

A typology of validity: content, face, convergent, discriminant, nomological and predictive validity

Journal of Trade
Science

155

Weng Marc Lim

*Sunway Business School, Sunway University, Sunway City, Malaysia;
School of Business, Law and Entrepreneurship, Swinburne University of Technology,
Hawthorn, Australia and
Faculty of Business, Design and Arts,
Swinburne University of Technology - Sarawak Campus, Kuching, Malaysia*

Received 22 March 2024
Revised 21 May 2024
Accepted 7 June 2024

Abstract

Purpose – Research serves to elucidate and tackle real-world issues (e.g. capitalizing opportunities and solving problems). Critical to research is the concept of validity, which gauges the extent to which research is adequate and appropriate in representing what it intends to measure and test. In this vein, this article aims to present a typology of validity to aid researchers in this endeavor.

Design/methodology/approach – Employing a synthesis approach informed by the 3Es of expertise, experience, and exposure, this article maintains a sharp focus on delineating the concept of validity and presenting its typology.

Findings – This article emphasizes the importance of validity and explains how and when different types of validity can be established. First and foremost, content validity and face validity are prerequisites assessed before data collection, whereas convergent validity and discriminant validity come into play during the evaluation of the measurement model post-data collection, while nomological validity and predictive validity are crucial in the evaluation of the structural model following the evaluation of the measurement model. Additionally, content, face, convergent and discriminant validity contribute to construct validity as they pertain to concept(s), while nomological and predictive validity contribute to criterion validity as they relate to relationship(s). Last but not least, content and face validity are established by humans, thereby contributing to the assessment of substantive significance, whereas convergent, discriminant, nomological and predictive validity are established by statistics, thereby contributing to the assessment of statistical significance.

Originality/value – This article contributes to a deeper understanding of validity's multifaceted nature in research, providing a practical guide for its application across various research stages.

Keywords Content validity, Face validity, Convergent validity, Discriminant validity, Nomological validity, Predictive validity, Construct validity, Criterion validity, Validity

Paper type General review

1. Introduction

The bedrock of scientific inquiry lies in research integrity, a principle that underscores the ethical, rigorous and transparent pursuit of knowledge (Moher *et al.*, 2020). Integral to upholding research integrity is the adherence to practices that ensure the reliability (consistency) and validity (accuracy) of research outcomes (Grey *et al.*, 2019). It is through this lens that this article approaches the concept of *validity*, a critical metric that gauges *the degree to which research adequately and appropriately represents what it aims to measure and test*.

© Weng Marc Lim. Published in *Journal of Trade Science*. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) licence. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and non-commercial purposes), subject to full attribution to the original publication and authors. The full terms of this licence may be seen at <http://creativecommons.org/licences/by/4.0/legalcode>

Data availability statement: Data sharing not applicable – no new data generated.



Journal of Trade Science
Vol. 12 No. 3, 2024
pp. 155-179
Emerald Publishing Limited
e-ISSN: 2755-3957
p-ISSN: 2815-5793
DOI 10.1108/JTS-03-2024-0016

Validity, in its essence, is not only a methodological consideration but also an indicator to the *applicability* and *authenticity* of research findings to real-world contexts. Validity assumes a central role in affirming the *strength* and *relevance* of conclusions drawn from data, serving as a linchpin that *connects theoretical constructs and relationships to empirical observations*. This distinction is crucial, as it ensures that the insights gleaned are not only *statistically significant*—indicating that the *results are unlikely to have occurred by chance*—but also *substantively significant*, reflecting *real-world importance* and *meaningful practical implications* of these findings. In this capacity, validity underpins the integrity of research by guaranteeing that findings are not merely mathematical artifacts but are also grounded in and reflective of theoretical and empirical reality. Through this dual contribution, validity ensures that scholarly investigations yield insights that are both methodologically sound and deeply impactful, bolstering the *credibility* and *utility* of research in addressing complex real-world issues. The importance of validity therefore extends beyond the confines of *academic discourse*, influencing *organizational strategies*, *policy-making* and *societal advancements* (Figure 1).

However, the pursuit of validity is fraught with challenges. Despite widespread acknowledgment of its importance, the practice of ensuring validity in research is often

