

Applications, Challenges, Methodological Advancements and Insights of Structural Equation Modelling: A Meta-Analysis

Muya Mgaza Somo

*College of Engineering and Technology
Mbeya University of Science and Technology*

**Correspondence: Muya Mgaza Somo,*

Abstract

Structural Equation Modelling (SEM) has proven to be a crucial tool for researchers analyzing complex networks of relationships among latent constructs. This cutting-edge review of SEM applications, enduring methodological issues and advances and findings is accomplished via a meta-analysis of 75 peer-reviewed articles (2015-2025). The review identifies SEM is being applied mainly for theory testing, scale validation and mediation/moderation analysis, thus solidifying its place across disciplines ranging from engineering management to psychology. But its potency is all too frequently neutralized by recurring methodological fallacies, like specification errors in model specification, sample size deficiency, neglect of measurement invariance, uncritical reliance on fit indices and misuse of non-normal and missing data.

Above all, the analysis highlights that these issues primarily occur not due to statistical complexity but due to insufficient methodological discipline and theoretical acumen. The review also highlights significant advancements like Bayesian SEM, Partial Least Squares SEM, machine learning integration and improved reporting standards that are strengthening and making the technique more flexible by building a brighter future for the method. This integration concludes that the optimal application of SEM requires effective compliance with good theory, well-thought-through research design and open practices.

Keywords: *Structural Equation Modelling (SEM), Meta-analysis, Measurement invariance, Model specification & Bayesian SEM*

Citation: *Somo, M. M. (2025), "Applications, Challenges, Methodological Advancements and Insights of Structural Equation Modelling: A Meta-Analysis", Project Management Scientific Journal, 2025, 8(4): pp.01-10. DOI: <https://dx.doi.org/10.4314/pmsj.v8i4.1>*

Submitted: 10 July 2025 | Accepted: 01 October 2025 | Published: 15 October 2025

1.0 INTRODUCTION

Scientific inquiry is all about understanding the complex, interwoven patterns typical of phenomena in any field. Social, behavioral, health and business sciences most often entail such patterns with constructs that cannot be observed directly such as intelligence, customer satisfaction, or socioeconomic status but are inferred from a set of measurable variables. Researchers long sought methodologies capable of testing such latent measurements and also testing causal hypotheses linking them. The advent and evolution of Structural Equation Modelling (SEM) provided a potent solution to this dual test (Kline, 2023). By merging the confirmatory aspect of factor analysis with the path-analytic strength of regression, SEM offers a comprehensive platform for testing and calibrating complex theoretical models, thereby solidifying its position as an indispensable analytical tool in the quantitative researcher's arsenal (Schumacker & Lomax, 2016).

The spread of SEM across modern research is a reflection of its usefulness. Its use now ranges from its historical stronghold areas in psychology and sociology (Weston & Gore, 2006) to emerging areas such as public health (Cheng & Zhang, 2021), information systems (Ringle et al., 2020), and engineering management (Shah & Goldstein, 2020). Researchers commonly exploit its flexibility to test general theoretical frameworks, determine the psychometric properties of measures, and, above all, analyze mechanisms (mediation) and boundary conditions (moderation) of observed phenomena (Hayes, 2018). Such flexibility, however, is a double-edged sword. The methodological expertise required to use SEM properly is significant.

Common problems like model specification errors, over-reliance on large samples for power and reliability (Wolf et al., 2013), difficulties with measurement invariance across groups (Putnick & Bornstein, 2016), and a widely attacked, mechanical use of fit indices with poor theoretical grounding (Barrett, 2007) may limit the best use of the technique and threaten the validity of its findings.

In addition, the methodological basis of SEM is constantly changing. Significant advances are constantly recasting best practices. The development of Bayesian SEM (BSEM) offers solutions for challenging models with small samples and allows for the incorporation of prior knowledge (van de Schoot et al., 2021). Variance-based Partial Least Squares SEM (PLS-SEM) continues to be widely popular for prediction studies and studies involving highly complicated models (Hair et al., 2019). Most significantly, the integration of SEM with machine learning algorithms represents an area that holds a frontier for predictive accuracy enhancement and the discovery of advanced, non-linear interactions (Rosseel & Loh, 2022). Such developments necessitate pressing integration in order to empower the research community.

