient data are available (i.e., meeting minimum sample size requirements), PLS-SEM can handle models of almost any size, including those with dozens of constructs and hundreds of indicator variables. As noted by Wold (1985), PLS-SEM is virtually without competition when path models with latent variables are complex in their structural relationships (Chapter 3). Box 1.3 summarizes rules of thumb for path model characteristics.

BOX 1.3 ■ Model Considerations When Choosing PLS-SEM

- PLS-SEM offers substantial flexibility in handling different measurement model setups. For example, PLS-SEM can handle both reflective and formative measurement models as well as single-item measures without additional requirements or constraints.
- The method allows for the specification of advanced model elements such as interaction terms and higher-order constructs.
- Model complexity is generally not an issue for PLS-SEM. As long as appropriate data meet minimum sample size requirements, the complexity of the structural model is virtually unrestricted.

GUIDELINES FOR CHOOSING BETWEEN PLS-SEM AND CB-SEM

To answer the question of when to use PLS-SEM versus CB-SEM, researchers should focus on the characteristics and objectives that distinguish the two methods. The primary distinguishing factor relates to researchers' assumptions about

the nature of the data, which is by definition unknown (Rigdon, Sarstedt, & Ringle, 2017). Whereas CB-SEM assumes the data stem from a common factor model population, PLS-SEM assumes composite model data. Researchers need to decide whether they adhere to the strict principles of common factor models or assume that theoretical concepts can, in principle, be approximated with composites. Note again that this model estimation perspective is unrelated to measurement-theoretic considerations; that is, whether the measures are operationalized reflectively or formatively—see Guenther, Guenther, Ringle, Zaefarian, and Cartwright (2025) for a discussion.

In addition to these model population considerations, the CB-SEM method, with its strong focus on model fit and its extensive data requirements, is primarily suitable for testing a theory within the confinement of a concise theoretical model. However, if the primary research objective is testing the predictive power of a model, whose structure is grounded in theory and logic (i.e., a causal-predictive analysis), PLS is the most appropriate SEM method (Hair, Hollingsworth, Randolph, & Chong, 2017; Hair, Sarstedt, & Ringle, 2019). Finally, CB-SEM emphasizes the explanation of relationships among all variables simultaneously and is not a method that focuses on prediction of dependent variable(s). In contrast, PLS-SEM is causal-predictive in nature. To assess the predictive power of their models, researchers using PLS-SEM can draw on the PLS_{predict} procedure and the cross-validated predictive ability test (Shmueli, Sarstedt, Hair, Cheah, Ting, Vaithilingam, & Ringle, 2019; Sharma, Liengaard, Hair, Sarstedt, & Ringle, 2023).

Summarizing the previous discussions and drawing on Hair, Risher, Sarstedt, and Ringle (2019), Box 1.4 displays the relevant rules of thumb when deciding whether to use CB-SEM or PLS-SEM. As can be seen, PLS-SEM is not recommended as a universal alternative to CB-SEM. Both methods differ from a statistical point of view, are designed to achieve different objectives, require different assumptions, and rely on diverse philosophies of measurement. Neither technique is generally superior to the other, and neither of them is appropriate for all situations (Hair, Sarstedt, Ringle, Sharma, & Liengaard, 2024; Peter, 2018; Sharma, Liengaard, Sarstedt, Hair, & Ringle, 2023). In general, the strengths of PLS-SEM are CB-SEM's limitations and vice versa. It is important for researchers, therefore, to understand the different applications each approach was developed for—and to apply them accordingly. Researchers need to choose the SEM technique that best suits their research objectives, data characteristics, and model setup (e.g., Guenther, Guenther, Ringle, Zaefarian, & Cartwright, 2023; Hair, Ringle, & Sarstedt, 2012; Rigdon, Sarstedt, & Ringle, 2017; Roldán & Sánchez-Franco, 2012; Sarstedt, Ringle, & Hair, 2014; Sarstedt, Ringle, Henseler, & Hair, 2014).

Researchers may also follow a multimethod approach to SEM (e.g., Guenther, Guenther, Ringle, Zaefarian, & Cartwright, 2025; Hair, Sharma, Chin, Sarstedt, Ringle, 2026; Sarstedt, Adler, et al., 2024; Sharma, Sarstedt, Ringle, Cheah, Herfurth, & Hair, 2024) that utilizes the specific methodological

strengths of each type (e.g., the demonstration of model fit on the one hand and predictive power of the model on the other) to verify the robustness of the model estimates across different methods. To reiterate, researchers should view CB-SEM and PLS-SEM not as competing but as complementary.

BOX 1.4 ■ Rules of Thumb for Choosing Between PLS-SEM and CB-SEM

Use PLS-SEM when

- composite models can be assumed to adequately approximate the theoretical concepts of interest;
- the analysis is concerned with testing a theoretical framework from a prediction perspective;
- the structural model is complex and includes many constructs, indicators, and/or model relationships;
- the research objective is to better understand increasing complexity by exploring theoretical extensions of established theories (exploratory research for theory development);
- the path model includes one or more formatively measured constructs;
- the research consists of financial ratios or similar types of artifacts;
- the research is based on secondary data, which may lack a comprehensive substantiation on the grounds of measurement theory;
- a small population restricts the sample size (e.g., business-to-business research), but PLS-SEM also works well with large sample sizes;
- distribution issues are a concern, such as a lack of normality; or
- the research requires latent variable scores for follow-up analyses.

Use CB-SEM when

- common factor models can be assumed to adequately approximate the theoretical concepts of interest;
- the goal is strictly theory testing and confirmation;
- error terms require additional specification, such as the covariation;
- the structural model has circular relationships; or
- the research requires a global goodness-of-fit criterion.

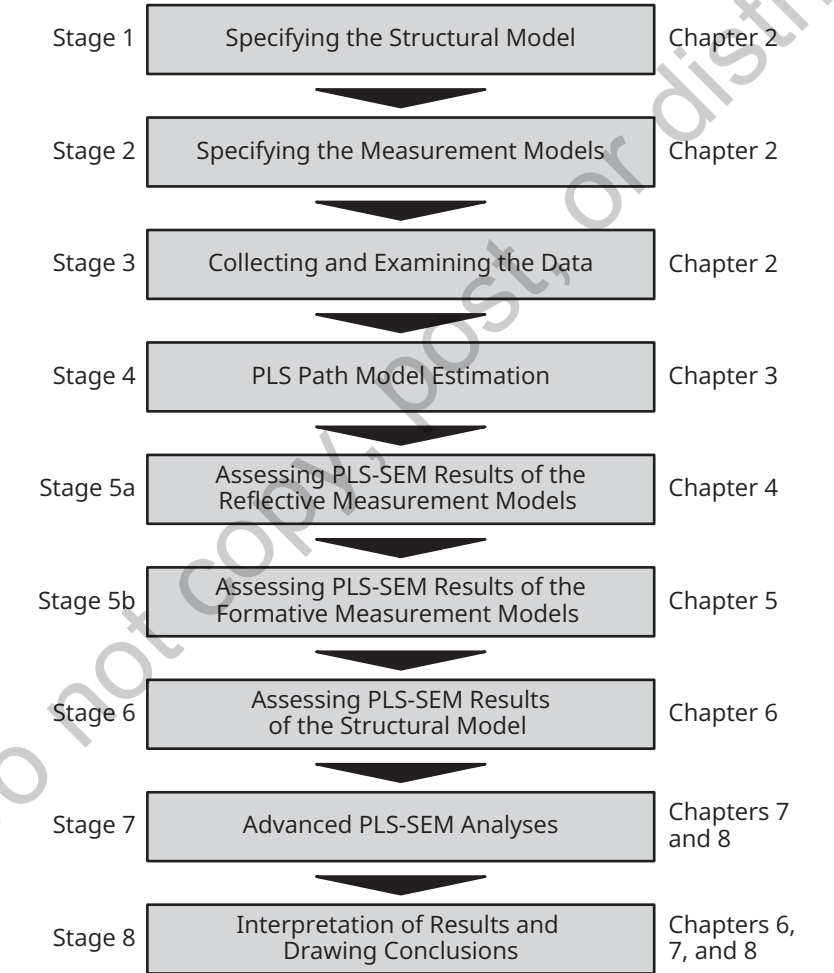
Source: Adapted from Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24.

Copyright © 2019 by Emerald Publishing. Reprinted with permission of the publisher (Emerald Publishing; <https://www.emeraldgroupublishing.com>).

ORGANIZATION OF REMAINING CHAPTERS

The remaining chapters provide more detailed information on PLS-SEM, including specific examples of how to use software to estimate simple and complex PLS path models. In doing so, the chapters follow a multistage procedure that should be used as a blueprint when conducting PLS-SEM analyses (Figure 1.3). Ideally, these and all related steps of the research process should be preregistered (Soliman et al., 2025).

FIGURE 1.3 ■ A Systematic Procedure for Applying PLS-SEM



Adler, Sharma, and Radomir (2023) introduced a preregistration template specifically tailored to PLS-SEM studies.

Specifically, the eight-stage process starts with the specification of structural and measurement models, followed by the examination of data (Chapter 2). Next, we discuss the PLS-SEM algorithm and provide an overview of important considerations when running the analyses (Chapter 3). On the basis of the results of the computation, researchers then evaluate the results. To do so, researchers must know how to assess both reflective and formative measurement models (Chapters 4 and 5). When the data for the measures are considered reliable and valid (based on established criteria), researchers can then evaluate the structural model (Chapter 6). Chapter 7 covers models with mediating and moderating effects whose analyses have become standard in PLS-SEM research. On the basis of the results of Chapters 6 and 7, researchers interpret their findings and draw their final conclusions. Finally, Chapter 8 offers a brief overview of advanced techniques.

Summary

- Understand the meaning of structural equation modeling (SEM) and its relationship to multivariate data analysis.** SEM is a second-generation multivariate data analysis method that facilitates analyzing the relationships among multiple endogenous and exogenous constructs, each measured by one or more indicator variables. The primary advantage of SEM is its ability to measure complex model relationships while accounting for measurement error inherent in the indicators. There are two types of SEM methods—CB-SEM and PLS-SEM. The two method types differ in the way they estimate the model parameters and their assumptions regarding the nature of data. Whereas CB-SEM considers the constructs as common factors, PLS-SEM considers the constructs as composites linearly formed by sets of indicator variables.
- Be able to describe the basic elements of a structural equation model.** Several considerations are necessary when applying multivariate analysis, including the following five: (1) composite variables, (2) measurement, (3) measurement scales, (4) coding, and (5) data distributions. A composite variable (also called a variate) is a linear combination of several indicators chosen based on the research problem at hand. Measurement is the process of assigning numbers to a variable based on a set of rules. Multivariate measurement involves using several variables to indirectly measure a concept to improve measurement accuracy. The anticipated improved accuracy is based on the assumption that using several variables (indicators) to measure a single concept is more likely to represent all the different aspects of the concept and thereby result in a more valid measurement of the concept. The ability to identify measurement error

(Continued)

(Continued)

using multivariate measurement also helps researchers obtain more accurate measurements and thereby more accurate results. Measurement error is the difference between the true value of a variable and the value obtained by a measurement. A measurement scale is a tool with a predetermined number of closed-ended responses that can be used to obtain an answer to a question. There are four types of measurement scales: nominal, ordinal, interval, and ratio. When researchers collect quantitative data using scales, the answers to the questions can be shown as a distribution across the available (predefined) response categories. The type of distribution must always be considered when working with SEM.

- **Comprehend the basic concepts of partial least squares structural equation modeling (PLS-SEM).** Path models are diagrams used to visually display the hypotheses and variable relationships that are examined when structural equation modeling is applied. Four basic elements must be understood when developing path models: (1) constructs, (2) measured variables, (3) relationships, and (4) error terms. Constructs (or latent variables) measure theoretical concepts that are not directly observable. They are represented in path models as circles or ovals. Measured variables are directly measured observations (raw data), generally referred to as either indicators or manifest variables, and are represented in path models as rectangles. Relationships represent hypotheses in path models and are shown as arrows that are single-headed, indicating a causal-predictive relationship between the constructs. These relationships are derived from structural theory and logic. Depending on their role in the model, constructs are either exogenous or endogenous. Error terms represent the unexplained variance when path models are estimated and are present for endogenous constructs and reflectively measured indicators. Exogenous constructs and formative indicators do not have error terms. Measurement theory specifies how the constructs (latent variables) are measured. Latent variables can be specified as either reflective or formative.
- **Gain the knowledge to explain the differences between covariance-based structural equation modeling (CB-SEM) and PLS-SEM and when to use each.** Compared to CB-SEM, PLS-SEM emphasizes prediction while simultaneously relaxing the assumptions regarding the data and specification of relationships. The statistical objective of PLS-SEM is maximizing the explained variance of the endogenous constructs in the structural model and the indicators in the measurement models by estimating partial model relationships in an iterative sequence of ordinary least squares regressions. In contrast, CB-SEM estimates model parameters such that the discrepancy between the estimated and sample covariance matrices is minimized. Instead of following a common factor model logic, as CB-SEM does, PLS-SEM calculates composites of indicators that serve as proxies for the theoretical concepts under research. The PLS-SEM method is not constrained by identification issues, even if the model becomes complex—a situation that typically

restricts CB-SEM use—does not require normally distributed data, and is robust when data is skewed (but outliers should be considered for removal). Moreover, PLS-SEM can handle formative measurement models much better and has advantages when sample sizes are relatively small and when analyzing secondary data. Researchers should consider the two SEM approaches as complementary and apply the SEM technique that best suits their research objective, data characteristics, and model setup.

Review Questions

1. What is multivariate analysis?
2. Describe the difference between first- and second-generation multivariate methods.
3. What is SEM?
4. What is the key difference between the common factor model and the composite model?
5. What is the value of SEM in understanding relationships between variables?

Critical Thinking Questions

1. When would SEM methods be more advantageous than first-generation techniques, such as multivariate regression in understanding relationships between variables?
2. Why should social science researchers consider using SEM instead of multiple regression?
3. What are the most important considerations in deciding whether to use CB-SEM or PLS-SEM?
4. Under what circumstances is PLS-SEM the preferred method over CB-SEM?
5. Why is an understanding of theory important when deciding whether to use PLS-SEM or CB-SEM?
6. Why is PLS-SEM's prediction focus a major advantage of the method?

Key Terms

10 times rule	Measurement model
Categorical scale	Measurement scale
Coding	Measurement theory
Common factor-based SEM	Metric scale
Composite-based SEM	Metrological uncertainty
Composite scores	Minimum sample size requirements
Composite variable	Multivariate analysis
Confirmatory	Nominal scale
Constructs	Ordinal scale
Covariance-based structural equation modeling (CB-SEM)	Outer model
Endogenous latent variables	Partial least squares structural equation modeling (PLS-SEM)
EP-mixed framework	Path model
Equidistance	PLS path modeling
Error terms	PLS regression
Exogenous latent variables	PLS-SEM bias
Explaining and predicting (EP) theories	R ² value
Exploratory	Ratio scale
Factor (score) indeterminacy	Reflective measurement model
First-generation techniques	Secondary data
Formative measurement model	Second-generation techniques
Indicators	Single-item constructs
Inner model	Statistical power
Interval scale	Structural equation modeling (SEM)
Inverse square root method	Structural model
Items	Structural theory
Latent variables	Sum scores
Manifest variables	Theory
Measurement	Variance-based SEM
Measurement error	Variate

Suggested Readings

Chin, W. W., Cheah, J.-H., Liu, Y., Ting, H., Lim, X.-J., & Cham, T. H. (2020). Demystifying the role of causal-predictive modeling using partial least squares structural equation modeling in information systems research. *Industrial Management & Data Systems*, 120(12), 2161–2209.

Cook, D. R., & Forzani, L. (2023). On the role of partial least squares in path analysis for the social sciences. *Journal of Business Research*, 167, 114132.

Guenther, P., Guenther, M., Ringle, C. M., Zaefarian, G., & Cartwright, S. (2023). Improving PLS-SEM use for business marketing research. *Industrial Marketing Management*, 111, 127–142.

Guenther, P., Guenther, M., Ringle, C. M., Zaefarian, G., & Cartwright, S. (2025). PLS-SEM and reflective constructs: A response to recent criticism and a constructive path forward. *Industrial Marketing Management*, 128, 1–9.

Hair, J. F., Sarstedt, M., Ringle, C. M., Sharma, N. P., & Liengaard, B. D. (2024). Going beyond the untold facts in PLS-SEM and moving forward. *European Journal of Marketing*, 58(13), 81–106.