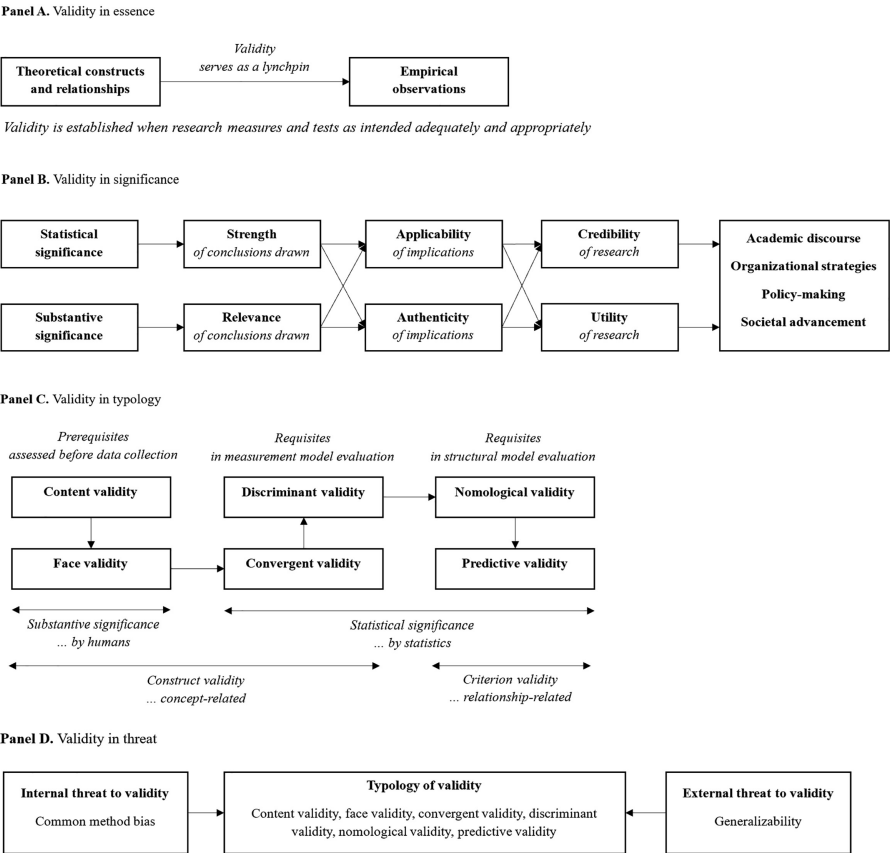


Figure 1.
Overview of validity

Source(s): Author's own illustration

compromised by a fragmented understanding of its multifaceted nature (Kock *et al.*, 2024). Validity is characterized by a plethora of types, each addressing different aspects of the research process—from the conceptualization and operationalization of constructs to the interpretation of findings. This diversity, while enriching, can lead to a piecemeal approach to validation efforts, wherein researchers might focus on certain types of validity at the expense of others.

This fragmented approach is further compounded by the absence of a cohesive framework to guide researchers in the systematic establishment of validity across different stages of their studies. The lack of a unified typology not only hampers the comprehensive assessment of validity but also undermines the integrity and impact of research. It is against this backdrop that the present article ventures to bridge the gap by presenting a comprehensive typology of validity. This typology encompasses content, face, convergent, discriminant, nomological and predictive validity, each serving a distinct yet interconnected role in bolstering the validity of research endeavors.

While this article focuses on the typology of validity, it complements existing methodological guides that address a range of research issues. Noteworthy, previous works have offered guidance on the basics of research philosophy and paradigms (Lim, 2023), the best practices in survey research and scale deployment (Hardy and Ford, 2014; Haws *et al.*, 2023; Hulland *et al.*, 2018; Robinson, 2018), the statistical and methodological issues commonly raised in the review process (Green *et al.*, 2016) and the widely-used structural equation modeling (Becker *et al.*, 2023; Memon *et al.*, 2017; Vaithilingam *et al.*, 2024). Other works have provided guidance on common method bias and data collection issues, offering insights into mitigating these threats (Memon *et al.*, 2023; Podsakoff *et al.*, 2024). The typology of validity presented in this article adds a valuable dimension to this arsenal of guides by offering a structured framework that specifically addresses the multifaceted nature of validity, ensuring that research findings are both methodologically robust and practically relevant.

Delineating the applications of each type of validity, this article aims to equip researchers with a holistic understanding of *how* and *when* to establish the validity of their work. In doing so, this article seeks to elevate the practice of research beyond the acknowledgment of validity as a conceptual ideal and toward its application as a fundamental aspect of research integrity.

While primarily oriented toward business disciplines such as management and marketing, the principles of this typology are also applicable across various domains. Business research itself encompasses a wide array of methodologies and intersects with numerous fields, highlighting the broad relevance of these validity types.

For example, in management, content validity ensures that a leadership assessment tool comprehensively covers all facets of leadership behavior, while face validity ensures that the tool appears sensible to those taking it. Convergent and discriminant validity are assessed through statistical analysis to ensure that the tool coherently measures leadership distinctively from other constructs like management style. Nomological validity involves examining the tool's placement within a framework, ensuring that leadership scores correlate appropriately with related constructs such as management style and organizational success. Predictive validity would assess if leadership scores forecast organizational success, thus overlapping with the nomological network. Similarly, in marketing, content validity ensures that a customer satisfaction survey covers all relevant aspects of customer experience, face validity ensures that respondents find the survey relevant, convergent and discriminant validity confirm that the survey coherently measures satisfaction as distinct from other constructs like loyalty, nomological validity checks if satisfaction fits within a network of related constructs such as loyalty (e.g. advocacy, purchase) and predictive validity assesses if satisfaction predicts loyalty. Extrapolating these examples to areas where management and

marketing intersect with other fields, the principles of validity are not confined to business alone; for instance, the management example can be applied in education (e.g. deans, principals), and the marketing example can be relevant in hospitality and tourism (e.g. hotel guest satisfaction, tourist experience).