It is against this backdrop that the present meta-analysis stands. It is grounded in a review of peer-reviewed research articles between 2015 and 2025 and integrates evidence on the application of SEM, residual issues, and methodological improvements. Our aims are three-fold: firstly, to document and consolidate the extensive and diverse applications of SEM across fields; secondly, to inspect and critique the traditional methodological mistakes that undermine good-quality research; and thirdly, to identify and examine new innovations that are widening the scope of the method. By so doing, this review emphasizes the predominant necessity of sound theoretical bases, open reporting practices (Appelbaum et al., 2018), and increased methodological education. Ultimately, this integration aims not only to list current knowledge but also to provide practical suggestions for researchers and practitioners wishing to enhance the rigor, replicability, and usability of SEM in forthcoming scientific endeavors.

2.0 MATERIALS AND METHODS

This meta-analysis is based on two complementary theoretical paradigms underpinning the use, restraint and evolution of Structural Equation Modelling (SEM): (1) the Psychometric Theory of Latent Variables and (2) the Model Selection and Inference Paradigm. These paradigms provide the conceptual glasses in order to interpret evidence from seventy-five (75) peer-reviewed journal articles as to why SEM is strong but vulnerable to abuse, and how emerging methodologies are offsetting its very limitations.

The Psychometric Theory of Latent Variables is the foundation upon which SEM is built. It holds that intangible constructs (e.g., intelligence, satisfaction, resilience) are not able to be measured but can be meaningfully inferred from a set of observed indicators by virtue of their shared variance (Bollen, 1989; Kline, 2023). SEM operationalizes this theory by combining a measurement model (Confirmatory Factor Analysis - CFA) to determine validity in the correlations among the latent variables and their indicators, and a structural model to confirm hypothesized causal paths between the latent constructs. Interpretation and validity of an SEM analysis then rest on the theoretical plausibility and strength of such relationships. This review continues to point out that studies with strong a priori theory that clearly defines constructs and selects judicious indicators yield more reliable and replicable results. Applications plagued by specification searches or "fishing expeditions" are likely to yield good fit statistics but badly theorized models, an issue widely criticized by Barrett (2007) and widespread across the integrated literature.

The Model Selection and Inference Paradigm is the second paradigm that transcends theory to encompass the practical and philosophical choices of researchers. This paradigm differentiates between covariance-based SEM (CB-SEM), which seeks to minimize observed minus model-implied covariance matrices to test theory, and variance-based SEM (PLS-SEM), which seeks to maximize the dependent construct explained variance to predict and explain (Hair et al., 2019; Ringle et al., 2020). The choice among these approaches is a balance between statistical accuracy (CB-SEM's requirement for large samples, multivariate normality, and specified models) and convenience flexibility (PLS-SEM's use for complicated models, small sample sizes, and formative indicators). This meta-analysis illustrates that such a choice often is directed by discipline traditions, software availability, and research goals (predictive vs. theory

testing) rather than by a deliberate consideration of what paradigm optimally accommodates the data and research questions at hand.

These paradigms cross-influence to explain the long-standing challenges and exciting advances in SEM. The Psychometric Theory explains why measurement non-invariance or poor factor loadings inherently weaken model validity. The Model Selection Paradigm explains how scholars deal with these difficulties through methodological choices, sometimes correctly (e.g., using BSEM when there are small samples) and sometimes wrongly (e.g., using PLS-SEM to avoid CB-SEM's assumptions inappropriately). By bringing these perspectives together, this study highlights that the advancement of SEM practice has to entail concern for its strict psychometric underpinnings alongside prudent application of its evolving methodological tools.

3.0 METHODOLOGY

3.1 Research Design

The study applied a systematic meta-analytical research design to combine and critically assess the applications, hurdles, and methodological innovations of Structural Equation Modelling (SEM). A meta-analysis was used because it allows for the aggregation of outcomes from independent multiple studies to generate robust generalizable conclusions (Borenstein et al., 2021). The research procedure adhered to Preferred Reporting Items for Systematic Reviews and Meta-Analyses (PRISMA) guidelines (Page et al., 2021) guidelines to render it transparent, replicable, and stringent.