Hair, J. F., Sharma, P. N., Sarstedt, M., Ringle, C. M., & Liengaard, B. D. (2024). The shortcomings of equal weights estimation and the composite equivalence index in PLS-SEM. *European Journal of Marketing*, 58(13), 30–55.

Jöreskog, K. G., & Wold, H. (1982). The ML and PLS techniques for modeling with latent variables: Historical and comparative aspects. In K. G. Jöreskog & H. Wold (Eds.), *Systems under indirect observation, Part I* (pp. 263–270). North-Holland.

Lohmöller, J. B. (1989). *Latent variable path modeling with partial least squares*. Physica.

Petter, S. (2018). “Haters gonna hate”: PLS and Information Systems Research. *ACM SIGMIS Database: The DATABASE for Advances in Information Systems*, 49(2), 10–13.

Rigdon, E. E. (2012). Rethinking partial least squares path modeling: In praise of simple methods. *Long Range Planning*, 45(5–6), 341–358.

Rigdon, E. E., Sarstedt, M., & Ringle, C. M. (2017). On comparing results from CB-SEM and PLS-SEM: Five perspectives and five recommendations. *Marketing ZFP*, 39(3), 4–16.

Ringle, C. M., Sarstedt, M., Sinkovics, N., & Sinkovics, R. R. (2023). A perspective on using partial least squares structural equation Modelling in data articles. *Data in Brief*, 48, 109074.

Sarstedt, M., Hair, J. F., Ringle, C. M., Thiele, K. O., & Gudergan, S. P. (2016). Estimation issues with PLS and CBSEM: Where the bias lies! *Journal of Business Research*, 69(10), 3998–4010.

Sharma, P. N., Sarstedt, M., Ringle, C. M., Cheah, J.-H., Herfurth, A., & Hair, J. F. (2024). A framework for enhancing the replicability of behavioral MIS research using prediction oriented techniques. *International Journal of Information Management*, 78, 102805.

Wold, H. (1985). Partial least squares. In S. Kotz & N. L. Johnson (Eds.), *Encyclopedia of statistical sciences* (pp. 581–591). Wiley.

Visit the companion site for this book at <https://www.pls-sem.net/>.

Do not copy, post, or distribute

2

SPECIFYING THE PATH MODEL AND EXAMINING DATA

LEARNING OUTCOMES

1. Understand the basic concepts of structural model specification, including mediation, moderation, nonlinear relationships, and the use of control variables.
2. Explain the differences between reflective and formative measurement models, and specify the appropriate measurement model.
3. Comprehend that the selection of the mode of measurement model and the indicators must be based on theoretical reasoning before data collection.
4. Explain the advantages of using differentiated indicator weights over sum scores.
5. Understand the nature of higher-order constructs.
6. Describe the data collection and examination considerations necessary to apply PLS-SEM.
7. Learn how to develop a PLS path model using the SmartPLS software.

CHAPTER PREVIEW

This chapter introduces the basic concepts of structural and measurement model specification when PLS-SEM is used. The concepts are associated with completing the first three stages in the application of PLS-SEM, as described in Chapter 1. To begin with, Stage 1 is specifying the structural model. Next, Stage 2 is selecting and specifying the measurement models. Stage 3 summarizes the major guidelines for data collection when the application of PLS-SEM is anticipated as well as the need to examine your data after they have been collected to ensure the results from applying PLS-SEM are valid and reliable. An understanding of these three topics will prepare you for Stage 4, estimating the model, which is the focus of Chapter 3.

STAGE 1: SPECIFYING THE STRUCTURAL MODEL

In the initial stages of a research project involving the application of SEM, an important first step is to prepare a diagram illustrating the research hypotheses and visually displaying the variable relationships that will be examined. This diagram is often referred to as a path model. Recall that a path model is a diagram that connects indicators and constructs based on theory and logic to visually display the hypotheses that will be tested (Chapter 1). Preparing a path model early in the research process enables researchers to organize their thoughts and visually consider the relationships among the variables of interest. Path models also are an efficient means of sharing ideas among researchers working on or reviewing a research project.

Path models are made up of two elements: (1) the structural model (also called the **inner model** in PLS-SEM), which describes the relationships between the latent variables, and (2) the measurement model (also called the outer model in PLS-SEM), which describes the relationships between the latent variable and its measures (i.e., its indicators). We discuss structural models first, which are developed in Stage 1. In the next section, we explain Stage 2, measurement models.

When a structural model is being developed, two primary issues need to be considered: the sequence of the constructs and the relationships among them. Both issues are critical to the concept of modeling because they represent the hypotheses and their relationship to the theory being tested.

The sequence of the constructs in a structural model is based on theory, logic, or practical experiences observed by the researcher. The sequence is typically displayed from left to right, with independent (predictor) constructs on the left side and dependent (outcome) variables on the right. That is, constructs on the left are assumed to precede and predict constructs on the right. Constructs that act only as independent variables are generally referred to as exogenous latent

variables, and are on the left side of the structural model. Those independent (i.e., exogenous) latent variables have arrows that only point out of them and never have arrows from other latent variables pointing into them. Constructs considered dependent in a structural model (i.e., those that have an arrow pointing into them from other latent variables) are called **endogenous latent variables** and are on the right side of the structural model. Constructs that operate as both independent and dependent variables in a structural model also are considered endogenous and appear in the middle of the diagram.

The structural model in Figure 2.1 illustrates the three types of constructs and the relationships among them. The reputation construct on the far left is an exogenous (i.e., independent) latent variable. It is modeled as explaining the satisfaction construct. The satisfaction construct is an endogenous latent variable that has a dual relationship as both independent and dependent. It is a dependent construct because it is explained by reputation. But it is also an independent construct because it explains loyalty. The loyalty construct on the right end is an endogenous (i.e., dependent) latent variable explained by satisfaction.

FIGURE 2.1 ■ Example of Path Model and Types of Constructs



Determining the sequence of the constructs is seldom an easy task because contradictory theoretical perspectives can lead to different sequencing of latent variables. For example, some researchers assume customer satisfaction precedes and predicts corporate reputation (e.g., Walsh, Mitchell, Jackson, & Beatty, 2009), whereas others argue corporate reputation predicts customer satisfaction (Damberg, Liu, & Ringle, 2024; Sarstedt, Wilczynski, & Melewar, 2013). Similarly, in the American Customer Satisfaction Index (ACSI) model, Fornell, Johnson, Anderson, Cha, and Bryant (1996; see also Fornell, Morgeson, Hult, & VanAmburg, 2020; Morgeson, Hult, Sharma, & Fornell, 2023) considered perceived value being a predecessor of customer satisfaction, whereas Olsen and Johnson (2003) proposed customer satisfaction as a predecessor of perceived equity in their model. Theory and logic should always determine the sequence of constructs in a structural model, but when the literature is inconsistent or unclear, researchers must use their best judgment to determine the sequence.

Acknowledging there is no single unique model that characterizes a phenomenon well, researchers can also establish and empirically compare theoretically justified alternative models (Burnham & Anderson, 2002). The models selected for comparison should be motivated by logic and theory from relevant fields in

line with PLS-SEM's "causal-predictive" nature (Jöreskog & Wold, 1982, p. 270; see also Wold, 1982). Because PLS path models focus on providing theoretical explanations, considering purely empirically motivated models would be akin to "snooping" and is not recommended for theoretical research that focuses on both explanation and prediction (Gregor, 2006). Establishing alternative models requires leveraging the existing literature to provide valid theoretical rationale for all the models being considered. In particular, you should be able to

- (1) describe the theoretical commonalities among the proposed alternative models (i.e., whether certain proposed effects are common across models),
- (2) contrast the models to highlight the differences in theoretical mechanisms being captured (such differences may manifest as additional/different paths or antecedents), and
- (3) explain why the commonalities and differences are important to consider in terms of the effect on the target variable for the population under study.

We discuss **model comparisons**, especially predictive model comparisons and selection in PLS-SEM (e.g., Lienggaard et al., 2021; Sharma, Shmueli, Sarstedt, Danks, & Ray, 2021), in the context of the structural model evaluation in Chapter 6. Sharma, Sarstedt, Shmueli, Kim, and Thiele (2019) introduced a five-step procedure for model comparison and inference in PLS-SEM. These authors also discuss possible misconceptions related to model comparisons.

Once the sequence of the proposed constructs has been decided, the relationships between them must be established by drawing arrows. The arrows are inserted with the arrowhead pointing to the right. This approach indicates the sequence and that the constructs on the left explain the constructs on the right. Although these relationships are sometimes referred to as **causal links**, researchers should be cautious in making causal claims because the results from any PLS-SEM analysis are correlational only. However, when following Sharma et al.'s (2024) EP-mixed framework—which combines explanatory and predictive perspectives in model design and evaluation—and if the structural theory offers strong support for the assumed relationships, causality can be assumed.

Causal processes are complex and may involve numerous constructs. Mapping all these constructs and their relationships is not only a difficult undertaking but may also have adverse consequences for the model's generalizability. The reason is that a complex model may represent the structure in a specific research situation well, as evidenced in high explanatory power. At the same time, however, the same model is unlikely to generalize to different research settings, as evidenced in low predictive power (Sarstedt & Danks, 2022; Shmueli, 2010). Expressed in more technical terms, the model overfits the data (see Chapter 6 for a discussion). Hence, in drawing arrows between the constructs, researchers always

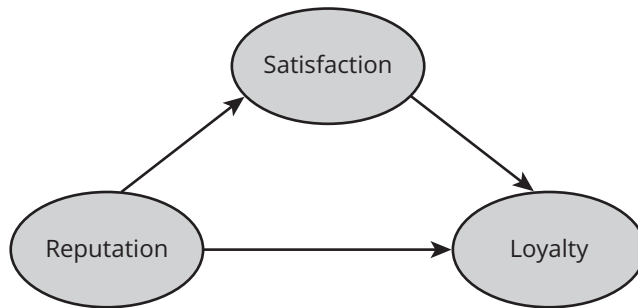
face a trade-off between theoretical soundness (i.e., including those relationships strongly supported by theory) and model parsimony (i.e., using fewer relationships). The latter should be of crucial concern because the most nonrestrictive statement, “everything is predictive of everything else,” is also the most uninformative. As pointed out by Falk and Miller (1992, p. 24), “A parsimonious approach to theoretical specification is far more powerful than the broad application of a shotgun.”

In most instances, researchers examine direct linear relationships between the constructs in a path model. Theory often suggests, however, that model relationships are more complex and may involve mediation, moderation, or nonlinear relationships. In addition, researchers commonly specify control variables that account for some of the variation in the endogenous constructs. In the following section, we briefly introduce these different relationship types. In Chapter 7, we explain how mediation and moderation—arguably the most common types of complex relationships—can be estimated and interpreted using PLS-SEM.

Mediation

A **mediating effect** is created when a third variable or construct intervenes between two other related constructs (Memon, Cheah, Ramayah, Ting, & Chuah, 2018; Nitzl, Roldán, & Cepeda-Carrión, 2016), as shown in Figure 2.2. To understand how mediating effects work, we will consider a path model in terms of direct and indirect effects. A **direct effect** is a relationship that links two constructs with a single arrow. An **indirect effect** is a relationship that involves a sequence of relationships with at least one intervening construct involved. Thus, an indirect effect is a sequence of two or more direct effects (compound path) that are represented visually by multiple arrows. This indirect effect is characterized as the mediating effect. In Figure 2.2, satisfaction is modeled as a possible mediator between reputation and loyalty.

FIGURE 2.2 ■ Example of a Mediating Effect



From a theoretical perspective, the most common application of mediation is to “explain” why an alternative indirect relationship exists between an exogenous and endogenous construct. For example, a researcher may observe a direct relationship between two constructs but may not be sure *why* the relationship exists or if the observed relationship is the only relationship between the two constructs. In such a situation, a researcher might posit an explanation of the relationship in terms of an intervening variable that operates by receiving the “inputs” from an exogenous construct and translating them into an “output” in the form of an endogenous construct. The role of the mediator variable then is to reveal the mechanism through which the independent constructs affect the dependent construct.

Consider the example in Figure 2.2, in which we want to examine the effect of corporate reputation on customer loyalty. On the basis of theory and logic, we know that a relationship exists between reputation and loyalty, but we are unsure how the relationship actually works (Eberl & Schwaiger, 2005; Schwaiger, 2004). As researchers, we might want to explain how companies translate their reputation into higher loyalty among their customers. We may observe that sometimes a customer perceives a company as being highly reputable, but this perception does not translate into high levels of loyalty. In other situations, we observe that some customers with lower corporate reputation assessments are highly loyal. These observations are confusing and lead to the question as to whether there is some other process going on that translates corporate reputation into customer loyalty.

In the diagram, the intervening process (mediating effect) is modeled via the construct satisfaction. If a respondent perceives a company to be highly reputable, this assessment may lead to higher satisfaction levels and ultimately to increased loyalty. In such a case, the relationship between reputation and loyalty may be explained by the reputation → loyalty sequence, or the reputation → satisfaction → loyalty sequence, or perhaps even by both sets of relationships (Figure 2.2). The reputation → loyalty sequence is an example of a direct relationship. In contrast, the reputation → satisfaction → loyalty sequence is an example of an indirect relationship. After empirically testing these relationships, the researcher would be able to explain how reputation is related to loyalty as well as the role that satisfaction might play in mediating that relationship. Chapter 7 offers additional details on mediation and explains how to test mediating effects in PLS-SEM (see also, e.g., Cheah, Nitzl, Roldán, Cepeda-Carrión, & Gudergan, 2021; Hayes, 2022; Sarstedt, Hair, Nitzl, Ringle, & Howard, 2020; Sarstedt & Moisesescu, 2024).

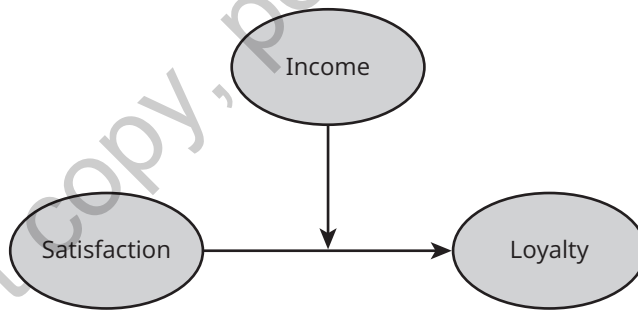
Moderation

Moderation is another important statistical analysis concept. In statistical **moderation**, a third variable directly affects the relationship between the exogenous and endogenous latent variables but in a different way from mediation. Referred to as a **moderator effect**, this situation occurs when the moderator (a variable or construct) changes the strength or even the direction of a relationship

between two constructs in the model (Becker Cheah, Gholamzade, Ringle, & Sarstedt, 2023; Becker, Sarstedt, & Ringle, 2018; Memon et al., 2019). The crucial distinction between moderation and mediation is that the moderator variable does not depend on the exogenous latent variable.

For example, income has been shown to significantly affect the strength of the relationship between customer satisfaction and customer loyalty (Homburg & Giering, 2001). In that context, income serves as a moderator variable on the satisfaction → loyalty relationship, as shown in Figure 2.3. Specifically, the strength of the relationship (as measured by the path coefficient) between satisfaction and loyalty has been shown to be weaker for people with high income than for people with low income. For higher-income individuals, there may be little or no relationship between satisfaction and loyalty. But for lower-income individuals, there often is a quite strong relationship between the two variables. As such, moderation may be understood as a way to account for **heterogeneity** in the theoretical model. Heterogeneity means that different types of effects can be expected for different groups of respondents (e.g., Haenlein & Kaplan, 2011; Ringle, Sarstedt, & Zimmermann, 2011). That is, instead of assuming the relationship between customer satisfaction and customer loyalty is the same for all respondents, we acknowledge that this effect is different for low- and high-income individuals.