Providing a structured framework for understanding and implementing validity ensures that researchers from diverse, quantitatively-oriented backgrounds (e.g. business, communication, education, psychology, sociology) can uphold research integrity and produce findings that are both statistically and substantively significant. Thus, this typology is not just an academic contribution but also a practical guide that is poised to enhance the quality and impact of empirical research across disciplines.

2. What validity is and is not

The concept of validity has evolved significantly since its early inception in the fields of psychology and education, where it initially served as a measure of the efficacy of psychological tests (Haertel and Herman, 2005). Historically, the roots of validity can be traced back to the early 20th century, with seminal contributions from scholars such as Terman (1919) and Thorndike (1910, 1918), who were among the pioneers in applying validity in the context of intelligence testing and educational assessment. Thorndike's works (1910, 1918), in particular, emphasized the importance of ensuring that tests genuinely measured what they purported to measure, laying the groundwork for subsequent discussions on validity.

Over time, the definition of validity has broadened and deepened, transcending its initial focus to encompass a wide range of applications in social sciences. The mid-20th century saw a significant expansion in the conceptualization of validity, with scholars like Messick (1989, 1995) alongside Cronbach and Meehl (1955) contributing to a finer-grained understanding of validity that went beyond the correlation between test scores and criteria. Messick (1989), in particular, advocated for a unified view of validity that integrated considerations of content, criterion and construct validity within a single framework, while Messick (1995), in his later work, extended construct validity, albeit confusingly, to encompass content, substantive, structural, generalizability, external and consequential aspects—this confusion will be addressed and reconciled through this article.

In empirical research, validity applies to the entire research process. In this regard, it is crucial to delineate what validity is and is not in this context:

- (1) *What validity is.* Validity is a multifaceted concept that assesses the adequacy and appropriateness of inferences drawn from research data. Validity encompasses several types, which this article unpacks later, and thus, is about the legitimacy and meaningfulness of inferences, ensuring that research findings genuinely reflect the phenomena under study.
- (2) *What validity is not.* Validity is not one-dimensional and thus cannot be assessed in isolation or through a single method. Validity is not about the consistency of measurements (that's reliability), nor is it merely about the statistical correlation between measures. Validity is not merely an inherent property of a measurement instrument (e.g. questionnaire) or a test (e.g. path analysis), but also a feature of the inferences made from data. Therefore, validity cannot be assumed but must be rigorously established and argued for in each specific research context.

Understanding validity in this way allows researchers to appreciate its complexity and the critical role it plays in ensuring research integrity. Recognizing what validity is and what it is not, researchers can better navigate the challenges of constructing and validating their

research, thereby contributing to the advancement of robust and relevant scientific knowledge. This perspective also underscores the dynamic nature of validity, encouraging ongoing reflection and adaptation in response to evolving research paradigms and societal needs [1].

3. A typology of validity

The essence of ethical, rigorous and transparent research lies in its ability to adequately represent and appropriately measure the phenomena it seeks to investigate (Moher *et al.*, 2020), a quality fundamentally anchored in the concept of validity. Validity, with its multifaceted dimensions, serves as a critical metric for evaluating the integrity of research findings. Understanding and applying the various types of validity necessitate a staged approach, reflecting the sequential phases of the research process (Lim, 2022; Lim and Koay, 2024; Rao *et al.*, 2024). This staged approach not only facilitates a systematic assessment of validity but also aligns the validation efforts with the natural progression of empirical inquiry.

At the outset of research, before any data collection commences, it is essential to establish the foundational aspects of validity. Content validity and face validity fall within this preliminary stage, serving as prerequisites for ensuring that the instruments and their measures (items, scales) are aptly designed [2], to capture and measure the constructs of interest. *Content validity* scrutinizes the extent to which the elements of a measure encompass the construct, based on expert judgment. Concurrently, *face validity* provides an initial check on the appropriateness of the measure, assessing whether it appears to measure what it is intended to, based on subjective judgment. Establishing validity at the outset of research, particularly through content and face validity, lays the groundwork for *construct validity* (as they pertain to constructs), contributing to *substantive significance*, as they are established by *humans*. Ensuring that instruments and measurement items (indicators) are thoughtfully designed to capture and measure the constructs of interest, researchers set a strong foundation for the subsequent phases of research.

As the research advances to data collection and subsequent analysis, the focus of validity assessment shifts toward the measurement model. It is within this post-data collection phase that convergent validity and discriminant validity become paramount. *Convergent validity* assesses the degree to which multiple indicators that are theoretically related to the same construct indeed converge or correlate highly with one another. In contrast, *discriminant validity* evaluates the distinctiveness of different constructs within the study, ensuring that measures that are not supposed to be related are indeed unrelated or minimally correlated. These types of validity are connected to *statistical significance*, as they are established through *statistics*, thereby reinforcing *construct validity* (as they relate to constructs). Researchers can therefore enhance the validity of their measurement model by rigorously establishing convergent and discriminant validity, contributing significantly to the robustness of their research.