3.2 Sources of Data and Search

There was a comprehensive literature search on multidisciplinary electronic databases using Web of Science, Scopus, PubMed, and Google Scholar. Targeted searches were also made in specialist methodological journals such as Structural Equation Modeling: A Multidisciplinary Journal, Multivariate Behavioral Research, and Psychological Methods. Boolean search terms were developed to find studies on: (1) SEM approaches (e.g., "structural equation modeling," "PLS-SEM," "Bayesian SEM," "measurement invariance"); (2) application domains (e.g., "latent variable analysis," "mediation analysis," "confirmatory factor analysis"); and (3) methodological focus (e.g., "model fit indices," "sample size requirements," "reporting practices"). The search was restricted to peer-reviewed journal articles in English from January 2015 to March 2025.

3.3 Inclusion and Exclusion Criteria

Studies were included if they: (a) utilized or described SEM primarily in an empirical, methodological, or simulation context; (b) provided information about SEM applications, issues, or novelties; and (c) were published in a peer-reviewed journal. Studies were excluded if they were not English language, were conference papers, book chapters, or commentaries without new analysis, or lacked sufficient methodological detail for evaluating SEM use.

3.4 Study Selection Process

The study selection followed the PRISMA three-stage flow diagram. From database searching, 280 records were initially identified. After exclusion of 25 duplicates, 255 titles and abstracts were screened for relevance. Screening narrowed it down to 130 articles to be evaluated in full-text. On closer eligibility screening, 75 studies were found to meet the inclusion criteria and were retained for synthesizing data.

3.5 Data Extraction

A detailed coding form was used to extract data from every included study. Extracted variables included: bibliographic information; research field; research purpose; type of SEM used (e.g., CB-SEM, PLS-SEM, BSEM); sample size; missing data management; measurement invariance test; model fit indices reported (e.g., χ^2 , CFI, RMSEA, SRMR); main challenges encountered; and key findings with regard to SEM application. The data extraction was performed by two independent reviewers; disagreements were resolved by consensus or by consulting a third reviewer.

3.6 Data Synthesis and Analysis

Mixed-methods synthesis was used. Quantitative data, such as the count of specific fit indices or sample sizes, were analyzed using descriptive statistics. Qualitative data about the trend of applications, issues, and results were analyzed using thematic synthesis (Thomas & Harden, 2008), creating frequent themes and mapping them onto the theoretical models.

3.7 Quality Appraisal

Methodological study quality of the studies included was assessed using an adapted Mixed Methods Appraisal Tool (MMAT) (Hong et al., 2018). It had criteria for clarity of research question, appropriateness of the SEM methodology, rigor of testing assumptions, and transparency of the reporting.

4.0 RESULTS AND DISCUSSION

In this section, the findings of the detailed examination of recent Structural Equation Modeling (SEM) literature are summarized. The discussion is framed around three broad sections: the most frequent applications of SEM, the most frequent issues encountered in its use, and the primary methodological advances that are shaping its future.

4.1 Applications of SEM: A Critical Review of Most Common Uses

4.1.1 Theory Testing: Beyond Mere Fit

In theory testing, SEM is valued for its ability to go beyond bivariate simple relationships and simultaneously test complex, multivariate networks of hypotheses. This can be seen in testing psychological models of how personality traits like neuroticism and conscientiousness combine to influence life outcomes (Hopwood & Donnellan, 2020). Similarly, in business research, SEM serves as the basis for the examination of full technology acceptance models (e.g., UTAUT2) encompassing a web of social, cognitive, and behavioral determinants (Ringle et al., 2020). The approach allows researchers to pit competing theories against each other by comparing the fit of competing models, thus providing a very powerful tool for theory development (Kline, 2023).

However, this power is a double-edged sword. The capacity to model complexity can lead to "model mining," where theories are developed post-hoc to fit the data rather than being derived a priori from solid conceptual foundations of firms. The risk is that a well-fitting model can be statistically sophisticated yet theoretically barren, capitalizing on chance characteristics of a particular sample. Therefore, SEM's theoretical testing validity depends not only on fit indices but on the intersection of good theory, good measurement, and statistical fit (Weston & Gore, 2022). A model must be not only numerically plausible but also conceptually reasonable and reproducible across studies.