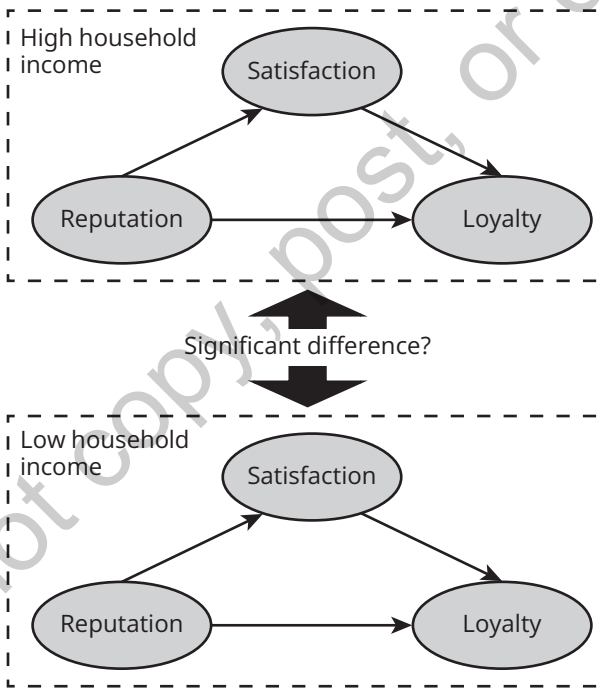
FIGURE 2.3 ■ Theoretical Model of a Continuous Moderating Effect



In the example displayed in Figure 2.3, income may be measured on a continuous scale, for example, using the respondents' annual household income. But a moderator variable can also be measured categorically; for example, high household income is > 160 thousand U.S. dollars a year, and low household income is ≤ 60 thousand U.S. dollars per year. Such a categorical variable may affect one or more direct relationships in the model but may also be used as a grouping variable that divides the data into subsamples. The same theoretical model is then estimated for each of the distinct subsamples,

assuming the categorical moderator variable affects all model relationships. Because researchers are usually interested in comparing the models and learning about significant differences between the subsamples, the model estimates for the subsamples are usually compared by means of **multigroup analysis** (Matthews, 2017). Specifically, multigroup analysis enables the researcher to test for differences between identical models estimated for different groups of respondents. The general objective is to see if there are statistically significant differences between the group-specific path coefficients. For example, we might be interested in evaluating whether the effects among reputation, satisfaction, and loyalty shown in Figure 2.2 are significantly different for high- versus low-income households (Figure 2.4); see also Homburg and Giering (2001).

FIGURE 2.4 ■ Example of a Multigroup Analysis



In Chapter 7, we discuss in greater detail how to use categorical and continuous variables for moderator analysis. Chapter 8 provides a brief overview of multigroup analysis.

Nonlinear Relationships

Standard conceptualizations that underpin cause–effect relationships in applications of PLS-SEM imply that constructs affect one another in a linear manner. In some instances, however, this assumption does not hold when relationships are nonlinear (e.g., Ahrholdt, Gudergan, & Ringle, 2019). When the relationship between two constructs is nonlinear, the size of the effect between two constructs depends not only on the magnitude of change in the exogenous construct but also on its value. For example, Steiner, Siems, Weber, and Guhl (2014) found a nonlinear relationship between quality satisfaction and customer retention. Specifically, the relationship shows decreasing returns to scale for increasing levels of quality satisfaction (i.e., a degressive shape), whereas for high levels of quality satisfaction, returns slightly increase again. Overall, the authors concluded that the relationship between quality satisfaction and customer retention is subject to a saturation effect, implying that customer retention cannot keep on increasing unlimitedly, even when customer satisfaction with regard to quality continues to increase.

To model such nonlinear effects in PLS-SEM, researchers specify higher-order terms of the independent constructs, such as squared, cubic, or higher powers, depending on the number of assumed turning points where the slope of the curve changes its signs. For example, a quadratic function has one turning point (e.g., the slope's negative sign changes to a positive one), whereas a cubic function has two turning points (e.g., with sign changes of the slope from a negative to a positive and again to a negative one). Modeling simple (i.e., quadratic) nonlinear effects is similar to a moderation analysis, where researchers include an auxiliary term (i.e., an interaction term), which captures the interplay between an independent construct and the moderator. In the case of a quadratic effect, this interaction term embodies the interaction of the independent construct with itself, which is why this type of analysis is also referred to as self-moderation. Hair, Sarstedt, Ringle, and Gudergan (2024, Chapter 3) and Basco, Hair, Ringle, and Sarstedt (2022) discuss the specification and estimation of nonlinear effects in PLS-SEM in greater detail. Sarstedt, Ringle, et al. (2020) and Vaithilingam, Ong, Moisesescu, and Nair (2024) have recommended researchers consider nonlinear relationships in the structural model for robustness checks in PLS-SEM. Also, the structural configuration search using the **fuzzy-set qualitative comparative analysis (fsQCA)** implicitly considers nonlinear effect (e.g., Acquah, Quaicoo, & Gatsi, 2024; Rasoolimanesh, Ringle, Sarstedt, & Olya, 2021).

Control Variables

When specifying theoretical models to be tested, researchers sometimes include control variables. The business disciplines of accounting, finance, international business, and management often include control variables in their research. **Control variables** are designed to measure the influence of independent variables

that are not of primary interest to the researcher (e.g., Bernerth & Aguinis, 2016; Klarmann & Feuer, 2018).

By adding control variables to the hypothesized structural model, researchers seek to account for other explanatory factors that potentially influence the endogenous constructs. For example, when estimating the impact of customer satisfaction on stock returns, researchers also need to account for the influence of several firm characteristics such as research and development intensity, marketing investments, and firm size (Raithel, Sarstedt, Scharf, & Schwaiger, 2012). Failure to account for these characteristics could lead to an overestimation of the effect of customer satisfaction on stock returns, potentially triggering a type I error (i.e., false positive).

From a statistical perspective, adding control variables to the model entails that the hypothesized effects are estimated at constant levels of the control variables. If the hypothesized relationships remain largely constant, researchers can rule out alternative explanations related to the control variables. As such, control variables help strengthen the causal inference of the main effects. In addition, adding control variables improves the precision of the model estimates because they reduce the statistical noise in the endogenous construct. This particularly holds when the control variables exhibit low correlations with the antecedents of the endogenous construct (Carlson & Wu, 2012; Memon, Thursamy, Ting, Cheah, & Chuah, 2024). Control variables should be considered only when there are meaningful theoretical or empirical reasons for their inclusion (Shiau, Chau, Thatcher, Teng, & Dwivedi, 2024; Spector & Brannick, 2011). In fact, the cure can be worse than the disease when adding control variables that may act as target variables themselves—such practice can distort results and bias coefficient estimates (Cinelli, Forney, & Pearl, 2024).

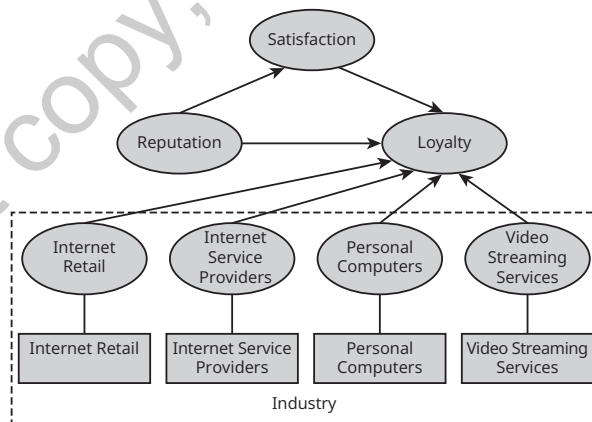
To add control variables into a PLS path model, researchers need to establish a separate exogenous construct for each control variable to be considered and link each new construct to the endogenous latent variable under consideration. For example, suppose a researcher wants to control for the impact of the respondent's age on loyalty when estimating the mediation model shown in Figure 2.2. To do so, the researcher needs to add a new construct into the model, measured with the single age item, and link this construct to the loyalty construct. By adding the age measure as a control, the effects of reputation on loyalty and customer satisfaction on loyalty will decrease, provided that age has an impact on the endogenous construct. The estimated effects of control variables are generally difficult to interpret in a causal-predictive sense (Hünemann & Louw, 2023), which is why researchers should focus on analyzing the effects of the focal constructs.

In some situations, researchers wish to control for the impact of categorical variables such as industry type. If the categorical variable has only two categories such as firm size (small or large), one uses a binary (dummy) variable and includes it as a single-item construct in the PLS path model. In this case,

zero becomes the reference category (e.g., small firms) and the relationship between the control variable and endogenous construct shows the effect of switching from the reference category to the other category (e.g., large firms). When the variable has more than two categories, the categorical variable needs to be recoded into a series of binary (dummy) variables. In PLS-SEM, each category (i.e., represented by a binary variable) is incorporated into the model as a single-item construct. Specifically, when the control variable has k categories, researchers need to create $k-1$ binary variables. The category that is left out is referred to as the reference category. To identify the reference category, all binary variables take the value zero. The values of the dummy variables other than zero (typically one) then denote the deviation from this reference category. When an observation falls into the reference category, which is typically the first category, all binary (dummy) variables are zero. When an observation falls into the second category, the first binary (dummy) variable is one, all others are zero, and so on.

Figure 2.5 illustrates the use of a categorical control variable, *Industry*, which comprises five categories: 1 = *Computer Software*, 2 = *Internet Retail*, 3 = *Internet Service Providers*, 4 = *Personal Computers*, and 5 = *Video Streaming Service*. In PLS-SEM, categorical variables must be recoded into binary (dummy) variables, with one category serving as the reference. In this example, the first category, *Computer Software*, is designated as the reference category and therefore is not

FIGURE 2.5 ■ Categorical Control Variable With Multiple Categories in PLS-SEM



Note: This figure does not show the measurement models and indicators of the constructs *Reputation*, *Satisfaction*, and *Loyalty*. The control variable *Industry* includes dummy-coded indicator variables of industries 2–5, where the first industry (i.e., *Computer Software*) represents the reference category.

included as a dummy indicator in the model. Each of the four remaining industry categories is represented in the model as a single-item construct. As shown in Figure 2.5, the two relationships in the structural model, namely, from *Reputation* to *Loyalty* and from *Satisfaction* to *Loyalty* are controlled by *Industry* (with *Computer Software* as reference category).

In conclusion, researchers are typically interested only in controlling for the impact of the control variables rather than explicitly hypothesizing and testing their impact. As a consequence, the path coefficients quantifying the effect of the control variables on the endogenous construct and their significances are not interpreted. When including control variables, researchers need to offer compelling theoretical arguments as to why these variables are important rather than following a kitchen sink approach, which considers all potential control variables available as meaningful (Spector & Brannick, 2011). Berneth and Aguinis (2016) and Becker (2005) provide best-practice recommendations that can be followed to make decisions on the appropriateness of including a specific control variable within a particular theoretical framework, research domain, and empirical study.

STAGE 2: SPECIFYING THE MEASUREMENT MODELS

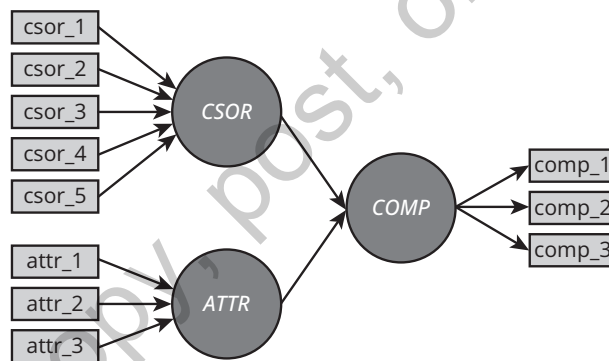
The structural model describes the relationships among latent variables (constructs). In contrast, the measurement models represent the relationships between constructs and their corresponding indicator variables (Sarstedt, Ringle, & Hair, 2025). The basis for determining these relationships is measurement theory. A sound measurement theory is a necessary condition to obtain useful results from any PLS-SEM analysis. Hypothesis tests involving the structural relationships among constructs will be only as reliable and valid as the measurement models are in explaining how these constructs are measured.

Researchers typically have several established measurement approaches to choose from, each a slight variant from the others. In fact, almost all social science researchers today use established measurement approaches published in prior research studies or scale handbooks (e.g., Bearden, Netemeyer, & Haws, 2011; Bruner, 2021; Zarantonella & Pauwels-Delassus, 2015) that performed well in the past (Ramirez, David, & Brusco, 2013). In some situations, however, researchers do not have access to an established measurement approach and must develop a new set of measures (or substantially modify an existing approach). A description of the general process for developing indicators to measure a construct can be long and detailed. Hair, Black, Babin, and Anderson (2019) described the essentials of this process. Likewise, Diamantopoulos and Winklhofer (2001), DeVellis and Thorpe (2021), and MacKenzie, Podsakoff, and Podsakoff (2011) offered thorough explanations of different approaches to measurement development. In

each case, decisions regarding how the researcher selects the indicators to measure a particular construct provide a foundation for the remaining analysis.

The path model shown in Figure 2.6 shows an excerpt of the path model we use as an example throughout this book. The model has two exogenous constructs—*corporate social responsibility* (*CSOR*) and *attractiveness* (*ATTR*)—and one endogenous construct, which is *competence* (*COMP*). Each of these constructs is measured by means of multiple indicators. For instance, the endogenous construct *COMP* has three indicator variables, *comp_1*, *comp_2*, and *comp_3*. Using a scale from 1 to 7 (*fully disagree* to *fully agree*), respondents had to evaluate the following statements: “[The company] is a top competitor in its market,” “As far as I know, [the company] is recognized worldwide,” and “I believe that [the company] performs at a premium level.” The answers to these three inquiries represent the measures for this construct. The construct itself is measured indirectly by these three indicator variables and, for that reason, is referred to as a latent variable.

FIGURE 2.6 ■ Example of a Path Model With Three Constructs



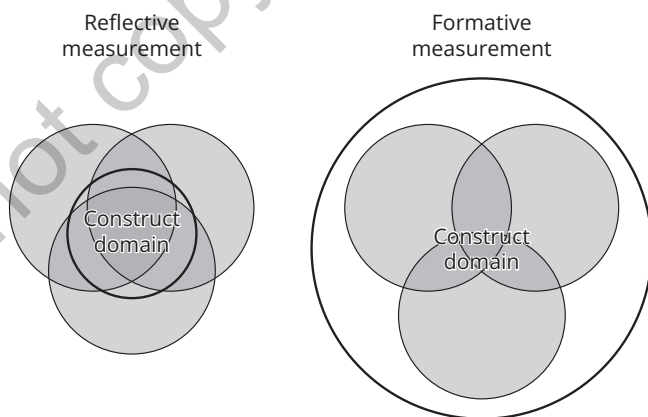
The other two constructs in the model, *CSOR* and *ATTR*, can be described in a similar manner. That is, the two exogenous constructs are measured by indicators that are each directly measured by responses to specific questions. Note that the relationship between the indicators and the corresponding construct is different for *COMP* compared with *ATTR* and *CSOR*. When you examine the *COMP* construct, the direction of the arrows goes from the construct to the indicators. This type of measurement model is referred to as *reflective*. When you examine the *ATTR* and *CSOR* constructs, the direction of the arrows is from the measured indicator variables to the constructs. This type of measurement model is called *formative*. As discussed in Chapter 1, an important characteristic of PLS-SEM is that the technique readily incorporates both reflective and formative measures. Likewise, PLS-SEM can easily be used when constructs are measured with only

a single item (rather than multiple items). Both of these measurement issues are discussed in the following sections.

Reflective and Formative Measurement Models

As indicated in the previous section, researchers must consider two broad types of measurement specification: reflective and formative measurement models. The **reflective measurement** model has a long tradition in the social sciences and is grounded in classical test theory. According to this theory, measures represent the effects (or manifestations) of an underlying construct. Therefore, causality is from the construct to its measures (*COMP* in Figure 2.6). Reflective indicators (sometimes referred to as **effect indicators** in the psychometric literature) can be viewed as a representative sample of all the possible items available within the conceptual domain of the construct (Nunnally & Bernstein, 1994). Therefore, because a reflective measure dictates that all indicator items are “caused” by the same construct (i.e., they stem from the same domain), indicators associated with a particular construct should be highly correlated with each other. In addition, individual items should be interchangeable, and any single item can generally be left out without changing the meaning of the construct as long as the construct has sufficient reliability. The fact that the relationship goes from the construct to its measures implies that if the evaluation of the latent trait changes (e.g., because of a change in the standard of comparison), all indicators will change simultaneously. A set of reflective measures is commonly called a **scale**.

FIGURE 2.7 ■ Conceptual Difference Between Reflective and Formative Measures



Note: The black circle represents the construct domain of interest and the gray-shaded circles the content domain captured by each indicator.

In contrast, **formative measurement** models are based on the assumption that the indicators form the construct by means of linear combinations. Therefore, researchers typically refer to this type of measurement model as being an **index**. An important characteristic of formative indicators is that they are not interchangeable, as is true with reflective indicators. Thus, each indicator for a formative construct captures a specific aspect of the construct's domain. Taken jointly, the items ultimately determine the meaning of the construct, which implies that omitting an indicator potentially alters the nature of the construct. As a result, breadth of coverage of the construct domain is extremely important to ensure the content of the focal construct is adequately captured (Diamantopoulos & Winklhofer, 2001).