The final stage of the research process involves the evaluation of the structural model, where nomological validity and predictive validity play crucial roles. *Nomological validity* examines the coherence of the conceptual model [3], ensuring that the relationships between constructs align with theoretical expectations grounded by established theories or schools of thoughts. *Predictive validity*, on the other hand, concerns the ability of the constructs to accurately predict outcomes, providing a direct measure of the practical utility of the research findings. The rigorous assessment of nomological and predictive validity at this stage by means of statistical significance also significantly contributes not only to *criterion validity* (when outcomes can be explained and predicted) but also to the *substantive significance* of the research when they are convincingly extrapolated (e.g. argument and evidence that the

recommendation is needed and that this recommendation will work in reality). Ensuring that theoretical relationships are logically sound and that constructs possess predictive power, these forms of validity enhance the *credibility* and *strength* of the research. This careful validation of the structural model also underpins the study's contribution to the body of knowledge and underscores its *applicability*, *authenticity*, *relevance* and *utility* in addressing real-world issues.

Adopting a staged approach to validity allows researchers to systematically address each aspect of validity in alignment with the appropriate phase of the research process (Table 1). This methodical progression ensures that validity is not an afterthought but a pivotal consideration integrated throughout the research, from conceptualization to conclusion.

3.1 Content validity

Content validity is a cornerstone in the research validation process, ensuring that instruments like questionnaires and their measurement items—whether adopted, adapted, or newly developed—*adequately* and *appropriately* represent the constructs they aim to measure. The establishment of content validity, particularly through *expert evaluation*, is paramount in ensuring that the research instrument is not only *suitable (conceptually)* but also provides *comprehensive coverage (operationally)* of the intended constructs from an expert perspective. This validation process ensures that the *measurement items in the instrument are aligned with the underpinnings of the constructs* that the study intends to capture and measure, thereby contributing to *construct validity* (as it pertains to constructs).

What content validity is and is not. At its core, content validity entails the assessment by subject-matter experts to verify that every aspect of the construct is reflected in or represented by the measurement items. This process extends beyond superficial appraisal, requiring a profound grasp of the construct's peculiarities. Unlike face validity, which assesses the instrument's outward appearance and intuitive alignment with the construct, content validity is focused on the representativeness and thoroughness of the measurement items. It is the initial step in affirming the instrument's conceptual and operational integrity, distinct from other validity forms that evaluate performance, such as construct relationships or predictive power. Laying a conceptually sound and operationally comprehensive foundation, content validity paves the way for subsequent validation efforts.

Establishing content validity (How). The journey to content validity often involves a *pretest*, where a cyclical process of critique and refinement, guided by feedback from a panel of experts, takes place. These experts, drawn from relevant academic and professional fields, possess the depth of knowledge necessary to assess the instrument's alignment with the construct. They scrutinize each measurement item to ensure the instrument's breadth and avoid the inclusion of extraneous elements. While a minimum of two experts is a good starting point, the ideal number varies with the construct's complexity and the study's breadth. This process can be iterative in nature, continuing until reaching data saturation—where no new modifications emerge from expert input—thereby ensuring a comprehensive and robust evaluation. Alternatively, a card sorting exercise can be undertaken, where experts, serving as judges, categorize and label each item without being provided with construct names. This process continues with a new set of judges until a desired correct hit ratio, typically above 90%, is achieved (Moore and Benbasat, 1991).

Timing for establishing content validity (When). As the foundational layer of the research validation process, content validity precedes all other forms, setting the stage before data collection. This preemptive approach ensures that the research instrument is both conceptually sound and operationally apt for the intended constructs. Expert consultation at this nascent phase is not just a procedural formality but also a critical step in refining the instrument, ensuring its readiness for both face validity assessment and the broader research

(continued)