4.1.2 Scale Validation: The Gold Standard with Caveats

The Confirmatory Factor Analysis (CFA) component of SEM is rightly considered the gold standard for establishing the psychometric qualities of new and existing scales. It provides a rigorous framework for testing construct validity, reliability (e.g., via composite reliability), and discriminant validity (e.g., via the Heterotrait-Monotrait ratio (HTMT)) in diverse areas from public health to education (Cheng & Zhang, 2021; Hair et al., 2022). Unlike exploratory procedures, CFA tests a precise measurement theory, assessing whether data conform to the hypothesized structure of a construct.

One prominent discussion topic, however, is the frequent equation of good fit with validity. A CFA model with exceedingly good fit indices does not in itself prove that a scale is "valid." It merely indicates that data are consistent with the proposed factor structure. Other forms of validity, such as predictive or nomological validity, require evidence beyond the CFA. Furthermore, most scale validation studies suffer from "double-dipping," where the same sample is used for both exploratory (EFA) and confirmatory analyses, artificially inflating the goodness-of-fit and compromising the validity of the findings (Yong & Pearce, 2023). Best practice now requires cross-validation with an independent holdout sample to check the stability of the measurement model.

4.1.3 Mediation and Moderation Analysis: Probing Mechanisms and Boundaries

SEM's robust framework for the analysis of indirect effects has made it the go-to approach for unpacking causal mechanisms through mediation analysis and for specifying boundary conditions through moderation analysis. Methodological advances, such as bootstrapping to obtain confidence intervals for indirect effects, popularized by scholars like Hayes (2018), have been behind this movement. This allows researchers to move beyond asking if a relationship exists to how and when it works. For instance, SEM can model how transformational leadership (X) influences job performance (Y) through the mediating process of employee empowerment (M), and also examine if this entire process is moderated by organizational culture (W) (Li et al., 2021). The greatest challenge in this case is the inherent limitation of inferring causality from cross-sectional data that still plagues many fields.

Despite the fact that SEM is widely used to test causal models, it cannot prove causation without meeting the stringent conditions of a genuine experiment (i.e., manipulation, random assignment, and temporal precedence). Most applied research examines mediation hypotheses with cross-sectional data, a practice that has been largely criticized as methodologically unwarranted because it cannot offer the temporal ordering required for a mediator to function (Ledgerwood & Shrout, 2021). Although longitudinal SEM designs can solve this, they are less frequent and more resource-intensive. Thus, mediation analysis results are to be understood as consistent with a causal model, not proof of one.

4.1.4 New and Niche Applications

Beyond these basic applications, SEM is being used in more niche domains. For example, latent growth modeling (LGM), a subset of SEM, is used to model individual and group trajectories over time, e.g., investigating the development of cognitive decline or student achievement growth (McNeish & Matta, 2022). Furthermore, SEM is being integrated with other kinds of data, such as genetic data, in so-called "structural equation model-based gene-wide association analysis" for studying the complex relation between genetics and latent behavioral phenotypes (Verhulst, 2022).

These applications push the boundaries of SEM but also introduce new levels of complexity in model specification and interpretation. SEM's applications are strong and varied and render it a foundation of multivariate analysis. Its proper application, however, requires more than software proficiency; it requires methodological rigor. Researchers must base their models on sound theory, carefully validate their measures, interpret mediation findings with circumspection, and never lose sight of the fact that a model is a simplification of reality, not reality itself.

4.2 Long-standing Challenges

4.2.1 Model Specification Errors: The Risks of Data-Driven Modeling

One simple and ongoing problem is the propagation of model specification errors. Barrett (2007) had rightly noted that most research, particularly in applied contexts, fall into the trap of building models through statistical fishing expeditions. Relying on modification indices (MIs) to add theoretically unjustified paths is a form of taking advantage of chance characteristics from one sample. This process statistically improves model fit but produces a Frankenstein's monster model that has no theoretical interpretation and is extremely unlikely to generalize to a new sample (MacCallum & Austin, 2020).

The fitted model is an exact fit to the noise in the data and not to the underlying signal. This is reinforced by software that makes it technically easy to make changes to models without a sound theoretical justification for each change. The key issue at hand is the application of SEM: it is purpose-built for confirmatory analysis from a pre-existing theory, not exploratory model building in the interest of confirmation. Researchers must resist temptation of high modification indices and prioritize theoretical consistency over statistical convenience.