Researchers distinguish between two types of indicators in the context of formative measurement: causal and composite indicators. Constructs measured with **causal indicators** have an error term, which implies that the construct has not been perfectly measured by its indicators (Bollen & Bauldry, 2011). More precisely, causal indicators show conceptual unity in that they correspond to the researcher's definition of the concept (Bollen & Diamantopoulos 2017). But researchers will not necessarily consider all indicators relevant for adequately capturing the construct's domain (e.g., Bollen & Lennox, 1991). The error term then captures all the other "causes" or explanations of the construct that the set of causal indicators do not capture (Diamantopoulos, 2006). The existence of an error term in causal indicator models suggests the construct can, in principle, be equivalent to the conceptual variable of interest, provided the model has perfect fit (e.g., Grace & Bollen, 2008). The use of causal indicators is prevalent in CB-SEM, which—at least in principle—facilitates explicitly defining the error term of a formatively measured latent variable. However, the nature of this error term is questionable because its magnitude partly depends on other constructs embedded in the model and their measurement quality (Aguirre-Urreta, Rönkkö, & Marakas, 2016).

Composite indicators constitute the second type of indicators associated with formative measurement models. Constructs measured with composite indicators also have an error term, which however, is set to zero. This implies the indicators define the construct in full (Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016). Composite indicators are frequently used to condense complex and multidimensional data into a simpler, more manageable form that can be easily analyzed and compared. For example, the Human Development Index combines data on life expectancy, education, and per capita income to rank countries by levels of human development. Similarly, a researcher may be interested in measuring the salient aspects of a company's corporate social responsibility using a set of composite indicators that capture the features relevant to the particular study. Finally, companies use composite indicators of employee engagement, satisfaction, and turnover rates to gauge overall workforce health. However, composite indicators can be used to measure any concept including attitudes, perceptions, and behavioral intentions (Nitzl & Chin, 2017), as long

as they operationally define the concept. But composite indicators are not a free ride for careless measurement. Instead, “as with any type of measurement conceptualization, researchers need to offer a clear construct definition and specify items that closely match this definition—that is, they must share conceptual unity” (Sarstedt, Hair, Ringle, Thiele, & Gudergan, 2016, p. 4002). Thus, composite indicator models view construct measurement as approximation of conceptual variables, acknowledging the practical problems that arise with measuring unobservable conceptual variables that populate theoretical models (Rigdon, Sarstedt, & Ringle, 2017).

In a nutshell, the distinction between causal and composite indicators relates to a difference in measurement philosophy. Causal indicators assume that a certain concept can—at least in principle—be fully measured using a set of indicators and an error term, an assumption that has been questioned on analytical and conceptual grounds (Rigdon, Becker, & Sarstedt 2019s). Composite indicators make no such assumption but view measurement explicitly as an approximation of a certain theoretical concept. The inclusion of an error term in causal indicator models appears appealing at first sight but because its magnitude depends on the measurement quality of downstream constructs, the error term’s value for judging the quality of the formative measurement model is ambiguous (Rigdon et al., 2014). In addition, by including an error term in a formative measurement model, CB-SEM treats the formative measurement as if it were a common factor model. PLS-SEM, in contrast, approximates the concepts of interest using composites (i.e., linear combinations of the indicator variables), regardless of whether the measurement models are specified reflectively or formatively.

Considering this, the distinction between causal and composite indicators in measurement appears rather artificial with little consequence for method choice. For the sake of simplicity and in line with seminal research in the field (e.g., Fornell & Bookstein, 1982), we therefore refer to formative indicators when assuming composite indicators (as used in PLS-SEM) in the remainder of this book. Similarly, we refer to formative measurement models to describe measurement models comprising composite indicators. Guenther, Guenther, Ringle, Zaefarian and Cartwright (2025), Rigdon, Sarstedt, and Ringle (2017), and Sarstedt, Hair, Ringle, Thiele, and Gudergan (2016) have provided further information on composite models as well as common factor models and their distinction.

Figure 2.7 illustrates the key difference between the reflective and formative measurement perspectives. The black circle illustrates the construct domain, which is the domain of content the construct is intended to measure. The gray circles represent the content domain that each indicator captures. Whereas the reflective measurement approach aims at maximizing the overlap between interchangeable indicators, the formative measurement approach tries to fully cover the domain of the latent concept under investigation (black circle) by the different formative indicators (gray circles), which should have small overlap.

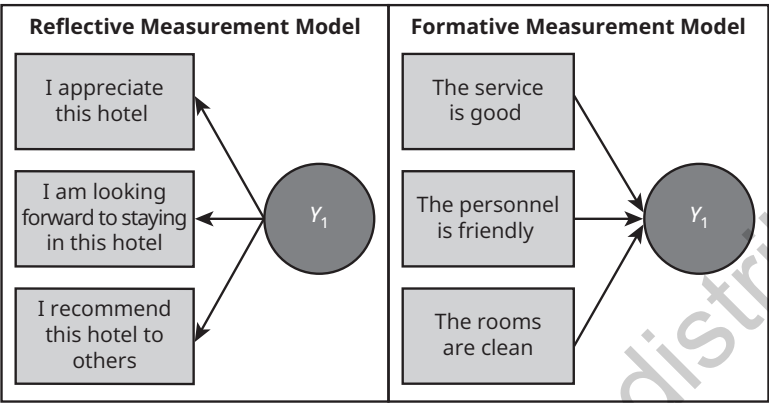
Unlike the reflective measurement approach, whose objective is to maximize the overlap between interchangeable indicators, there are no specific expectations about patterns or the magnitude of intercorrelations among formative indicators (Diamantopoulos, Riefler, & Roth, 2008). Because there is no “common cause” for the items in the construct, there is not any requirement for the items to be correlated, and they may be completely independent. In fact, collinearity among formative indicators can present significant problems because the weights linking the formative indicators with the construct can become unstable and nonsignificant. Furthermore, formative indicators have no individual measurement error terms. That is, they are assumed to be error-free in a conventional sense. These characteristics have broad implications for the evaluation of formatively measured constructs, which rely on a totally different set of criteria compared with the evaluation of reflective indicators (Chapter 5). For example, analyzing the internal consistency reliability of a formatively measured construct could suggest that individual indicators need to be removed because of low inter-item correlations. However, such a step would decrease the content validity of the measurement approach (Diamantopoulos & Siguaw, 2006). Broadly speaking, researchers need to pay closer attention to the content validity of the measures by determining how well the indicators represent the domain (or at least its major aspects) of the latent concept under research (e.g., Bollen & Lennox, 1991).

But when do we measure a construct reflectively or formatively? There is not a definite answer to this question because constructs are not inherently reflective or formative. Instead, the specification depends on the construct conceptualization and the objective of the study. Consider Figure 2.8, which shows how the construct “satisfaction with hotels” (Y_1) can be operationalized in both ways (Albers, 2010).

The left side of Figure 2.8 shows a reflective measurement model setup. This type of model setup is likely to be more appropriate when a researcher wants to test theories with respect to satisfaction. In many managerially oriented business studies, however, the aim is to identify the most important drivers of satisfaction that ultimately lead to customer loyalty. In this case, researchers should consider the different facets of satisfaction, such as satisfaction with the service or the personnel, as shown on the right side of Figure 2.8. In the latter case, a formative measurement model specification is more promising because it allows identifying distinct drivers of satisfaction and thus deriving more nuanced recommendations. This especially applies to situations where the corresponding constructs are exogenous. However, formative measurement models may also be used on endogenous constructs when measurement theory supports such a specification.

Apart from the role a construct plays in the model and the recommendations the researcher wants to give based on the results, the specification of the content of the construct (i.e., the domain content the construct is intended to capture) primarily guides the measurement perspective. Still, the decision as to which measurement model is appropriate has been the subject of considerable debate in a variety of disciplines and is not fully resolved. In Table 2.1, we present

FIGURE 2.8 ■ Satisfaction as a Formatively and Reflectively Measured Construct



Adapted from source: Albers, S. (2010). PLS and success factor studies in marketing. In V. Esposito Vinzi, W. W. Chin, J. Henseler, & H. Wang (Eds.), *Handbook of partial least squares: Concepts, methods and applications in marketing and related fields* (pp. 409–425). Springer.

a set of guidelines that researchers can use in their decision of whether to measure a construct reflectively or formatively. Note that there are also empirical means to determine the measurement perspective. Drawing on Bollen and Ting (1993, 2000), Gudergan, Ringle, Wende, and Will (2008) proposed the **confirmatory tetrad analysis in PLS-SEM (CTA-PLS)**, which allows testing the null hypothesis that the construct measures are reflective in nature (e.g., Troiville, 2024; Wilden, Gudergan, Nielsen, & Lings, 2013). We discuss the CTA-PLS technique, which is, for instance, available in the SmartPLS software, in greater detail in Chapter 8 (see also Cefis, Angelelli, Carpita & Ciavolino, 2025; Hair, Sarstedt, Ringle, & Gudergan, 2024, Chapter 3). Clearly, a purely data-driven perspective needs to be supplemented with theoretical considerations based on the guidelines summarized in Table 2.1.

Finally, and as indicated in Chapter 1, researchers need to carefully distinguish the measurement-theoretic perspective from the way PLS-SEM approximates theoretical concepts (Chapter 3). For example, indicators in a measurement model might be conceptually (and empirically) highly correlated, giving rise to a reflective specification. However, the mental model guiding respondents' answering behavior frequently follows a composite logic. This is because attitudes expressed at a given moment are the result of a constructive process (Zaller & Feldman, 1992). Rather than representing a single, fixed perspective on an object or issue, attitudes are dynamic, real-time constructs shaped by accessible information, current emotional states, and contextual cues (e.g., Fazio, Lenn, & Effrein, 1984); see Guenther, Guenther, Ringle, Zaefarian and Cartwright (2025) for a discussion.

TABLE 2.1 ■ Guidelines for Choosing the Measurement Model Mode

Criterion	Decision	Reference
What is the causal priority between the indicator and the construct?	- From the construct to the indicators: reflective - From the indicators to the construct: formative	Diamantopoulos and Winklhofer (2001)
Is the construct a trait explaining the indicators or rather a combination of the indicators?	- If trait: reflective - If combination: formative	Fornell and Bookstein (1982)
Do the indicators represent consequences or causes of the construct?	- If consequences: reflective - If causes: formative	Rossiter (2002)
Is it necessarily true that if the assessment of the trait changes, all items will change in a similar manner (assuming they are equally coded)?	- If yes: reflective - If no: formative	Chin (1998)
Are the items mutually interchangeable?	- If yes: reflective - If no: formative	Jarvis, MacKenzie, and Podsakoff (2003)

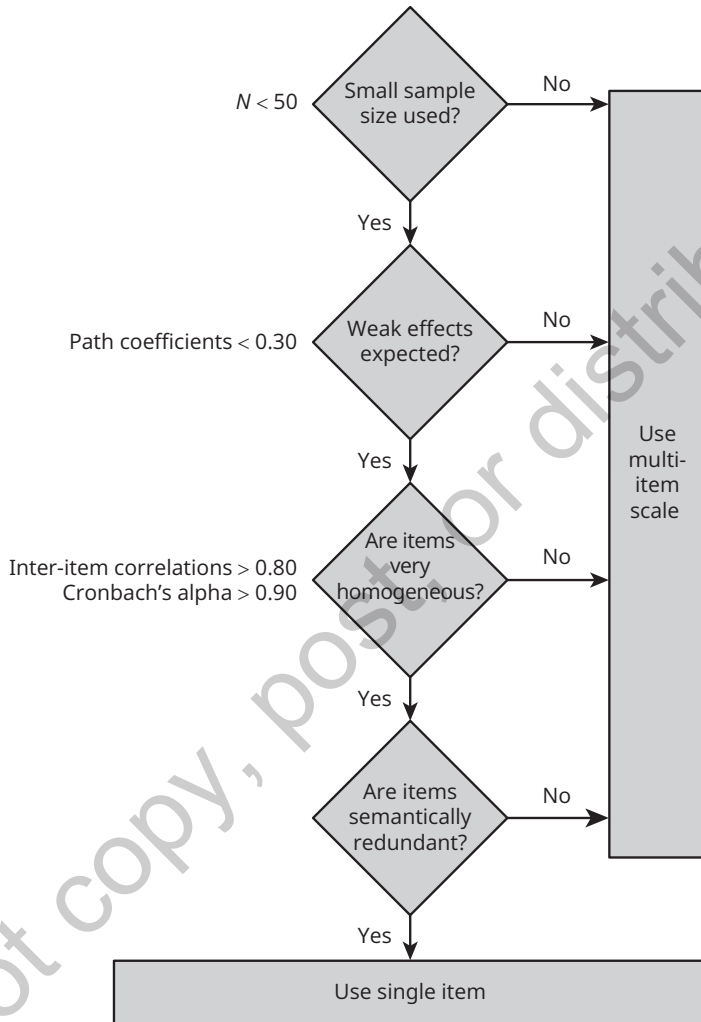
Single-Item Measures

Rather than using multiple items to measure a construct, researchers sometimes choose to use a single item. PLS-SEM proves valuable in this respect because the method is not characterized by identification problems when using fewer than three items in a measurement model, as is the case with CB-SEM (Hair, Black, Babin, & Anderson, 2019). Single items have practical advantages such as ease of application, brevity, and lower costs associated with their use. Unlike long and complicated scales, which often result in a lack of understanding and mental fatigue for respondents, single items promote higher response rates because the questions can be easily and quickly answered (Fuchs & Diamantopoulos, 2009; Sarstedt & Wilczynski, 2009). However, single-item measures do not offer more for less. For instance, when partitioning the data into

groups, researchers have fewer options because scores from only a single variable are available to partition the data. Similarly, information is available from only a single measure instead of several measures when using imputation methods to deal with missing values.

More importantly, from a psychometric perspective, single-item measures do not allow for the removal of measurement error (as is the case with multiple items), which generally decreases their reliability. Note that, contrary to commonly held beliefs, single-item reliability can be estimated (e.g., Cheah, Sarstedt, Ringle, Ramayah, & Ting, 2018; Loo, 2002; Wanous, Reichers, & Hudy, 1997)—see Box 5.1 in Chapter 5 for details. In addition, opting for single-item measures in most empirical settings is a risky decision when it comes to predictive validity considerations. Specifically, the set of circumstances that would favor the use of single-item over multi-item measures is unlikely to be encountered in practice. According to the guidelines by Diamantopoulos, Sarstedt, Fuchs, Kaiser, and Wilczynski (2012), single-item measures should be considered only in situations when (1) small sample sizes are present (i.e., $N < 50$), (2) path coefficients (i.e., the coefficients linking constructs in the structural model) of 0.30 and lower are expected, (3) items of the originating multi-item scale are highly homogeneous (i.e., inter-item correlations > 0.80 , Cronbach's alpha > 0.90), and (4) the items are semantically redundant (Figure 2.9). For further discussions on the efficacy of single-item measures, see Kamakura (2015).

When setting up measurement models, this purely empirical perspective should be complemented with practical considerations. Some research situations call for or even necessitate the use of single items. Respondents frequently feel they are over-surveyed, which contributes to lower response rates (Eggleston, 2024). The difficulty of obtaining large sample sizes in surveys, often due to a lack of willingness to take the time to complete questionnaires, leads to the necessity of considering reducing the length of construct measures where possible. Therefore, if the population being surveyed is small, or only a limited sample size is available (e.g., due to budget constraints, difficulties in recruiting respondents, or dyadic data), the use of single-item measures may be a pragmatic solution. Even if researchers accept the consequences of lower predictive validity and use single-item measures anyway, one fundamental question remains: What should this item be? Unfortunately, research clearly points to severe difficulties when choosing a single item from a set of candidate items, regardless of whether this selection is based on statistical measures or expert judgment (Sarstedt, Diamantopoulos, Salzberger, & Baumgartner, 2016). Against this background, we clearly advise against the use of single items for construct measurement unless indicated otherwise by Diamantopoulos, Sarstedt, Fuchs, Kaiser, and Wilczynski's (2012) guidelines. Finally, it is important to note that these issues must be considered for the measurement of unobservable phenomena, such as perceptions or attitudes. But single-item measures are clearly appropriate when used to measure observable characteristics such as sales, quotas, profits, and so on.

FIGURE 2.9 ■ Guidelines for Single-Item Use

Source: Diamantopoulos, A., Sarstedt, M., Fuchs, C., Kaiser, S., & Wilczynski, P. [2012]. Guidelines for choosing between multi-item and single-item scales for construct measurement: A predictive validity perspective. *Journal of the Academy of Marketing Science*, 40(3), 434–449.