Validity → Characteristic ↓	Content validity	Face validity	Convergent validity	Discriminant validity	Nomological validity	Predictive validity
What it is	Validity ensuring that the indicators (items) adequately operational coverage) and appropriately conceptually sound represent the constructs they aim to measure	Validity ensuring that the items (indicators) appear (initial impression, intuitive understanding) to measure what they intend to measure (clear, comprehensible, relevant)	Validity ensuring that items statistically relate to the same construct	Validity ensuring that indicators (items) statistically distinct and not highly correlated to indicators statistically relating to another construct	Validity ensuring that constructs and relationships are theoretically consistent (coherent), if not theoretically explainable (logical), within the model	Validity ensuring that constructs forecast outcomes in a theoretically consistent, if not in a theoretically explainable, manner
How to establish	- Expert evaluation by academic and/or industry experts, minimum two experts and up to the point of data saturation in feedback - Experts could also conduct a card sorting exercise, categorizing and labeling items without construct names until a correct hit ratio above 90% is achieved	- Subjective evaluation by potential participants from the target population, minimum two participants and up to the point of data saturation in feedback - Cross-check by asking respondents to explain what they understand by the question to avoid miscomprehension	For reflective indicators - Factor loading (≥ 0.708), or report as is if lower with justification (evolving context, newness) - Average variance extracted (≥ 0.50), or reassess in consultation with the literature on whether (i) the indicators are indeed reflective of the construct, (ii) the indicators are still relevant (e.g. current content or era), or (iii) the indicators may be better suited as formative elements that contribute to the construct, or report as is with explanation (evolving context, newness) For formative indicators - Regression analysis (bootstrap), where the statistical contribution of indicator weight (≥ 0.50) and significance is considered alongside the substantive contribution of the indicator to the theoretical breadth and depth of the concept - Redundancy analysis, where the formative construct correlates more than 50% (≥ 0.708) with the alternative (reflective) construct, or report as is with explanation (evolving context, newness)	For reflective indicators - The square root of the AVE for each construct is greater than the construct's highest correlation with any other construct For formative indicators - Variance inflation factor (VIF) of < 3 for a maximum, < 5 for a moderate maximum, and < 10 for a liberal maximum threshold For reflective and formative indicators - HTMT values < 0.85 for conceptually distinct constructs and < 0.90 for conceptually similar constructs - Correlation values < 0.70 for conceptually distinct constructs and < 0.80 for conceptually similar constructs	For co-variance-based structural equation modeling (CB-SEM) - Goodness-of-fit index (GFI), adjusted goodness-of-fit index (AGFI), comparative fit index (CFI), normed-fit index (NFI), and Tucker-Lewis index (TLI) ≥ 0.90 - Root mean square error of approximation (RMSEA) < 0.08 - Relative chi-square (χ^2/df) < 3 - For partial least squares structural equation modeling (PLS-SEM) - Standardized root mean square residual (SRMR) < 0.08 For predictive capability of the model - Coefficient of determination (R^2) at high values: Comparable, if not better than alternative models or existing studies - Predictive relevance (Q^2) > 0 - Constructs and relationships are theoretically consistent (as hypothesized), if not theoretically explainable (when not as hypothesized)	Coefficient, showing direction (+/-) and strength of the relationship Statistical significance, showing that the relationship is unlikely to occur due to chance - p-value: 95% confidence level ($p < 0.05$) or 99% confidence level ($p < 0.01$) - t-value: 95% confidence level ($t > 1.96$ (two-tailed; no direction specified) or $t > 1.645$ (one-tailed; direction specified) or 99% confidence level ($t > 2.576$ (two-tailed; no direction specified) or $t > 2.326$ (one-tailed; direction specified)) - Confidence interval: Does not include zero (+ and - or - and -) to be considered significant, narrower intervals indicate greater precision Effect size, showing the magnitude of the relationship and how important it is in practical terms - Cohen's f^2 for regression and Pearson's r for correlations, with standard benchmarks categorizing effect sizes as small ($f^2 \geq 0.02$, $r \approx 0.1$ to 0.3), medium ($f^2 \geq 0.15$, $r \approx 0.3$ to 0.5), or large ($f^2 \geq 0.35$, $r > 0.5$)

Table 1.
Typology of validity

Table 1.

Validity →Characteristic ↓	Content validity	Face validity	Convergent validity	Discriminant validity	Nonological validity	Predictive validity
When to establish	Pretest before data collection (1st)	Pilot study before data collection (2nd)	Main study after data collection as part of measurement model evaluation	Main study after data collection (4th) as part of measurement model evaluation	Main study after data collection (5th) as part of structural model evaluation	Main study after data collection (6th) as part of structural model evaluation
Significance contribution	Substantive significance	Substantive significance	Statistical significance	Statistical significance	Statistical significance	Statistical significance
Validity contribution	Construct validity	Construct validity	Construct validity	Construct validity	Criterion validity	Criterion validity
Note(s): Construct validity = Validity pertaining to construct (concept-focused). Criterion validity = Validity relating to explanation and prediction of outcomes (relationship-focused)						
Source(s): Author's own illustration						

endeavor. Establishing content validity first enables researchers to solidify the instrument's conceptual and operational framework, enhancing the overall validity of the research.

3.2 Face validity

Following the establishment of content validity, face validity serves as the next critical checkpoint in the research validation process. Face validity assesses whether the research *instrument and its items appear to measure the constructs they are intended to measure*. This form of validation is important because it addresses the *initial impressions* and *intuitive understandings* of the instrument by the *target population*, which can significantly influence participation rates and the quality of responses. The logical progression from content to face validity is rooted in the need to first ensure that the instrument adequately and appropriately covers the conceptual and operational aspects of the constructs (content validity) before gauging its *clarity, comprehensibility and relevance* (face validity). This sequence ensures that the instrument is not only theoretically robust but also accessible to participants, fostering better engagement and more valid responses.

What face validity is and is not. Face validity is concerned with the apparent clarity, comprehensibility and relevance of the research instrument to those with no specialist knowledge of the underlying constructs. It relies on the subjective perceptions of potential participants to gauge the instrument on these criteria. Face validity is not about the theoretical correctness or the empirical performance of the instrument, which are covered by content validity and subsequent forms of validity. It is more about the immediate, intuitive response the instrument elicits, making it a crucial step in ensuring the instrument's approachability and user-friendliness.