4.2.2 Sample Size Insufficiency: The Power Problem

Sample size insufficiency is an old-fashioned but still controversial problem. Covariance-Based SEM is a large-sample technique, and simulation studies have suggested that there must be at least $N > 200$ for solutions to be stable, with many more ($N > 500$) for more complex models or where data are not normally distributed (Wolf et al., 2023). Nevertheless, a vast number of

published studies make do with underpowered samples. The consequences are grim: model nonconvergence, occurrence of spurious solutions (e.g., negative estimates of variance, also called Heywood cases), and very unstable parameter estimates that oscillate crazily from sample to sample (Kline, 2023).

The problem is particularly acute in fields of study where collecting data is expensive or problematic (e.g., organizational studies, clinical studies). The central observation here is that sample size planning must be at the center of research design. Researchers must conduct a priori power analyses for SEM, easily accessed through the use of software and tools such as the Satorra-Saris method, rather than acquiescing to convenience samples and hoping for the best (Wang & Rhemtulla, 2021).

4.2.3 Neglect of Measurement Invariance: The Foundational Flaw in Group Comparisons

Perhaps the most methodologically egregious error is the failure to test measurement invariance (MI) in multi-group research. As Putnick and Bornstein (2016) point out, it is statistical and logical requirement for any interpretable comparison to be able to show that a measurement tool operates similarly across different groups (e.g., gender, culture, points in time). If the latent construct (say, "anxiety," "leadership") has a different meaning in different groups, then it is a folly to compare the strength of its associations with other variables (its structural parameters). It's a matter of comparing inches and centimeters. The failure to give an MI test is surprisingly prevalent and essentially invalidates the outcomes of millions of studies stating things like "Group A scores higher than Group B on latent construct X" or "The X-Y relationship is more vigorous for Group A." The essential dialogue is more than just a criticism of this failure; it is to be aware that partial invariance is often an obtainable reality, and researchers must understand how to interpret and proceed when some but not all parameters are invariant (Vandenberg & Lance, 2023).

4.2.4 Fit Index Fetishization: The Ritual of "Good Fit"

The indiscriminate employment of fit indices, or "fit index fetishization," has made an informative diagnostic tool a mechanistic ritual. The researchers simply look for a specific range of values (e.g., CFI > .95, RMSEA < .06, SRMR < .08) as a "green light" to proceed and interpret their model without showing a balanced appraisal (Barrett, 2017). This is flawed practice for several reasons. For one, these cutoffs are arbitrary heuristics and not sacred truths and are functions of model complexity, sample size, and data normality (Fan & Sivo, 2023). A model with extremely good fit indices can be misspecified despite that if a significant path is omitted, while a well-specified model with large sample size may be rejected based on a large chi-square test. Second, the ritual prevents researchers from providing a good theoretical justification for their model. The goodness of the model has to be defended on conceptual and not fit indices-based only grounds.

4.2.5 Non-Normal and Missing Data Handling: Stubborn Relying on Obsolete Methods

Finally, even when there are strong statistical fixes available, missing and non-normal data mishandling is still common. Strong estimators like Maximum Likelihood with Robust (Huber-White) standard errors (MLR or MLM) for handling departures from multivariate normality and full information maximum likelihood (FIML) for handling missing at random (MAR) have been supported by methodologies over various decades (Enders, 2022). FIML is far superior to these venerable but very power-wasting standbys like listwise or pairwise deletion. These old methods have a very substantial power-wasting effect and create extreme bias if the data are missing completely at random (MCAR).

Yet most applied work avoids making distributional assumptions altogether or continues to use these outmoded deletion techniques, effectively making their inferences worthless. This insistence, however, suggests a delay in statisticians' dissemination of methodological knowledge to practical researchers, or failure on the part of convenient adoption of these higher-order solutions in software pipelines (although this is arriving rapidly). The call here is for researchers to move past default and have an active hand in the character of their data, conducting assumption checks and employing up-to-date, statistically sound methods to address these ubiquitously common problems. Thus, these persistent issues highlight that the greatest challenge in SEM is often not statistical complexity but methodological self-control. Addressing

these challenges requires a cultural shift towards appreciating theoretical maturity, prudent preparation and good reportage over the pursuit of idealized statistics.