Sum Scores

Some scholars recommend the use of sum scores over differentiated indicator weights as produced by PLS-SEM (e.g., Rönkkö, Lee, Evermann, McIntosh, & Antonakis, 2023). Instead of explicitly estimating their varying relationships with

the construct, the sum scores approach uses the average value of the indicators to compute latent variable scores. Sum scores, therefore, are a simplification of PLS-SEM in that all indicator weights in the measurement model are weighted equally (see Box 1.1 in Chapter 1). The use of sum scores is problematic, however, because it ignores the effect of measurement error inherent in each indicator and conceals measurement model issues as indicators with low loadings or weights remain unidentified. In contrast, the individual weighting of the indicators in a PLS-SEM analysis accounts for measurement error, thereby increasing the reliability and validity of the model estimates (Yuan, Wen, & Tang, 2020). For example, Hair, Hult, Ringle, Sarstedt, and Thiele (2017) have shown that sum scores can produce substantial parameter biases and lower levels of statistical power compared to PLS-SEM. Similarly, Hair, Sharma, Sarstedt, Ringle, and Liengaard (2024) have shown that sum scores almost always lead to lower out-of-sample predictive power compared to PLS-SEM-based differentiated weights (see also McNeish, 2024). Apart from these reliability- and validity-related concerns of sum scores this approach does not inform researchers which indicators have higher or lower relative importance (Hair, Sarstedt, Ringle, Sharma, & Liengaard, 2024).

McNeish and Wolf (2020) summarized the conceptual and empirical shortcomings of sum scores as follows (see Hair, Sharma, Sarstedt, Ringle, & Liengaard, 2024):

- Sum scores produce unrealistic expectations about the population model by enforcing unnatural constraints on the empirical model, which prove particularly problematic in the context of measurement invariance assessment,
- they hinder rigorous and accurate psychometric assessments by ignoring measurement theory in its entirety,
- they adversely affect construct validity and reliability,
- they can result in vastly different conclusions due to inaccurate coefficient estimation, and
- because virtually all the psychometric scales used in management and social sciences research have been validated under the assumption of differentiated weights, using equal weights when applying these scales is categorically inappropriate because doing so enforces a model different from the one validated.

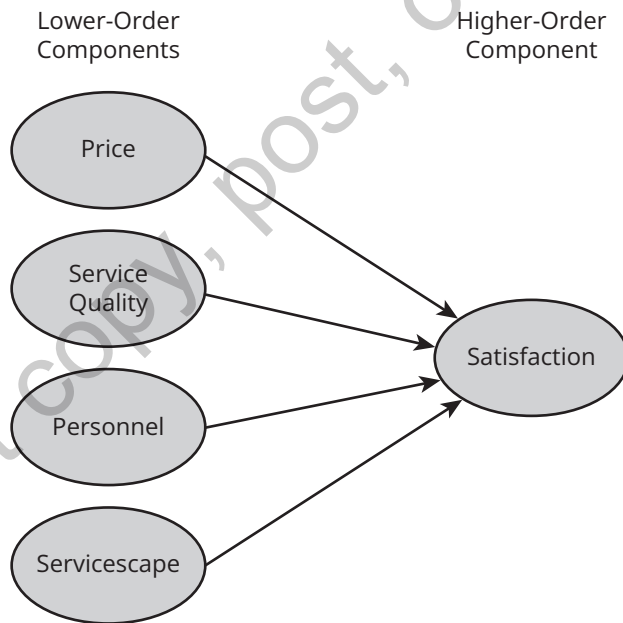
We therefore clearly advise against the use of sum scores in favor of differentiated weights estimated through PLS-SEM.

Higher-Order Constructs

Thus far, we have considered constructs that are measured on a single layer of abstraction. That is, we measured each construct with a set of indicators that are

similar in terms of their concreteness. However, PLS-SEM also allows researchers to model a construct on multiple layers of abstraction simultaneously. **Higher-order constructs**, also referred to as **higher-order models** or **hierarchical component models** in the context of PLS-SEM (Lohmöller, 1989; Sarstedt, Hair, Cheah, Becker, & Ringle, 2019), enable researchers to specify a single construct on a more abstract dimension and more concrete subdimensions at the same time (Cheah et al., 2019; Hair, Sarstedt, Ringle, & Gudergan, 2024, Chapter 2; Wetzels, Odekerken-Schröder, & van Oppen, 2009). For example, the construct satisfaction can be represented by a number of more concrete aspects, measured by **lower-order components** that capture separate attributes of satisfaction. In the context of services, these might include satisfaction with the quality of the service, the service personnel, the price, or the service scape. These lower-order components might form the more abstract **higher-order component** satisfaction, as shown in Figure 2.10.

FIGURE 2.10 ■ Example of a Higher-Order Construct



Instead of modeling the attributes of satisfaction as drivers of the respondent's overall satisfaction on a single construct layer, the higher-order construct summarizes the lower-order component into a single multidimensional construct.

This modeling approach leads to more parsimony and reduces model complexity. Theoretically, this process can be extended to any number of multiple layers, but researchers usually restrict their modeling approach to two layers of abstraction (i.e., one higher-order component and several lower-order components). Constructs with two layers of abstraction are also referred to as **second-order constructs**. Chapter 8 offers more details on higher-order constructs (see also Becker et al., 2023; Hair, Sarstedt, Ringle, & Gudergan, 2024, Chapter 2; Hauff, Richter, & Ringle, 2025).

At this point, you should be able to create a path model. Box 2.1 summarizes some key guidelines you should consider when preparing your path model. The next section continues with collecting the data needed to empirically test your PLS path model.

BOX 2.1 ■ Guidelines for Preparing Your PLS Path Model

Structural model

- The constructs considered relevant to the study must be clearly identified and defined.
- The structural model discussion states how the constructs are related to each other, that is, which constructs are dependent (endogenous) or independent (exogenous). This may also involve considering more complex relationships such as mediators, moderators, and nonlinear effects or including control variables.
- If possible, the nature (positive or negative) of the relationships as well as the direction is hypothesized on the basis of theory, logic, previous research, or researcher judgment.
- There is a clear explanation of why you expect these relationships to exist. The explanation cites theory, qualitative research, business practice, or some other credible source.
- A theoretical model or framework is prepared to clearly illustrate the hypothesized relationships.

Measurement model

- The measurement model discussion states whether constructs are conceptualized as regular or higher-order constructs.
- The measurement specification (i.e., reflective vs. formative) must be clearly stated and motivated. A construct's conceptualization and the aim of the study guide this decision.
- Single-item measures should be used only if indicated by Diamantopoulos, Sarstedt, Fuchs, Kaiser, & Wilczynski's (2012) guidelines.
- Do not use the sum scores approach to measurement.

STAGE 3: DATA COLLECTION AND EXAMINATION

Application of PLS-SEM requires quantitative data to be available. Social science researchers typically use primary data that have been collected for a specific research project, commonly using questionnaires (Sarstedt & Mooi, 2019; Chapter 3). However, researchers are increasingly turning their attention to secondary data that are available from databases or come in the form of website tracking information; social media, geospatial, and sensor data; as well as other information obtained through scraping and similar data collection methods (Hulland, Baumgartner, & Smith, 2018).

When empirical data are collected using questionnaires, typically data collection issues must be addressed after the data are collected. The primary issues that need to be examined include missing data, suspicious response patterns (straight lining or inconsistent answers), outliers, and data distribution. We briefly address each of these on the following pages. The reader is referred to more comprehensive discussions of these issues in Hair, Black, Babin, and Anderson (2019).

Missing Data

Researchers often must deal with missing data. There are two levels at which missing data occur:

- Entire surveys are missing (survey non-response).
- Respondents have not answered all the items (item nonresponse).

Survey nonresponse (also referred to as unit nonresponse) occurs when entire surveys are missing. Survey nonresponse is common because only 5–25% of surveys are typically filled out. Item nonresponse occurs when respondents do not provide answers to certain questions. There are different forms of missingness, including people not filling out or refusing to answer questions. Item nonresponse is common, and 2–10% of questions usually remain unanswered. However, this number greatly depends on factors such as the subject matter, the length of the questionnaire, and the method of administration. Nonresponse can be much higher in respect to questions that people consider sensitive and varies from country to country. In some countries, for instance, reporting income is a sensitive issue (Sarstedt & Mooi, 2019; Chapter 3). As a rule of thumb, when the amount of missing data for a specific respondent exceeds 15%, the observation should be removed from the data set. Similarly, we recommend excluding an indicator from the analysis if it has more than 15% missing values.

Once observations with too many missing responses have been removed, the next step is to decide how to deal with the remaining missing values in the data

set. Research has proposed **missing value treatment** options, some of which have been implemented in the SmartPLS 4 software.

In **mean value replacement**, the missing values of an indicator variable are replaced with the mean of valid values of that indicator. Although easy to implement, mean value replacement decreases the variability in the data and likely reduces the possibility of finding meaningful relationships. It should therefore be used only when the data exhibit very low levels of missing data.

Researchers can also choose to remove all cases from the analysis that include missing values in any of the indicators used in the model (referred to as **casewise deletion** or **listwise deletion**). However, when using casewise deletion researchers run the risk of systematically deleting a certain group of respondents. For example, market researchers frequently observe that affluent respondents are more likely to refuse answering questions related to their income. Running casewise deletion would systematically omit this group of respondents and therefore yield erroneous conclusions. Second, using casewise deletion can dramatically diminish the number of observations in the data set. It is therefore crucial to carefully check the number of observations used in the final model estimation when this type of missing value treatment is used.

Instead of discarding all observations with missing values, **pairwise deletion** uses all observations with complete responses in the calculation of the model parameters. For example, assume we have a measurement model with three indicators (x_1 , x_2 , and x_3). To estimate the model parameters, all valid values in x_1 , x_2 , and x_3 are used in the computation. That is, if a respondent has a missing value in x_3 , the valid values in x_1 and x_2 are still used to calculate the model.

Which of these missing value treatment options should be used? Research doesn't offer a clear answer to this question. Amusa and Hossana (2024) have compared the options in a PLS-SEM context and found only minor difference among the options (in terms of biases) when the data are missing at random (MAR)—even when the share of missing data is extremely high. MAR means that the probability that a data point is missing depends on the respondents' traits (e.g., age, gender, or income) or their answering behavior regarding other variables in the data set (Sarstedt & Mooi, 2019; Chapter 5). However, biases produced by the different treatment options are more pronounced when the data are missing not at random (MNAR). MNAR means that the probability of a data point missing depends on the variable itself. For example, affluent and poor people are generally less likely to indicate their income (Sarstedt & Mooi, 2019; Chapter 5). In this missingness pattern, researchers should prefer regression imputation over simple mean imputation. Finally, regardless of whether the data are MAR or MNAR, researchers should avoid using casewise deletion, where all observations with missing values are removed from the analysis. This treatment option has produced considerably higher biases in most settings. Contrary to this recommendation, Grimm and Wagner (2020) showed that PLS-SEM estimates are stable when using casewise deletion when up to 9% of the data are missing completely at random (MCAR). When data are MCAR, the probability of a data

point being missing is entirely independent of both the respondents' traits and the variable with missing values. Liu, Chin, Cheah, Hair, and Lyu (2025) found similar results, showing that parameter biases are low with up to 5% of MCAR data. Finally, most researchers advise against the use of pairwise deletion, sometimes even calling it "unwise deletion," because the different calculations in the analysis may be based on different sample sizes, which can bias the results (Hair, Black, Babin, & Anderson, 2019). Exceptions are situations in which many observations have missing values—thus hindering the use of mean replacement and especially casewise deletion—and the aim of the analysis is to gain first insights into the model structure.

Considering these findings, we recommend to generally avoid using pairwise deletion. Researchers should rely on mean value replacement when there are less than 5% values missing per indicator. In the rare cases where data are MCAR, casewise deletion may be used, provided that the share of missing data is low. In case more data points are missing, more complex imputation routines may be used (Little & Rubin, 2002; Schafer & Graham, 2002), which at the time of this writing, need to be executed outside the SmartPLS 4 software. Liu, Chin, Cheah, Hair, and Lyu (2025) have evaluated the efficacy of more complex options and found that a combined EM imputation with inverse probability weighting proves particularly valuable, regardless of whether the data are MAR or MNAR.

Suspicious Response Patterns

Before analyzing their data, researchers should also examine response patterns. In doing so, they are looking for a pattern often described as straight lining. **Straight lining** occurs when a respondent marks the same response for a high proportion of the questions. For example, if a 7-point scale is used to obtain answers and the response pattern is all 4s (the middle response), then that respondent in most cases should be deleted from the data set. Similarly, if a respondent selects only 1s or only 7s, then that respondent should in most cases be removed. Other suspicious response patterns are **diagonal lining** and **alternating extreme pole responses**. A visual inspection of the responses or the analysis of descriptive statistics (e.g., mean, variance, and distribution of the answers per respondent) allows identifying suspicious response patterns.

Inconsistency in answers may also need to be addressed before analyzing the data. Many surveys start with one or more screening questions. The purpose of a screening question is to ensure that only individuals who meet the prescribed criteria complete the survey. For example, a survey of mobile phone users may screen for individuals who own an Apple iPhone. But a question later in the survey is posed, and the individual indicates they use an Android device. This respondent would therefore need to be removed from the data set. Surveys often ask the same question with slight variations, especially when reflective indicators are used. If a respondent gives a different answer to the same question asked in a slightly different way, this too raises a red flag and suggests the respondent was not reading the

questions closely or simply was marking answers to complete and exit the survey as quickly as possible. Finally, researchers sometimes include specific questions to assess the attention of respondents. For example, in the middle of a series of questions, the researcher may instruct the respondent to check only a 1 on a 7-point scale for the next question. If any answer other than a 1 is given for the question, it is an indication the respondent is not closely reading the question.

Outliers

An **outlier** is an extreme response to a particular question or extreme responses to all questions. Outliers must be interpreted in the context of the study, and this interpretation should be based on the type of information they provide. Outliers can result from data collection or entry errors (e.g., the researcher coded “77” instead of “7” on a 1–9 Likert scale). However, exceptionally high or low values can also be part of reality (e.g., an exceptionally high income). Finally, outliers can occur when combinations of variable values are particularly rare (e.g., spending 80% of annual income on holiday trips). The first step in dealing with outliers is to identify them. Standard statistical software packages offer a multitude of univariate, bivariate, or multivariate graphs and statistics, which allow identifying outliers. For example, when analyzing box plots, one may characterize responses as extreme outliers, which are three times the interquartile range below the first quartile or above the third quartile. Moreover, IBM SPSS Statistics has an option called Explore that develops box plots and stem-and-leaf plots to facilitate the identification of outliers by respondent number (Sarstedt & Mooi, 2019; Chapter 5).

Once the outliers are identified, the researcher must decide what to do. If there is an explanation for exceptionally high or low values, outliers are typically retained because they represent an element of the population. However, their impact on the analysis results should be carefully evaluated. That is, one should run the analyses with and without the outliers to ensure that a few (extreme) observations do not influence the results substantially. If the outliers are a result of data collection or entry errors, they are always deleted or corrected (e.g., the value of 55 on a 7-point scale). If there is no clear explanation for the exceptional values, outliers should be retained—see Sarstedt and Mooi (2019; Chapter 5) for more details about outliers.

Outliers can also represent a unique subgroup of the sample. There are two approaches to use in deciding if a unique subgroup exists. First, a subgroup can be identified based on prior knowledge, for example, based on observable characteristics such as gender, age, or income. Using this information, the researcher partitions the data set into two or more groups and runs a multigroup analysis to disclose significant differences in the model parameters. The second approach to identifying unique subgroups is the application of latent class techniques. Latent class techniques allow researchers to identify and treat unobserved heterogeneity, which cannot be attributed to a specific observable characteristic or a combination

of characteristics. Several latent class techniques have recently been proposed that generalize finite mixture modeling, iteratively reweighted least squares, hill-climbing approaches, and genetic algorithms to PLS-SEM (Sarstedt, Ringle, & Hair, 2017). In Chapter 8, we discuss several of these techniques in greater detail.

Data Distribution

PLS-SEM is a nonparametric statistical method. Different from CB-SEM, which draws on a maximum likelihood estimator that requires normally distributed data, PLS-SEM does not make any distributional assumptions (Hair, Ringle, & Sarstedt, 2011). Nevertheless, it is important to verify that the data are not too far from normal because extremely nonnormal data prove problematic in the assessment of the parameters' significances. Specifically, extremely non-normal data inflate standard errors obtained from bootstrapping (see Chapter 5 for more details; e.g., Beasley et al., 2007; Wood, 2005) and thus trigger type II errors (i.e., false negatives).