Establishing face validity (How). The assessment of face validity typically involves a *pilot study*, soliciting feedback from a sample of potential participants who mirror the profile of individuals targeted in the main study. This feedback focuses on the measurement items' clarity, comprehensibility and relevance. A rule of thumb for the number of participants to consult is at least two, but the process often benefits from a larger pool to ensure a comprehensive assessment. The goal is to reach a point of data saturation, where no new significant feedback is obtained, indicating that the instrument's face validity is well established. Achieving a consensus among this sample group on the instrument's face validity can often necessitate multiple rounds of feedback and revision, aiming for clarity and intuitive understanding across the board. A good cross-check is to ask respondents to explain what they understand by the question, thus avoiding miscomprehension (Hardy and Ford, 2014).

Timing for establishing face validity (When). Face validity is assessed after content validity has been firmly established and before the main data collection phase begins. This timing ensures that the instrument is not only theoretically sound but also appears clear, comprehensible and relevant to potential participants. Establishing the face validity of the instrument following the establishment of content validity allows researchers to refine their tools and ensure they are both conceptually robust and operationally relevant for participants. This step is pivotal in preparing the instrument for the main study, ensuring that it is not only valid in a theoretical sense but also accessible and comprehensible to those who will be providing the data.

3.3 Convergent validity

With the foundational aspects of content and face validity established, attention shifts to convergent validity in the post-data collection phase of the research process. Convergent validity evaluates *the extent to which different indicators (items) statistically relate to the same construct*. This form of validity is essential for confirming that various indicators of a construct cohesively measure the same underlying concept. Establishing convergent validity is particularly vital in research involving complex constructs that can be operationalized in

multiple ways, ensuring that these diverse indicators are not diverging but rather coalescing around the same construct.

What convergent validity is and is not. Convergent validity is concerned with the empirical evidence demonstrating that multiple measures of the same construct are indeed related to the construct, as they theoretically should be. Yet, it is not sufficient for measures to be highly associated to the construct in any manner; these associations must be theoretically justified and empirically supported through rigorous statistical analysis. Furthermore, unlike content and face validity, which rely on qualitative assessments like expert judgment or subjective evaluation, convergent validity demands quantitative validation through empirical data analysis. Moreover, convergent validity is distinct from reliability, which assesses the internal consistency of responses across indicators within a single measure. Therefore, convergent validity focuses on the convergence among different indicators of the same construct, ensuring they all reflect the same theoretical dimension.

Establishing convergent validity (How). The assessment of convergent validity depends on the nature of the indicators within the measurement model.

For *reflective indicators*, convergent validity is assessed through factor loadings in exploratory factor analysis (EFA) for newly developed items and confirmatory factor analysis (CFA) for adopted, adapted and newly developed items. Factor loadings represent the degree to which the variance in an observed indicator (like a survey question) is accounted for by its association with an underlying construct (such as a psychological trait). High factor loadings indicate that a significant portion of the variance in each indicator can be explained by the construct it is intended to measure, affirming that the indicators are indeed reflecting the same construct. The construct is therefore considered to explain the variance in the indicators, not the other way around, in a reflective measurement model.

A factor loading of 0.7 or higher is considered strong, indicating that a substantial portion of the variance in the indicator is explained by the construct. The 0.7 threshold is generally recommended as it implies that about 50% (since 0.7 squared is 0.49) of the variance in the indicator is accounted for by the construct, signifying a strong relationship (Hair *et al.*, 2014). More specifically, squaring the factor loading of an indicator on a construct gives us the proportion of variance in that indicator that is explained by the construct. This is often referred to as the “explained variance” or “R-squared” value for that indicator. A factor loading of 0.7, when squared, becomes 0.49, implying that approximately half ($\approx 50\%$) of the variance in the indicator is explained by the construct. Ideally, this should be more than 50%, to demonstrate that the majority of the indicator’s variance is indeed attributable to the construct, reinforcing the validity of the measurement. Therefore, an ideal factor loading should be at least 0.708 (since 0.708 squared is 0.501 > 0.50), a threshold that receives support from the literature (Hair *et al.*, 2021, 2022). However, in exploratory research or with new constructs, lower loadings can sometimes be acceptable. In such instances, it is best to report the factor loading(s) as is with evidence showing that the concept is new (as in the case of scale development studies) or remains evolving in the sample population, while also discussing potential reasons for the lower loadings and suggesting directions for future research to refine and validate the measures further.

Another pivotal metric in the assessment of convergent validity for reflective indicators is the average variance extracted (AVE), which gauges the extent of variance captured by a construct relative to the variance attributable to measurement error. The established threshold for AVE is 0.5 or higher (Fornell and Larcker, 1981), indicating that the construct explains over half of the variance in its indicators, thus providing strong evidence of convergent validity. Achieving an AVE above this threshold not only affirms that a substantial proportion of indicator variance is construct-driven, but it also establishes a solid foundation for the discriminant validity of the construct by ensuring that the construct is

sufficiently distinct and not overly influenced by measurement error or overlap with other constructs.