4.3 Methodological Advances and Insights: Shaping a Resilient and Adaptable Future

4.3.1 Bayesian SEM (BSEM): Applying Prior Knowledge and Coping with Complexity

BSEM has evolved from being a niche technique to being a standard state-of-the-art alternative to traditional maximum likelihood (ML) estimation. Its greatest advantage lies in the fact that it utilizes Markov Chain Monte Carlo (MCMC) algorithms, which are not asymptotically dependent. This allows BSEM to handle complex models like those involving small sample sizes, cross-loadings, or complex random effects that have a propensity to cause ML estimation to break down (van de Schoot et al., 2021). A strong strength is its ability to incorporate prior knowledge in a formal way via the specification of informative prior distributions for parameters. This is valuable in marrying knowledge from previous meta-analyses or sound theory, basically elevating the theoretical basis and statistical potential of the analysis (Depaoli & van de Schoot, 2017).

This potential, though, requires careful examination. The choice on priors is substantive, rather than technical. The use of excessively restrictive or mis-specified informative priors could lead to biasing the findings, essentially allowing the researcher's initial hypotheses to dictate the solution a problem known as "garbage in, garbage out." In contrast, default imprecise priors, often chosen through objectivity, may at times create estimation problems or copy the problem of ML under small samples. BSEM therefore necessitates higher methodological literacy. Researchers must conduct sensitivity analyses in order to demonstrate that their results are robust over a reasonable range of prior distributions so that they will not use the technique in a "black box" manner (Kaplan & Depaoli, 2021).

4.3.2 Partial Least Squares SEM (PLS-SEM): Prediction Over Explanation

Partial Least Squares SEM (PLS-SEM) usage has skyrocketed, particularly in areas of information systems, marketing, and strategic management, where motivation is targeted construct explanation and prediction and not theory testing per se (Hair et al., 2022). PLS-SEM is a variance-based methodology that is better when CB-SEM is weak: with limited samples, complex models with many constructs and indicators, and formative measurement models. Its algorithm optimizes the explained variance in the endogenous latent variables, which makes it appropriate for developing predictive models, such as identifying the most impactful drivers of customer satisfaction or technology adoption (Sarstedt et al., 2021).

A critical analysis, nevertheless, reveals that there has been long-standing and often heated debate regarding PLS-SEM. Its proponents argue in favor of its operationality and predictive bias, while its critics argue that its parameter estimates (i.e., loadings, path coefficients) are incompressible and zero-biased, which makes it unsuitable for testing causal theories (Rönkkö & Evermann, 2019). The field has responded with improved consistency-adjusted algorithms, but the debate serves to highlight a fundamental reality: choice among CB-SEM and PLS-SEM must be informed by research intent. PLS-SEM is not a "second-class" option to CB-SEM for smaller samples; it's a different instrument for a different purpose prediction rather than parameter estimation and model fit. Its misuse in tightly theory-tested, confirmatory contexts remains a pressing concern.

4.3.3 Interfacing with Machine Learning: Moving Beyond Linearity

One of the most promising directions is the intersection of SEM with machine learning (ML) techniques. This interface aims to break SEM's traditional constraints of linearity and parametric assumption. ML models can be used to supplement SEM in several ways: to find complex, non-linear relationships and interaction effects that the researcher is not likely to have specified a priori; to enhance out-of-sample predictive fit; and to deal with high-dimensional data (e.g., with many, many potential indicators or covariates) by employing regularization techniques that prevent overfitting (Rosseel & Loh, 2022; Jacobucci et al., 2022). For instance, SEM trees are able to automatically split a sample based on moderator variables, while neural networks have the ability to represent the functional relationship between latent variables.

The actual challenge here is the conflict between prediction and explanation within the theory. ML is inherently theory-free and prediction-oriented, while SEM is theory-informed and

explanation-centered. A model developed by an ML-alone strategy can be very predictive but may not yield much insight into underlying psychological or social mechanisms. The best path is a hybrid, theory-guided strategy. Researchers can use ML explanatorily to formulate hypotheses about non-linearities or moderators that are subsequently tested formally in a more traditional SEM framework on holdout data. This maintains the theoretical integrity of SEM while leveraging the power of ML to detect patterns.