Researchers can draw on a wide range of statistical tests to evaluate whether the data are (non)normal. These tests emphasize either skewness, kurtosis, or both, and therefore differ in their performance (Yazici & Yolacan, 2007; Yap & Sim, 2011). One such test is the **Cramér-von Mises test**, which Becker, Proksch, and Ringle (2022) identified as generally useful for determining (non)normality. However, with larger sample sizes, normality tests like the Cramér-von Mises test often yield statistically significant results even when deviations from normality are practically negligible (i.e., they are subject to type I errors with growing sample size). Researchers should therefore complement the Cramér-von Mises test results with graphical inspections (e.g., histograms, Q–Q plots) and consider distributional measures such as skewness and kurtosis to comprehensively assess potential deviations from normality.

Skewness assesses the extent to which a variable's distribution is symmetrical. If the distribution of responses for a variable stretches toward the right or left tail of the distribution, then the distribution is characterized as skewed. A negative skewness indicates a greater number of larger values, whereas a positive skewness indicates a greater number of smaller values. As a general guideline, a skewness value between -1 and $+1$ is considered excellent, whereas a value between -2 and $+2$ is generally considered acceptable. Values beyond -2 and $+2$ are considered indicative of substantial nonnormality. **Kurtosis** is a measure of whether the distribution is too peaked (a narrow distribution with most of the responses in the center). A positive value for the kurtosis indicates a distribution more peaked than normal. In contrast, a negative kurtosis indicates a shape flatter than normal. Analogous to the skewness, the general guideline is that if the kurtosis is greater than $+2$, the distribution is too peaked. Likewise, a kurtosis of less than -2 indicates a distribution that is too flat. When both skewness and kurtosis are close to zero, the pattern of responses is considered a normal distribution (George & Mallery, 2019).

Serious effort, considerable amounts of time, and a high level of caution are required when collecting and analyzing the data you need for carrying out multivariate techniques. Always remember the garbage in, garbage out rule. All your analyses are meaningless if your data are inappropriate. Box 2.2 summarizes some key guidelines you should consider when examining your data and preparing them for PLS-SEM. For more detail on examining your data, see Chapter 2 of Hair, Black, Babin, and Anderson (2019).

BOX 2.2 ■ Guidelines for Examining Data Used With PLS-SEM

- Missing data must be identified. When missing data per observation or per indicator exceed 15%, they may be removed from the data set. Other missing data should be dealt with before running a PLS-SEM analysis. When less than 5% of values *per indicator* are missing, use mean replacement. Otherwise, use more complex treatment options or casewise deletion, provided that the deletion of observations does not occur systematically and that enough observations remain for the analysis. Generally avoid using pairwise deletion.
- Suspicious and inconsistent response patterns typically justify removing a response from the data set.
- Outliers should be identified before running PLS-SEM. Subgroups that are substantial in size should be identified based on prior knowledge or by statistical means (e.g., using a latent class analysis).
- Lack of normality in variable distributions can distort the results of multivariate analysis. This problem is much less severe with PLS-SEM, but researchers should still examine PLS-SEM results carefully when distributions deviate substantially from normal. Besides considering, for instance, the Cramér-von Mises normality test, researchers should graphically inspect a variable's distribution and assess if absolute skewness and kurtosis values are greater than 2, which indicates nonnormal data.

CASE STUDY ILLUSTRATION— SPECIFYING THE PLS PATH MODEL

The most effective way to learn how to use a statistical method is to apply it to a data set. In this book, we use a single example that allows you to apply the method in practice using the SmartPLS 4 software (Ringle, Wende, & Becker, 2024)—see also the SmartPLS review articles by Cheah, Magno and Cassia (2024), Memon et al. (2021), and Sarstedt and Cheah (2019), or textbooks based on the SmartPLS software, for example, by Chua (2024), Ramayah, Cheah, Chuah, Ting, and Memon (2018), and Wong (2019). We start the example with a simple model, which we expand in Chapter 5. For our initial model, we hypothesize a

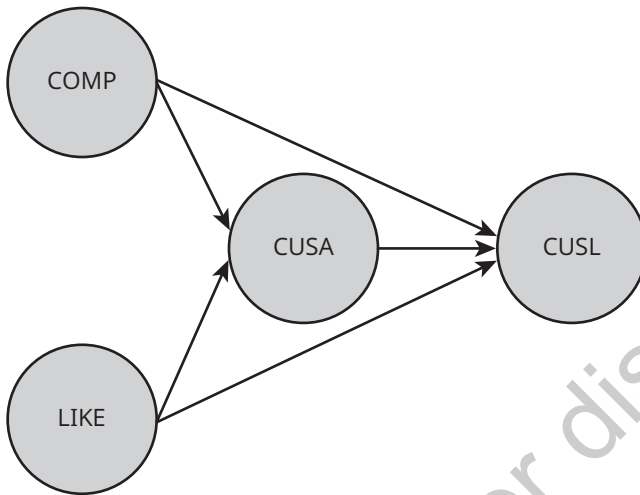
path model to estimate the relationships among corporate reputation, customer satisfaction, and customer loyalty. The example will provide insights on (1) how to develop the structural model representing the underlying concepts or theory, (2) the setup of measurement models for the constructs, and (3) the structure of the empirical data used. Then, our focus shifts to setting up the SmartPLS for PLS-SEM.

Application of Stage 1: Structural Model Specification

To specify the structural model, we begin with some fundamental assumptions about theoretical models. The corporate reputation model by Eberl (2010) is the basis of our theory. The goal of the model is to explain the effects of corporate reputation on customer satisfaction (*CUSA*) and, ultimately, customer loyalty (*CUSL*). Corporate reputation represents a company's overall evaluation by its stakeholders (Helm, Eggert, & Garnefeld, 2010). It is measured using two dimensions. One dimension represents cognitive evaluations of the company, and the construct is the company's competence (*COMP*). The second dimension captures affective judgments, which determine the company's likeability (*LIKE*). This two-dimensional approach to measure reputation was developed by Schwaiger (2004). It has been applied in various research studies (e.g., Damberg, Schwaiger, & Ringle, 2022; Eberl & Schwaiger, 2005; Radomir & Moisesescu, 2019; Radomir & Wilson, 2018; Raithel & Schwaiger, 2015; Raithel, Wilczynski, Schloderer, & Schwaiger, 2010; Sarstedt & Schloderer, 2010; Schloderer, Sarstedt, & Ringle, 2014; Schwaiger, Raithel, & Schloderer, 2009; Yun, Kim, & Cheong, 2020) and validated in different countries (e.g., Damberg, Liu, & Ringle, 2024; Eberl, 2010; Zhang & Schwaiger, 2012) as well as other contexts (e.g., Damberg, 2023; Sarstedt, Ringle, & Iuklanov, 2023). Research also shows that the approach performs favorably (in terms of convergent validity and predictive validity) compared with alternative reputation measures (Sarstedt, Wilczynski, & Melewar, 2013).

Schwaiger (2004) further identified four antecedent dimensions of reputation—quality, performance, attractiveness, and corporate social responsibility—measured by a total of 21 formative indicators. These driver constructs of corporate reputation are components of the more complex example that we will use in Chapter 5 and beyond. Likewise, we do not consider more complex model setups such as mediation or moderation effects yet. These aspects will be covered in the case studies in Chapter 7.

In summary, the simple corporate reputation model has two main theoretical components: (1) the target constructs of interest—namely, *CUSA* and *CUSL* (endogenous constructs)—and (2) the two corporate reputation dimensions *COMP* and *LIKE* (exogenous constructs), which represent key determinants of the target constructs. Figure 2.11 shows the constructs and their relationships, which represent the structural model for the PLS-SEM case study.

FIGURE 2.11 ■ Example of a Theoretical Model (Simple Model)

The structural model displays the theoretical framework with its key elements (i.e., constructs) and cause-effect relationships (i.e., paths). Researchers typically develop hypotheses for the constructs and their path relationships in the structural model. For example, consider Hypothesis 1 (H₁): Customer satisfaction has a positive effect on customer loyalty. PLS-SEM enables statistically testing the significance of the hypothesized relationship (Chapter 6). When conceptualizing the theoretical constructs and their hypothesized structural relationships for PLS-SEM, it is important to make sure the model has no circular relationships (i.e., causal loops). A circular relationship would occur if, for example, we reversed the relationship between *COMP* and *CUSL* because this would yield the causal loop *COMP* → *CUSA* → *CUSL* → *COMP*.

Application of Stage 2: Measurement Model Specification

Because the constructs are not directly observed, we need to specify a measurement model for each construct. The specification of the measurement models (i.e., multi-item vs. single-item measures and reflective vs. formative measures) draws on prior research studies by Schwaiger (2004) and Eberl (2010).

In our simple example of a PLS-SEM application, we have three constructs (*COMP*, *CUSL*, and *LIKE*) measured by multiple items (Figure 2.12). All three constructs have reflective measurement models as indicated by the arrows pointing from the construct to the indicators. For example, *COMP* is measured by means of the three reflective items *comp_1*, *comp_2*, and *comp_3*, which relate to

the following survey questions (Table 2.2): “[The company] is a top competitor in its market,” “As far as I know, [the company] is recognized worldwide,” and “I believe that [the company] performs at a premium level.” Respondents had to indicate the degree to which they (dis)agree with each of the statements on a 7-point scale from 1 = *fully disagree* to 7 = *fully agree*.

FIGURE 2.12 ■ Types of Measurement Models in the Simple Model

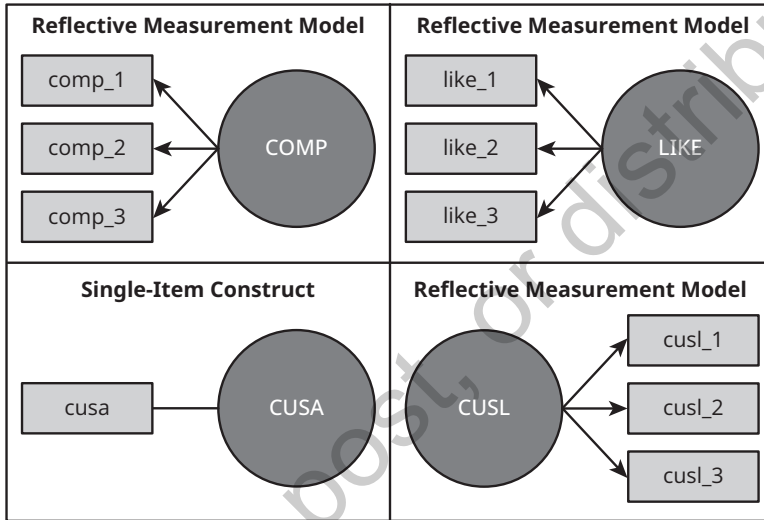


TABLE 2.2 ■ Indicators for Reflective Measurement Model Constructs

Competence (COMP)	
comp_1	[The company] is a top competitor in its market.
comp_2	As far as I know, [the company] is recognized worldwide.
comp_3	I believe that [the company] performs at a premium level.
Likeability (LIKE)	
like_1	[The company] is a company that I can better identify with than other companies.
like_2	[The company] is a company that I would regret more not having if it no longer existed than I would other companies.
like_3	I regard [the company] as a likeable company.

(Continued)

TABLE 2.2 ■ (Continued)	
Customer Loyalty (CUSL)	
cusl_1	I would recommend [the company] to friends and relatives.
cusl_2	If I had to choose again, I would choose [the company] as my mobile phone services provider.
cusl_3	I will remain a customer of [the company] in the future.

Note: For data collection, the actual name of the company was inserted in the bracketed space.

Different from *COMP*, *CUSL*, and *LIKE*, the customer satisfaction construct (*CUSA*) is operationalized by a single item (*cusa*) associated with the following question in the survey: “If you consider your experiences with [company], how satisfied are you with [company]?” The single indicator is measured with a 7-point scale indicating the respondent’s degree of satisfaction (1 = *very dissatisfied*; 7 = *very satisfied*). The single item was used due to practical considerations in an effort to decrease the overall number of items in the questionnaire. As customer satisfaction items are usually highly homogeneous, the loss in predictive validity compared with a multi-item measure is not considered severe. Moreover, because *cusa* is the only item measuring customer satisfaction, construct and item are equivalent (as indicated by the fact that the relationship between construct and single-item measure is always one in PLS-SEM). Therefore, the choice of the measurement perspective (i.e., reflective vs. formative) is of no concern, and the relationship between construct and indicator is undirected.

Application of Stage 3: Data Collection and Examination

To estimate the PLS-SEM, data were collected using computer-assisted telephone interviews (Sarstedt & Mooi, 2019; Chapter 4) that asked about the respondents’ perception of and their satisfaction with four major mobile network providers in Germany’s mobile communications market. Respondents rated the questions on 7-point Likert scales with higher scores denoting higher levels of agreement with a particular statement. In the case of *cusa*, higher scores denote higher levels of satisfaction. Satisfaction and loyalty were measured with respect to the respondents’ own service providers. The data set used in this book is a subset of the original set and has a sample size of 344 observations. The data have been collected using a quota sampling approach (Sarstedt, Bengart, Shaltoni, & Lehmann, 2018) by a professional market research company in the German market.

Table 2.3 shows the data matrix for the model. The 10 columns represent a subset of all variables (i.e., specific questions in the survey as described in the previous section) that have been surveyed, and the 344 rows (i.e., cases) contain the

answers of every respondent to these questions. For example, the first row contains the answers of Respondent 1, whereas the last row contains the answers of Respondent 344. The columns show the answers to the survey questions. Data in the first nine columns are for the indicators associated with the three constructs, and the 10th column includes the data for the single indicator for *CUSA*. The data set contains further variables that relate to, for example, the driver constructs of *LIKE* and *COMP*. We will cover these aspects in Chapter 5.

TABLE ■ 2.3 Data Matrix for the Indicator Variables

Case Number	Variable Name										
	comp_1	comp_2	comp_3	like_1	like_2	like_3	cusl_1	cusl_2	cusl_3	cusa	
1	6	7	6	6	6	6	7	7	7	7	...
2	4	5	6	5	5	5	7	7	5	6	...
...	...	...	...	...	...	...	...	...	...	...	...
344	6	5	6	6	7	5	7	7	7	7	...

If you are using a data set in which a respondent did not answer a specific question, you need to insert a number that does not appear otherwise in the responses to indicate the missing values. Researchers commonly use -99 to indicate missing values, but you can use any other value that does not normally occur in the data set. In the following, we will also use -99 to indicate missing values. If, for example, the first data point of *comp_1* was a missing value, the -99 value would be inserted into the space as a missing value space holder instead of the value of 6 that you see in Table 2.3. Missing value treatment procedures (e.g., mean replacement) could then be applied to these data (e.g., Hair, Black, Babin, & Anderson, 2019). Again, if the number of missing values in your data set per indicator is relatively small (i.e., less than 5% missing per indicator), we recommend mean value replacement instead of casewise deletion to treat the missing values when running PLS-SEM. Furthermore, we need to ascertain that the number of missing values per observation does not exceed 15%. If this is the case, the corresponding observation should be eliminated from the data set.

The data set used in the book's example has few missing values. None of the observations has more than 15% missing values, so we can proceed analyzing all 344 respondents. Focusing on the indicators, we find that *cusa* has one missing value (0.29%), *cusl_1* and *cusl_3* have three missing values (0.87%), and *cusl_2* has four missing values (1.16%). Because the missing values per indicator are less than 5%, mean value replacement can be used.

To run outlier diagnostics, we compute a series of box plots using IBM SPSS Statistics—see Chapter 5 in Sarstedt and Mooi (2019) for details on how to run

these analyses in IBM SPSS Statistics. The results indicate some influential observations but no outliers.

The *Data view* in SmartPLS, which we will discuss next, offers normality diagnostics. Here we find that the Cramér-von Mises test produces p values less than 0.01 for all variables in the data set and, thus, rejects the null hypothesis of normality. However, upon closer inspection, we find that nonnormality of data regarding skewness and kurtosis is not a serious issue. The kurtosis and skewness values of all the indicators are within the -2 and $+2$ range.