In practice, achieving the AVE threshold can pose challenges, particularly with complex constructs or when indicators demonstrate diverse associations with the construct. This may signal a need to reassess whether the indicators are indeed reflective of the construct, whether the indicators are still relevant (e.g. current context or era), or whether the indicators may be better suited as formative elements that contribute to the construct [4]. This re-evaluation is a thoughtful process that involves revisiting the theoretical underpinnings of the construct and its indicators. It should be noted that this reassessment is not about adjusting to meet certain criteria but about ensuring the measurement model aligns with the true nature of the construct as supported by theory and empirical evidence. Such a critical review might lead to a re-operationalization of the construct, transforming the measurement model to more accurately reflect the construct's dimensions. This rigorous approach not only enhances the validity of the measurement but also contributes to a deeper understanding and more precise operationalization of the construct. More importantly, transparently reporting the process toward achieving (or not achieving) the AVE threshold, including the theoretical and empirical rationales for any adjustments made, is crucial. This transparency in the validation process underscores the scholarly integrity of the research and invites constructive dialog within the academic community. This also highlights the iterative nature of research as a progressive endeavor that, through challenges, fosters deeper insights and reinforces the validity of the constructs under study.

For *formative indicators*, where each indicator contributes to forming the construct rather than reflecting it, convergent validity is approached differently. In formative measurement models, convergent validity can be established through a two-step approach: a regression analysis and a redundancy analysis. The initial step, involving regression analysis, allows for the evaluation of each indicator's contribution and significance to the formative construct, ensuring that each indicator's unique aspect is accounted for in the construct's formation. The subsequent step, redundancy analysis, serves to validate the convergence of the formative construct with an alternative, conceptually identical or similar construct, typically measured reflectively. This two-step approach ensures a thorough validation by first establishing the contribution and relevance of individual indicators to the formative construct and then confirming that the construct as a whole converges (with the other conceptually identical or similar construct), thereby reinforcing the construct's validity.

In the first step, a regression analysis can be employed to robustly evaluate the weights and significance of the indicators where the formative construct is regressed on its indicators. Bootstrapping, a resampling method, provides a more reliable estimate of the weights and their statistical significance by repeatedly sampling from the data set and calculating the statistic across these samples. While larger absolute values (≥ 0.5) of weights indicate a stronger contribution of the indicator to the construct, indicators with smaller weights or those that are not statistically significant might still be retained if they are deemed theoretically important—a practice that is supported by notable scholars (Hair *et al.*, 2021, 2022). This is because each formative indicator contributes distinct content to the construct, and its exclusion could overlook essential facets of the construct's conceptual domain. Therefore, the decision to retain indicators is informed not only by their statistical contribution of weight and significance but also by their substantive contribution to the breadth and depth of the construct, ensuring the construct's comprehensive representation.

In the second step, a redundancy analysis (Chin, 1998) can be performed to shed light on the extent to which a construct with formative indicators correlates with an alternative construct (e.g. a construct with reflective indicators) of the same concept. This approach necessitates advance planning in the research design (e.g. questionnaire development for primary data and data pooling for secondary data) to ensure that both formative and

alternative measures of the concept are included. Alternative measures could come in the form of a reflective construct in the case of primary data (Hair *et al.*, 2021, 2022) and a similar construct in the case of secondary data (Houston, 2004). The work of Cheah *et al.* (2018) evidences that a global (single) indicator [5] suffices as a measure for an alternative construct in the case of small samples (≤ 300) while multiple indicators work better for an alternative (reflective) construct in the case of larger samples (> 300). Leveraging on the same threshold logic for factor loading as discussed above, convergent validity is established when the correlation between the formative and alternative constructs of the concept is 0.708 or higher, implying that the formative construct explains more than 50% of the variance in the alternative, conceptually identical or similar construct—a practice that receives support from the literature (Hair *et al.*, 2021, 2022).

Timing for establishing convergent validity (When). Convergent validity assessment is strategically positioned in the research validation timeline immediately following the data collection phase and after the establishment of content and face validity. This sequencing ensures that the foundational qualitative aspects of the research instrument—its comprehensive coverage of the constructs (content validity) and its clarity and apparent relevance to participants (face validity)—are firmly in place before examining the quantitative validation of how well the constructs are measured. Situating convergent validity after content and face validity but before discriminant validity, researchers can systematically build a robust case for their instrument's validity. The post-data collection phase is ideal for convergent validity assessment because it requires empirical data to perform the necessary statistical analyses, such as EFA, CFA and AVE for reflective indicators and regression and redundancy analyses for formative indicators, to evaluate the strength of the relationships between each indicator and its associated construct. The successful establishment of convergent validity sets a strong foundation for the subsequent assessment of discriminant validity, where the distinctiveness of the constructs is examined. This logical progression ensures that the constructs are not only well-defined and understood at the outset but are also empirically robust, enhancing the overall validity of the research.