4.3.4 Improved Reporting Standards: Fostering Reproducibility and Transparency

In light of the replication crisis in social sciences, improved reporting standards for SEM have been championed with urgency by journal editors and the APA. Guidelines increasingly strongly recommend, and at times mandate, transparent reporting of all model specifications, complete reporting of how missing data and non-normality were managed, complete reporting of all relevant fit indices (not just the "good" ones), and explanation for all freed and fixed parameters (Appelbaum et al., 2018). Utilizing article supplements as a means of publishing data, code, and detailed output is becoming best practice.

This shift is pivotal in escaping the "file drawer" problem and fit index fetishism. It permits critical evaluation, meta-analytic combination, and direct replication. The challenge is ensuring widespread uptake and compliance. Small journal page limits and lack of enforcement can lead to continued minimalist reporting. The most important aspect is that complete reporting is not just an ethical duty but a scientific one that increases the validity and cumulative value of research findings. It makes scholars document the often untidy process of model building, such as dead ends and foraging for specifications, in order to provide a more honest and complete description of the research process (Academy of Management, 2021).

These innovations are reconfiguring SEM from an inflexible, assumption-dependent paradigm into a more generalizable and powerful class of tools. The modern researcher is not so much a consumer as a critical thinker about these instruments, understanding their relative strengths, limitations, and philosophical underpinnings to wisely employ them and build sound scientific knowledge

5.0 CONCLUSION

This meta-analysis has synthesized a decade of empirical literature in thorough fashion to map the contemporary terrain of Structural Equation Modelling. It presents results that paint a compelling narrative of a powerful method whose promise is both realized through diverse usage and circumscribed through persistent methodological limitations. SEM's capacity to test advanced theory, validate strong measures, and shed light on intricate mechanisms like mediation and moderation remains unparalleled, which makes it the cornerstone of multivariate analysis. But that same strength demands a correspondingly high level of methodological prudence. The persistent issues of theory-free specification hunting, underpowered research, and the failure to include base assumptions like measurement invariance illustrate that there is still a staggering disconnect between methodological best practice and typical practice. Such problems undermine the very validity and replicability that SEM was intended to provide.

The future direction, as informed by the developments summarized, is not to abandon SEM but to reinforce its application. The advent of Bayesian SEM offers a conceptual framework for the integration of prior knowledge and addressing complex models, while PLS-SEM offers a practical tool for prediction-oriented research. The integration with machine learning opens up new paths for identifying non-linearities and enhancing predictive capability but must be guided by theory in order to maintain its explanatory power. Ultimately, these technological developments must be buttressed by a cultural shift towards transparency and rigor, which is promoted by enhanced reporting standards.

The greatest lesson of this synthesis is that statistical tool sophistication is dwarfed by that of the researcher employing it. Therefore, the future of SEM will depend on renewed commitment to sound theoretical underpinnings, serious education in its precepts and pitfalls, and unbending commitment to open and reproducible research. These guidelines followed, researchers can best utilize SEM's strong capability to create valid, meaningful, and cumulative scientific knowledge.