PATH MODEL CREATION USING THE SMARTPLS SOFTWARE

The SmartPLS 4 software (Ringle, Wende, & Becker, 2024) is used to execute all the PLS-SEM analyses in this book. The discussion includes an overview of the software's functionalities. The student version of the software is available free of charge at <https://www.smartpls.com>. The student version offers practically all functionalities of the full version but is restricted to data sets with a maximum of 100 observations. However, because the data set used in this book has more than 100 observations (344 to be precise), you should use the professional version of SmartPLS, which is available as a 30-day trial version at <https://www.smartpls.com>. After the trial period, a license fee applies. Licenses are available for different periods of time (e.g., 1 month, 1 year, 2 years, etc.) and can be purchased through the SmartPLS website. The SmartPLS website includes a download area and many additional resources such as short explanations of PLS-SEM and software-related topics, a list of recommended literature, answers to frequently asked questions, tutorial videos for getting started using the software, and the SmartPLS forum, which enables you to discuss PLS-SEM topics with other users.

Now run the SmartPLS software by clicking on the desktop icon that is available after the software installation on your computer. Alternatively, go to the folder where you installed the SmartPLS software on your computer. Click on the icon that runs SmartPLS to start the software. When starting SmartPLS for the first time, the software will ask you to select a *Workspace*. The workspace is a folder on your computer in which your projects including PLS path models and data sets will be stored. You can choose your SmartPLS workspace folder by clicking on the corresponding button in the main screen on the right of the SmartPLS *Workspace view*.

A good starting point is to explore the ready-to-run SmartPLS sample projects. Under *Sample Projects* on the right-hand side of the screen, you will find several *Sample Projects*, which you can activate by ticking the checkbox next to their label. Locate the section labeled *PLS-SEM*, and click on the checkbox next to *Corporate reputation–PLS-SEM book (primer)*. The project will now appear in the *Workspace* at the left-hand side of the screen. After successfully importing,

double-click *1 Simple Model* in the project. The path model shown in Figure 2.16 will then open automatically in the SmartPLS *Modeling view*. Select the *PLS-SEM algorithm* under *Calculate* in the menu bar to directly estimate the model using the Corporate Reputation data set.

However, instead of working with the ready-to-run reputation model, we will describe how to set up this model step-by-step using the SmartPLS software. Before starting with the creation of a new project, you need to have data that serve as the basis for running the model. SmartPLS 4 supports data imported from file formats such as Microsoft Excel (.xls or .xlsx), SPSS (.sav), comma-separated values (.csv), and text (.txt). A key aspect we have to pay attention to is that the first row contains the variable names in text format and otherwise only numerical values (no text or special characters; also no numerical values in scientific format, e.g., 10 E-7). SmartPLS cannot interpret string elements. For example, when you have a variable indicating a respondent's gender (e.g., female, male, other), you must code the categories with numerical values such as 1, 2, and 3 instead of using text labels. Similarly, SmartPLS interprets single dots (such as those produced by IBM SPSS Statistics in case an observation has a system-missing value) as string elements. The data we will use with the reputation model can be downloaded either as comma-separated value (.csv) or text (.txt) data sets in the download section of this book's web page at the following URL: <https://www.pls-sem.net>.

To do so, go to the *Workspace view*. In case you have opened the sample PLS path model (as previously described), go back to this window by clicking *Back* in the menu bar. Now create a new project by clicking on *New project* or *Files → Create new → New project* in the menu bar. First type a name for the project into the *Name* box (e.g., *Corporate reputation*). After clicking *Create*, the new project is created and appears right under *Workspace* on the left side of your screen. All previously created SmartPLS projects also appear in this window. On top of the main screen of the SmartPLS *Workspace view*, you see the path to your workspace folder (e.g., on your computer) where your SmartPLS projects are saved. You can select or create a new *Workspace* (e.g., on your computer or on a cloud drive) by clicking the *Switch* button and then by selecting a previously used workspace from the list or by clicking on the link with the label *switch or create another workspace*.

Next, you need to assign a data set to the project, in our case, *Corporate reputation data.txt* (or whatever name you gave to the data you downloaded). To do so, click on *Import data file* below the project you created, find and highlight your data folder, and click *Open*. A window opens (Figure 2.13), offering a general overview of the data set, including a list of all variables, their measurement scales as identified by the SmartPLS software, as well as the observed minimum and maximum values. You can specify the *Delimiter character* to determine the separation of the data values in your data set (i.e., comma, semicolon, tabulator, space), the *Escape character* (i.e., none, single quote, double quote) in case the values use quotations (e.g., "7"), and the *Locale* (i.e., United States with a dot as decimal separator or Europe with a comma as decimal separator). SmartPLS

distinguishes between the observed minimum and maximum values and the scale's minimum and maximum values. Whereas the former relates to the actual values as observed in the data set, the latter refers to the theoretical minimum and maximum values that could be present in the data. For example, although an item may have been measured on a 1–7 scale, respondents might not have used the full range. The distinction between the observed and theoretical min and max values is important for some advanced analyses, which we will not cover in this book. You can change the minimum and the maximum value for each variable manually by entering the correct value under *Min* and/or *Max* in the *Change settings* window. For simultaneous changes you can use the *Bulk change* option.

When your data includes missing values, SmartPLS will automatically identify empty cells as missing data points. However, if you use a certain number for missing values such as –99 or –9999, it is important to provide SmartPLS with this information. Otherwise, the software uses this number as an actual response value, which will distort your results. In our data set, missing values are coded as –99. This is also indicated by the minimum value of –99 in some variables (Figure 2.13), where respondents could indicate their responses on a scale from 1 to 7. To indicate that –99 represents missing values in this data set, click on *missing value marker* at the bottom of the window. A box will then open, allowing you to enter the number that represents the missing values. Type –99 into the text field, and click on the *OK* button. The text at the bottom of the window now indicates that –99 matches 11 missing values that SmartPLS identified in the data set. In addition, SmartPLS now identifies 1 as the minimum value for all variables, which previously had –99 as minimum. The only exception is for variables where none of the respondents selected 1 as a response (e.g., *att_global*); in such cases, the minimum value is identified as 2.

Finally, at the top of the window (Figure 2.13), you can select the target *Project* in the *Workspace* (e.g., *Corporate reputation*) and change the *File name* of the data set. In this example, we use the original name (i.e., *Corporate reputation data*) and proceed by clicking *Import*. SmartPLS will automatically open the *Data* view (Figure 2.14), which provides information on the data set and its format.

Below the menu bar, on the right side of the screen, a list appears showing all variables, the number of missing values, the scale type, and basic descriptive statistics (e.g., mean, median, scale minimum, and scale maximum values; observed minimum and observed maximum values; and standard deviation, excess kurtosis, skewness, and Cramér-von Mises *p* values). Note that *Scale min* and *Scale max* refer to the theoretical range of the scale, whereas *Observed min* and *Observed max* relates to the actual range as present in the data set. In case you need to adjust the data set settings (e.g., missing value marker or minimum value of indicator scales), click on *Setup* in the menu bar. Any changes that you make will dynamically update the values in the *Data* view. For example, adjusting the minimum value of the *attr_global* variable from 2 to 1 will change the value displayed in the *Scale min* column accordingly while keeping the value of 2 in the *Observed min* column. In addition, the *Derive indicators* button in the menu bar enables you to

FIGURE 2.13 ■ CSV Import Window in SmartPLS

CSV import
✕

Project

Example - Corporate reputation (primer)

Delimiter character Semicolon

Escape character None

File name

Corporate reputation data

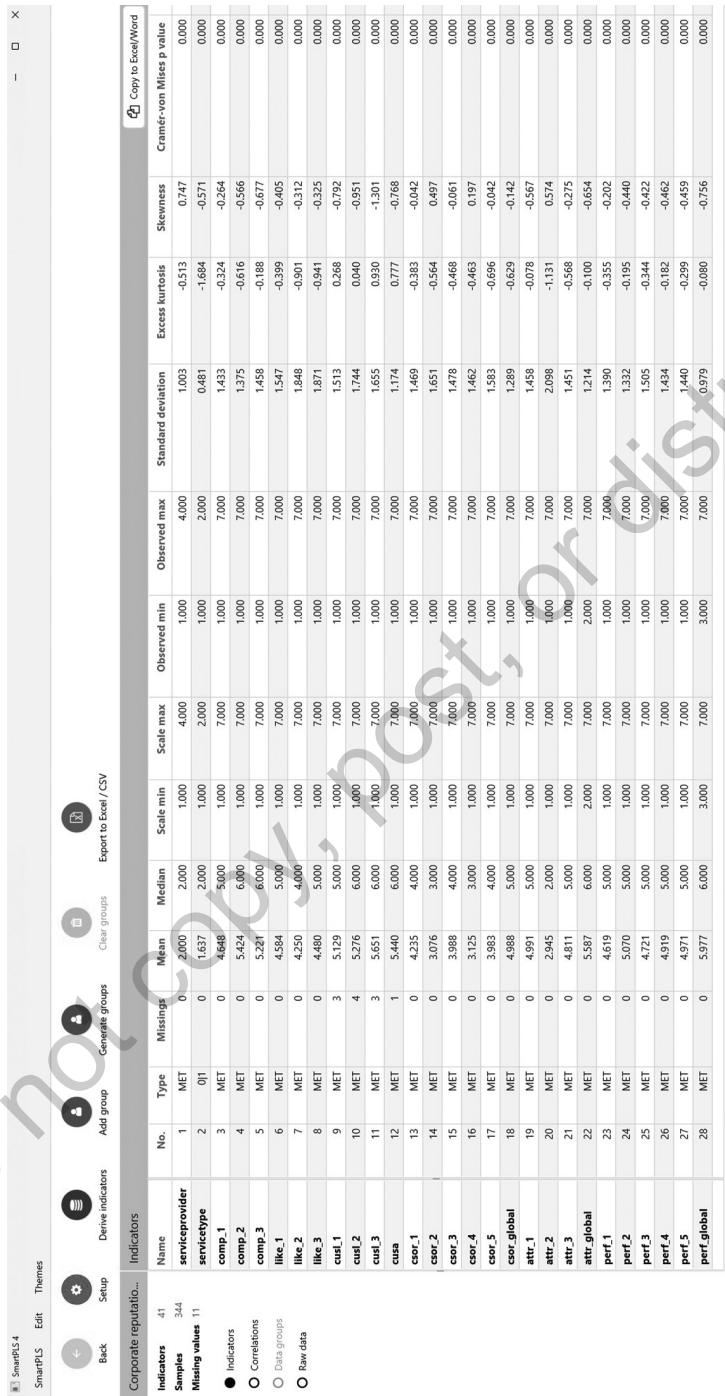
Bulk change

Name	Missing	Scale	Min	Max
✓ serviceprovider	0	Metric	<>	1.0000
✓ servicetype	0	Binary	<>	1.0000
✓ comp_1	0	Metric	<>	1.0000
✓ comp_2	0	Metric	<>	1.0000
✓ comp_3	0	Metric	<>	1.0000
✓ like_1	0	Metric	<>	1.0000
✓ like_2	0	Metric	<>	1.0000
✓ like_3	0	Metric	<>	1.0000
✓ cusl_1	0	Metric	<>	-99.0000
✓ cusl_2	0	Metric	<>	-99.0000
✓ cusl_3	0	Metric	<>	-99.0000

Cases: -1 Missing: 0

Define a missing value marker if your data file contains missing values.

FIGURE 2.14 Data View in SmartPLS



Indicators

41

344

Missing values

11

Indicators

Correlations

Data groups

Raw data

use your existing variables in the data set to create new ones (e.g., by using the *Reverse scales*, *Recode value*, *Combine indicators*, and *Convert categorical variable to dummy variables* options).

The window left of the *Data* view shows the number of indicators, the sample size, and the number of missing values. Here, you can also navigate to the indicator correlation matrix (*Correlations*), *Data groups*, and *Raw data*. Note that to view the data groups, you need to first specify a grouping. This can be done by clicking on *Add group* or *Generate groups* in the menu bar at the top of the screen. However, as the specification of grouping variables is primarily relevant for more advanced analyses, we skip this aspect for now. At this point, you can click on the orange arrow button with the label *Back* in the menu bar and return to the *Workspace* view. Note that you can always reopen the *Data* view by double-clicking on the data set (i.e., *Corporate reputation data*) in the *Workspace* view.

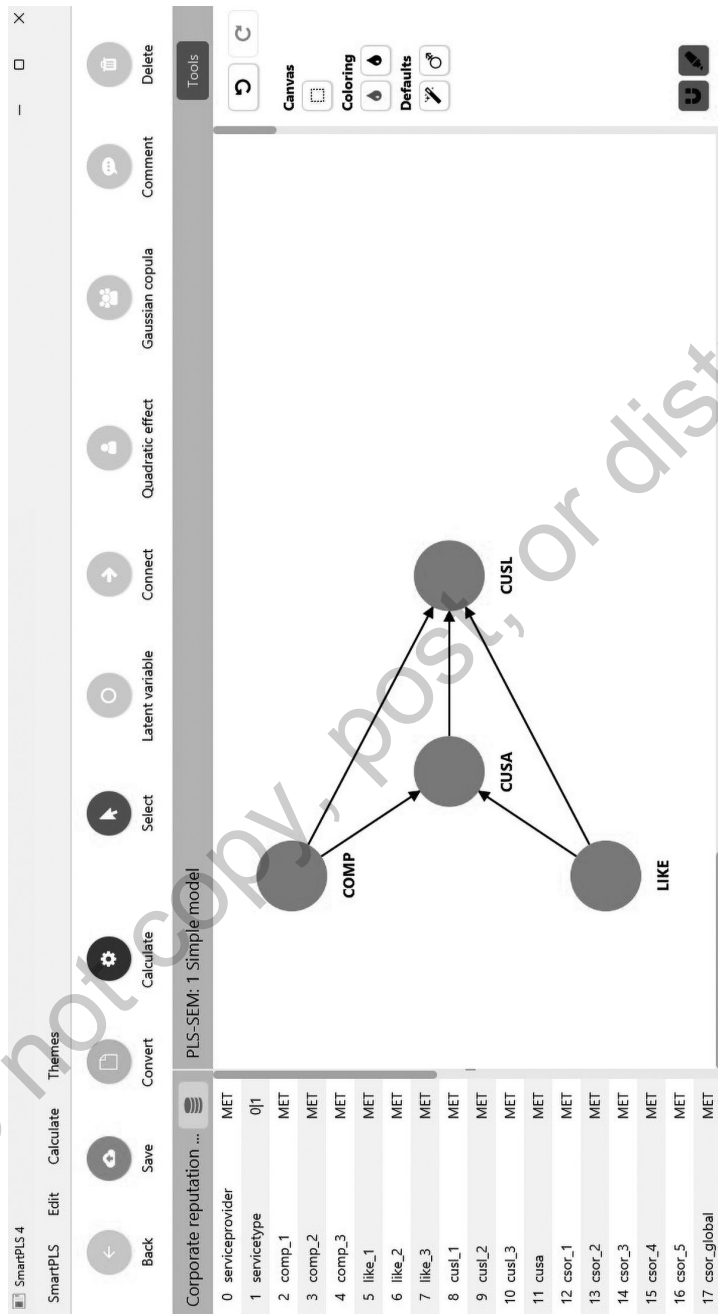
Each project can have one or more path models, data sets in different file formats, and saved reports. To create the first model, click on *Create model* right below your newly created model. In the combo box that opens, choose *PLS-SEM* under *Model type*, enter the *Model name* (e.g., *1 Simple model*), and click on *Save*. Note that you can also create a model by right-clicking on the Project name and choosing the *New PLS-SEM model* option from the dialog that opens.

SmartPLS now automatically opens a new window (or after double-clicking on *1 Simple model*), which shows the indicator list on the left and the *Modeling window* on the right, where you can create a path model. Because we initiated a new project (as opposed to working with a saved or sample project), the *Modeling window* is empty. You can start creating the path model shown in Figure 2.15. To do so, select *Latent variable* in the menu bar, and left-click in the *Modeling window*. A text box is shown, which suggests a name for the first latent variable (*Alpha*), which you should change into *COMP*. When clicking on *OK* or pressing the *Enter* key on your keyboard, SmartPLS will place the construct labeled *COMP* into the *Modeling window*. Each time you left-click in the *Modeling window*, a new text box for entering the name of the new construct will appear. After pressing the *Enter* key, a new construct represented by a red circle will be added.

Once you have included all your constructs, click on *Select* in the menu bar. The select mode allows you to adjust the elements in the *Modeling window*. For example, you can left-click on any of the constructs to select, resize, or move it. Similarly, you can rename a construct by right-clicking on it and selecting *Rename* in the window that opens.

To connect the latent variables with each other (i.e., to draw path relations), left-click on *Connect* in the menu bar. Next, left-click on an exogenous (independent) construct, and move the cursor over the target endogenous (dependent) construct. Now left-click on the endogenous construct, and a path relation (directional arrow) will be inserted between the two constructs. Repeat the same process, and connect all the constructs based on your structural theory. When you finish, the path model should look like the one in Figure 2.15.

FIGURE 2.15 ■ Initial Model



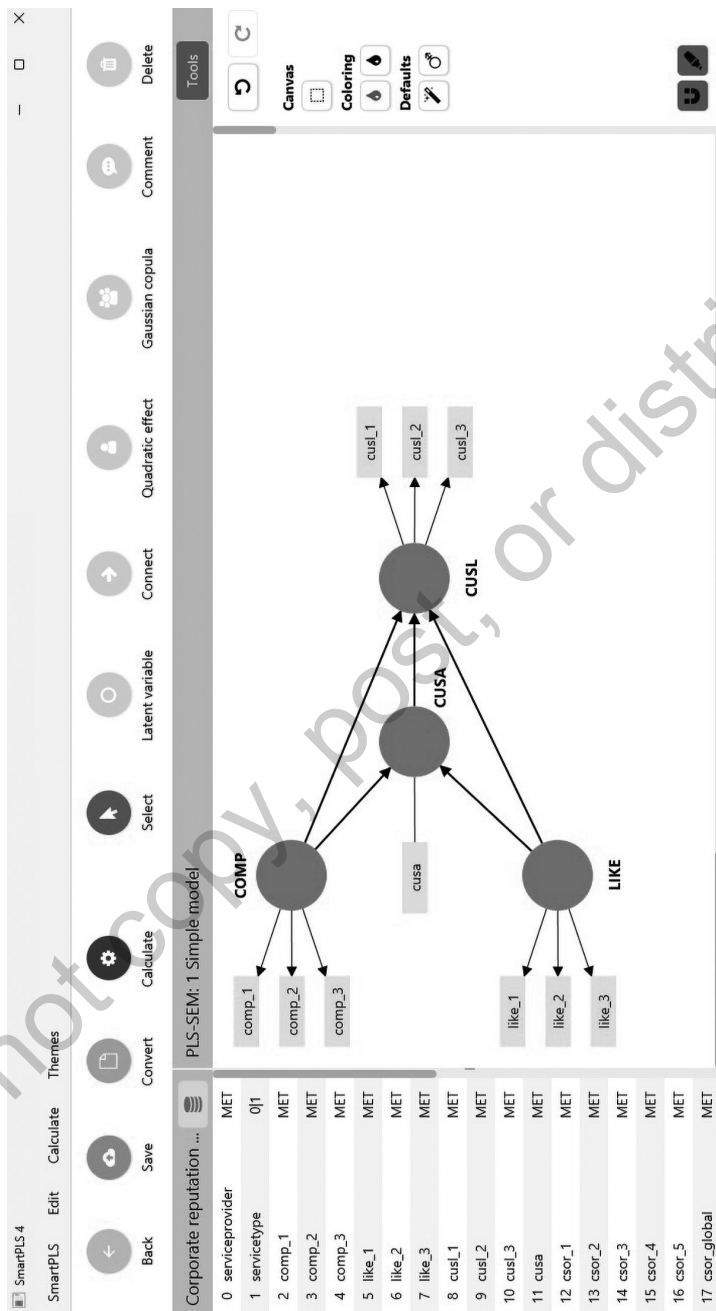
Next, you need to assign indicators to each of the constructs. On the left side of the screen, there is an *Indicators window* that shows all the indicators that are in your data set along with some basic descriptive statistics that you can request by left-clicking on an indicator. Start with the *COMP* construct by dragging the first competence indicator *comp_1* from the *Indicators window* and dropping it on the construct (i.e., left-click the mouse, and hold the button down, then move it until over a construct, then release). After assigning an indicator to a construct, it appears in the *Modeling window* as a yellow rectangle attached to the construct. Assigning an indicator to a construct will also turn the color of the construct from red to blue. Continue until you have assigned all the indicators to the constructs as shown in Figure 2.16.

Go back to the *Select* mode by clicking on the corresponding option in the menu bar. You can now move the indicator around, but it will remain attached to the construct (unless you delete it). By right-clicking on the construct and choosing one of the options under *Align* (e.g., *Align indicators to the top*), you can align the indicator(s). You can also access the align indicators option via the *Tools* on the right-hand side of the *Modeling window*. Please note that you can also move the constructs by left-clicking and dragging them, using the colored horizontal and vertical guidelines to help align them accurately. You can also reposition a construct's label by left-clicking and dragging it. Make sure to save the model by clicking on *Save* in the menu bar.

Right-clicking on selected construct(s) in the *Modeling window* opens a menu with several options. Apart from renaming the constructs, you can, for example, change the measurement model from reflective to formative, and vice versa (*Invert measurement model*), *Hide indicators*, and *Show indicators* of the construct. Additionally, when double-left-clicking on a construct, a different menu opens that allows you to select an indicator weighting mode per construct (i.e., *Automatic*, *Mode A*, *Mode B*, *Equal weights*, and *Pre-defined*) and prespecify the construct reliability of formatively measured constructs. To add a note to your *Modeling window*, left-click on the *Comment* button in the menu bar. Clicking on the right mouse button while the cursor is placed over other elements also opens a menu with additional functions.

Now that we have created our PLS path model, let's explore other options in the *Workspace* view by clicking on *Back* in the menu bar. If you place the cursor in the *Workspace* and right-click on the project name, you can create a new model (e.g., *New PLS-SEM model*, *New REGRESSION model*, *New PROCESS model*, *New CB-SEM model*, and *New GSCA model*), create a new project (*New project*), or import a new data set (*Import data file*). Moreover, you can select the *Copy resource*, *Paste resource*, and *Delete resource* options for projects and models that appear in the *Workspace*. For example, the *Duplicate* option is useful when you would like to modify a PLS path model but want to keep your initial model setup. Using the *Copy resource* option, you can copy and paste a PLS path model from one project to the other.

FIGURE 2.16 ■ Simple Model With Names and Data Assigned



When right-clicking on a project in the *Workspace*, you can choose the *Export project* from the dialog that appears to externally save the project with all its files (i.e., models and data sets) and SmartPLS settings as a zip file in a folder of your choice. This option is useful for sending the project to collaborators as well as for creating a backup. Similarly, you can also directly import such previously exported projects by going to *Files* → *Import project from backup file*. For instance, you can use this option to import the project that includes the PLS-SEM example on corporate reputation. The file name is *Corporate Reputation.zip*. This project is ready to download on your computer system in the download section at <https://www.pls-sem.net>. Download this file, and save it on your computer system. Then, go to *Files* → *Import project from backup file*. SmartPLS allows you to browse your computer and select the downloaded project *Corporate Reputation.zip* for import. Alternatively, for the corporate reputation model example used in this book, you can simply import the sample project *Corporate reputation – PLS-SEM book (primer)* by ticking the checkbox to the left of this label, as discussed earlier in this section.

Summary

- Understand the basic concepts of structural model specification, including mediation, moderation, nonlinear relationships, and the use of control variables.** This chapter includes the first three stages in the application of PLS-SEM. Building on an established theory, prior research, and logic, the model specification starts with the structural model (Stage 1). Each element of the theory represents a construct in the structural model. Moreover, assumptions for the causal relationships between the elements must be considered. The relationships between the constructs are directed (i.e., the arrows linking the constructs go from one construct to the next), but they can also be more complex and contain mediating or moderating relationships. Although these relationships are linear in nature, researchers may also hypothesize and test nonlinear relationships in their PLS-SEM analyses. In addition, researchers frequently specify control variables to account for the impact of other characteristics or phenomena that are not part of the primary theoretical model being tested. The goal of the PLS-SEM analysis is to empirically test the theory or a certain element thereof in the form of the structural model.
- Explain the differences between reflective and formative measurement models, and specify the appropriate measurement model.** Stage 2 focuses on selecting a measurement model for each construct in the structural model to obtain reliable and valid measurements. Generally, there are two types of measurement models: reflective and formative. In reflective measurement models, relations go from the construct to its indicators, suggesting that the measures represent the effects (or manifestations) of an underlying construct

(Continued)

(Continued)

and are therefore highly correlated. In contrast, in formative measurement models, arrows point from the indicators in the measurement model to the constructs. Hence, all indicators together form the construct, and all major elements of the domain must be represented by the selected formative indicators. Because formative indicators represent independent sources of the construct's content, they do not necessarily need to be correlated (in fact, they shouldn't be highly correlated). This measurement-theoretic perspective needs to be distinguished from the way PLS-SEM estimates the model relationships.

- **Comprehend that the selection of the mode of measurement model and the indicators must be based on theoretical reasoning before data collection.** A reflective specification would use different indicators than a formative specification of the same construct. Researchers typically use reflective constructs as target constructs of the PLS path model, whereas formative constructs may be particularly valuable as explanatory sources (i.e., exogenous constructs) or drivers of these target constructs. During the data analysis phase, the measurement mode of the constructs can be empirically tested by using confirmatory tetrad analysis.
- **Explain the advantages of using differentiated indicator weights over sum scores.** Some scholars advocate for using sum scores, which average indicator values, thereby assigning equal weights, over the differentiated weights estimated by PLS-SEM. Although sum scores simplify the model estimation, this approach overlooks measurement error, conceals issues with low-loading indicators, and diminishes the reliability and validity of model estimates. Research has demonstrated that sum scores can result in parameter biases, lower statistical power, and reduced out-of-sample predictive power compared to PLS-SEM. Additionally, sum scores fail to reveal the relative importance of indicators, which is crucial for nuanced analysis. Thus, we recommend sum scores not be used in your research.
- **Understand the nature of higher-order constructs.** Higher-order constructs, also referred to as hierarchical component models, are used to specify a single construct on a more abstract dimension and more concrete subdimensions at the same time. Higher-order constructs have become increasingly popular in research because they offer a means of establishing more parsimonious path models. Researchers often specify and estimate higher-order constructs with two layers of abstraction, also referred to as second-order constructs.
- **Describe the data collection and examination considerations necessary to apply PLS-SEM.** Stage 3 underlines the need to examine your data after they have been collected to ensure that the results from the methods application

are valid and reliable. The primary issues that need to be examined include missing data, suspicious response patterns (straight lining or inconsistent answers), and outliers. Distributional assumptions are of less concern because of PLS-SEM's nonparametric nature. However, because highly skewed data can cause issues in the estimation of significance levels, researchers should ensure that the data are not too far from normal. As a rule of thumb, always remember the garbage in, garbage out rule. All your analyses are meaningless if your data are inappropriate.

- **Learn how to develop a PLS path model using the SmartPLS software.** The first three stages of conducting a PLS-SEM analysis are explained by conducting a practical exercise. We discuss how to draw a PLS path model focusing on corporate reputation and its relationship with customer satisfaction and loyalty. We also explain several options that are available in the SmartPLS software. The outcome of the exercise is a PLS path model drawn using the SmartPLS software that is ready to be estimated.

Review Questions

1. What is a structural model?
2. What is a reflective measurement model?
3. What is a formative measurement model?
4. What is a single-item measure?
5. When do you consider data to be “too nonnormal” for a PLS-SEM analysis?

Critical Thinking Questions

1. How can you decide whether to specify a construct reflectively or formatively?
2. Which research situations favor the use of reflective and formative measures?
3. Explain the difference between multi-item and single-item measures, and assess when to use each measurement type.
4. Create your own example of a PLS path model (including the structural model with latent variables and the measurement models).

(Continued)

(Continued)

5. Why is it important to carefully analyze your data prior to analysis? What problems do you encounter when the data set has relatively large amounts of missing data per indicator (e.g., more than 5% of the data are missing per indicator)?

Key Terms	
Alternating extreme pole responses	Index
Casewise deletion	Indirect effect
Causal indicators	Inner model
Causal links	Kurtosis
Composite indicators	Listwise deletion
Confirmatory tetrad analysis in PLS-SEM (CTA-PLS)	Lower-order components
Control variables	Mean value replacement
Cramér-von Mises test	Mediating effect
Diagonal lining	Missing value treatment
Direct effect	Model comparisons
Effect indicators	Moderation
Endogenous latent variables	Moderator effect
Formative measurement	Multigroup analysis
Fuzzy-set qualitative comparative analysis (fsQCA)	Outlier
Heterogeneity	Pairwise deletion
Hierarchical component models	Reflective measurement
Higher-order component	Scale
Higher-order constructs	Second-order constructs
Higher-order models	Skewness
	Straight lining

Suggested Readings

Becker, J.-M., Cheah, J. H., Gholamzade, R., Ringle, C. M., & Sarstedt, M. (2023). PLS-SEM's most wanted guidance. *International Journal of Contemporary Hospitality Management*, 35(1), 321–346.

Becker, J.-M., Ringle, C. M., & Sarstedt, M. (2018). Estimating moderating effects in PLS-SEM and PLSc-SEM: Interaction term generation*data treatment. *Journal of Applied Structural Equation Modeling*, 2(2), 1–21.

Bollen, K. A. (2011). Evaluating effect, composite, and causal indicators in structural equation models. *MIS Quarterly*, 35(2), 359–372.

Cheah, J.-H., Magno, F., & Cassia, F. (2024). Reviewing the SmartPLS 4 software: The latest features and enhancements. *Journal of Marketing Analytics*, 12, 97–107.

Guenther, P., Guenther, M., Ringle, C. M., Zaefarian, G., & Cartwright, S. (2025). PLS-SEM and reflective constructs: A response to recent criticism and a constructive path forward. *Industrial Marketing Management*, 128, 1–9.

Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., & Thiele, K. O. (2017). Mirror, mirror on the wall: A comparative evaluation of composite-based structural equation modeling methods. *Journal of the Academy of Marketing Science*, 45(5), 616–632.

Hair, J. F., Sarstedt, M., Ringle, C. M., Sharma, N. P., & Liengard, B. D. (2024). Going beyond the untold facts in PLS-SEM and moving forward. *European Journal of Marketing*, 58(13), 81–106.

Hair, J. F., Sharma, P. N., Sarstedt, M., Ringle, C. M., & Liengard, B. D. (2024). The shortcomings of equal weights estimation and the composite equivalence index in PLS-SEM. *European Journal of Marketing*, 58(13), 30–55.

Liu, Y., Chin, W. W., Cheah, J.-H., Hair, J. F., & Lyu, C. (2025). Tackling missing data in PLS-SEM: Strategies and insights for business research. *Journal of Business Research*, 201, 115739.

Lohmöller, J.-B. (1989). *Latent variable path modeling with partial least squares*. Physica.

Memon, M. A., Cheah, J.-H., Ramayah, T., Ting, H., & Chuah, F. (2018). Mediation analysis: Issues and recommendations. *Journal of Applied Structural Equation Modeling*, 2(1), i–ix.

Memon, M. A., Cheah, J.-H., Ramayah, T., Ting, H., Chuah, F., & Cham, T. H. (2019). Moderation analysis: Issues and guidelines. *Journal of Applied Structural Equation Modeling*, 3(1), i–ix.

Memon, M. A., Thurasamy, R., Ting, H., Cheah, J.-H., & Chuah, F. (2024). Control variables: A review and proposed guidelines. *Journal of Applied Structural Equation Modeling*, 8(2), 1–18.

Nitzl, C., Roldán, J. L., & Cepeda-Carrión, G. (2016). Mediation analyses in partial least squares structural equation modeling: Helping researchers to discuss more sophisticated models. *Industrial Management & Data Systems*, 116(9), 1849–1864.

Sarstedt, M., Hair, J. F., Cheah, J.-H., Becker, J.-M., & Ringle, C. M. (2019). How to specify, estimate, and validate higher-order models. *Australasian Marketing Journal*, 27(3), 197–211.

[Continued]

(Continued)

Sarstedt, M., Hair, J. F., & Ringle, C. M. (2023). "PLS-SEM: Indeed a silver bullet"—retrospective observations and recent advances. *Journal of Marketing Theory & Practice*, 31(3), 261–275.

Sharma, P. N., Sarstedt, M., Ringle, C. M., Cheah, J.-H., Herfurth, A., & Hair, J. F. (2024). A framework for enhancing the replicability of behavioral MIS research using prediction oriented techniques. *International Journal of Information Management*, 78, 102805.

Yuan, K.-H., Wen, Y., & Tang, J. (2020). Regression analysis with latent variables by partial least squares and four other composite scores: Consistency, bias and correction. *Structural Equation Modeling: A Multidisciplinary Journal*, 27(3), 333–350.

Visit the companion site for this book at <https://www.pls-sem.net/>.