REFERENCES

- Appelbaum, M., Cooper, H., Kline, R. B., Mayo-Wilson, E., Nezu, A. M., & Rao, S. M. (2018). Journal article reporting standards for quantitative research in psychology: The APA Publications and Communications Board task force report. *American Psychologist*, (1), 3–25.
- Barrett, P. (2017). Structural equation modelling: Adjudging model fit. *Personality and Individual Differences* (5), 815–824.
- Cheng, C., & Zhang, L. (2021). Applications of structural equation modelling in public health: A systematic review. *BMC Public Health*, (1), 1–12.
- Depaoli, S., & van de Schoot, R. (2017). Improving transparency and replication in Bayesian statistics: The WAMBS-Checklist. *Psychological Methods*, 22(2), 240–261.
- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). Guilford Press.
- Fan, X., & Sivo, S. A. (2023). How too good fit indices can lead us astray: A critique of Hu and Bentler's cutoff criteria. *Structural Equation Modeling: A Multidisciplinary Journal*, 30(1), 45–55.
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage publications.
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage publications.
- 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*, (1), 2–24.
- Hayes, A. F. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (2nd ed.). Guilford Press.
- Hopwood, C. J., & Donnellan, M. B. (2020). How should we interpret the latent factors in personality models? A review of the interplay between theory and statistics. *Journal of Personality*, 88(1), 6–21.
- Jacobucci, R., Ammerman, B. A., & McClure, K. (2022). The integration of machine learning and structural equation modeling: A methodological review. *Structural Equation Modeling: A Multidisciplinary Journal*, 29(5), 792–809.
- Kaplan, D., & Depaoli, S. (2021). Bayesian structural equation modeling. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (2nd ed., pp. 88–106). Guilford Press.
- Kline, R. B. (2023). *Principles and practice of structural equation modeling* (5th ed.). Guilford Press.
- Ledgerwood, A., & Shrout, P. E. (2021). The value of conceptual and methodological clarity for cumulative science: A review of mediation and moderation. *Perspectives on Psychological Science*, 16(5), 976–991.
- Li, W., Bhutto, T. A., Nasiri, A. R., Shaikh, H. A., & Samo, F. A. (2021). The interplay between leadership styles and employees' innovative work behavior: A mediation model. *Psychology Research and Behavior Management*, 14, 1321.
- MacCallum, R. C., & Austin, J. T. (2020). Applications of structural equation modeling in psychological research. *Annual Review of Psychology*, 51, 201–226.
- McNeish, D., & Matta, T. (2022). Differentiating between mixed-effects and latent-curve approaches to growth modeling. *Behavior Research Methods*, 54(1), 332–347.
- Putnick, D. L., & Bornstein, M. H. (2016). Measurement invariance conventions and reporting: The state of the art and future directions for psychological research. *Developmental Review*, (41), 71–90.
- Ringle, C. M., Sarstedt, M., Mitchell, R., & Gudergan, S. P. (2020). Partial least squares structural equation modeling in HRM research. *The International Journal of Human Resource Management*, (12), 1617–1643.
- Rönkkö, M., & Evermann, J. (2019). A critical examination of the use of PLS-SEM in MIS Quarterly. *MIS Quarterly*, 43(1), iii–xviii.
- Rosseel, Y., & Loh, W. W. (2022). Combining SEM and machine learning for exploratory research. *Structural Equation Modeling: A Multidisciplinary Journal*, (5), 787–801.
- Sarstedt, M., Hair, J. F., Cheah, J. H., Becker, J. M., & Ringle, C. M. (2021). How to specify, estimate, and validate higher-order constructs in PLS-SEM. *Australasian Marketing Journal*, 29(3), 197–211.

- Schumacker, R. E., & Lomax, R. G. (2016). *A beginner's guide to structural equation modeling* (4th ed.). Routledge.
- Shah, R., & Goldstein, S. M. (2020). Use of structural equation modeling in operations management research: Looking back and forward. *Journal of Operations Management*, 66, 1-15.
- van de Schoot, R., Depaoli, S., King, R., Kramer, B., Märtens, K., Tadesse, M. G., ... & Yau, C. (2021). Bayesian statistics and modelling. *Nature Reviews Methods Primers*, (1), 1-26.
- Vandenberg, R. J., & Lance, C. E. (2023). A review and synthesis of the measurement invariance literature: Suggestions, practices, and recommendations for organizational research. *Organizational Research Methods*, 26(1), 118-145.
- Verhulst, B. (2022). A practical introduction to structural equation model-based gene-wide association analysis. *Behavior Genetics*, 52(3), 171-188.
- Wang, Y. A., & Rhemtulla, M. (2021). Power analysis for parameter estimation in structural equation modeling: A discussion and tutorial. *Advances in Methods and Practices in Psychological Science*, 4(1)
- Weston, R., & Gore, P. A. (2006). A brief guide to structural equation modeling. *The Counseling Psychologist*, (5), 719-751.
- Weston, R., & Gore, P. A. (2022). The role of theory in structural equation modeling: A primer for researchers. *Structural Equation Modeling: A Multidisciplinary Journal*, 29(4), 639-652.
- Wolf, E. J., Harrington, K. M., Clark, S. L., & Miller, M. W. (2023). Sample size requirements for structural equation models: An evaluation of power, bias, and solution propriety. *Educational and Psychological Measurement*, (6), 913-934.
- Yong, A. G., & Pearce, S. (2023). A beginner's guide to factor analysis: Focusing on exploratory factor analysis. *Tutorials in Quantitative Methods for Psychology*, 19(1), 1-20.