3.4 Discriminant validity

Following the establishment of convergent validity, discriminant validity becomes the next critical focus in the construct validation sequence. Discriminant validity assesses the extent to which measures of different constructs are *distinct* and *not highly correlated*, ensuring that each construct is unique and captures phenomena not represented by other constructs in the study. This form of validity is pivotal for clarifying the conceptual boundaries between constructs, especially in complex research models where constructs may be theoretically related yet empirically distinct.

What discriminant validity is and is not. Discriminant validity is the evidence that constructs are unrelated in the empirical data. Discriminant validity is established when the measures of different constructs do not demonstrate high correlations among each other, signifying that the constructs are distinct and capture different dimensions. Discriminant validity is the opposite of convergent validity: while convergent validity focuses on the unity among measures of the same construct, discriminant validity emphasizes the separateness and specificity of different constructs.

Establishing discriminant validity (How). Like convergent validity, the assessment of discriminant validity is also contingent on the nature of the indicators within the measurement model.

For *reflective constructs*, discriminant validity can be established by comparing the square root of the AVE for each construct with the correlations between that construct and all other constructs in the model (Fornell and Larcker, 1981). Discriminant validity is supported when

the square root of the AVE for each construct is greater than the construct's highest correlation with any other construct, indicating that the construct is more closely related to its own indicators than to those of any other construct.

For *formative constructs*, the assessment process involves ensuring each indicator uniquely contributes to the construct without undue collinearity (redundancy). High collinearity among indicators can inflate standard errors and lead to misinterpretation of indicator weights, potentially affecting the construct's validity. The variance inflation factor (VIF) is commonly used to assess collinearity, with values above 10 (or even 5 or 3, in more conservative cases) indicating critical collinearity (Becker *et al.*, 2015; Hair *et al.*, 2021, 2022). Researchers may need to consider reducing collinearity through methods such as eliminating or merging indicators or constructing higher-order constructs, especially if unexpected sign changes in indicator weights occur [6], which could confound the interpretation of the formative model.

For both *reflective and formative constructs*, the heterotrait-monotrait (HTMT) ratio of correlations can be used to assess discriminant validity (Henseler *et al.*, 2015). An HTMT value less than 0.85 generally indicates adequate discriminant validity for conceptually distinct constructs, signifying sufficient distinctiveness between them. While a threshold of less than 0.90 has been suggested for conceptually similar constructs, this should be applied with caution and justified within the specific theoretical context of the study. Additionally, inspecting a correlation matrix can offer additional insights, with correlations between different constructs ideally being less than 0.70 for conceptually distinct constructs and 0.80 for conceptually similar constructs to support discriminant validity [7]. This traditional approach serves as an initial check, complementing more sophisticated methods like the HTMT ratio.

Timing for establishing discriminant validity (When). The assessment of discriminant validity is ideally conducted after convergent validity has been confirmed and is typically the final step in the construct validation process before proceeding to test the research hypotheses. This timing ensures that each construct has already been shown to be internally coherent (via convergent validity) and that the research instrument is ready for rigorous hypothesis testing. Establishing discriminant validity at this stage, researchers solidify the distinctiveness and specificity of their constructs, laying the groundwork for accurate and meaningful interpretation of the relationships among constructs.

3.5 Nomological validity

As the validation process progresses beyond the establishment of construct validity through content, face, convergent and discriminant assessments, attention turns to nomological validity. This type of validity, a specific form of criterion validity that is *relationship-focused*, is crucial for affirming the *theoretical consistency*, if not *theoretical explainability*, within the model of the study. Nomological validity evaluates the extent to which the predicted relationship(s) among constructs, as delineated by the theoretical underpinnings of the study, is (are) empirically supported.

What nomological validity is and is not. Nomological validity centers on the coherence and logical consistency of the inter-construct relationship(s) within a model (or framework), adhering to established theories or schools of thoughts underpinning proposed hypothesis(es). This goes beyond identifying association(s) or correlation(s) among constructs; rather, it demands that association(s) or correlation(s) conform(s) to a theoretically grounded network of relationship(s), thus ensuring the constructs' integration into a coherent and logical theoretical structure. While demonstrating nomological validity can reinforce construct validity by evidencing that constructs behave as theoretically expected, its essence lies beyond evaluating individual constructs. Instead, nomological

validity focuses on the dynamics among multiple constructs, assessing their interactions and alignment with theoretical predictions to ensure a coherent and theoretically grounded model. In this regard, the establishment of nomological validity is fundamentally anchored in the scrutiny of the nomological network, which represents the complex net or web of theoretical relationships among constructs. This examination is pivotal for ensuring that the empirical evidence aligns with theoretical expectations. It is important to note, however, that while comprehensive nomological networks can encompass multiple constructs and relationships (Cronbach and Meehl, 1955; Preckel and Brunner, 2017) [8], the minimum criterion for establishing such a network involves the examination of a single relationship between two constructs. This foundational requirement highlights the principle that even a singular theoretically grounded and empirically supported relationship can serve as a testament to the model's nomological validity [9]. This approach underscores the importance of each relationship's theoretical justification and empirical validation, regardless of the network's complexity. Starting with the fundamental unit of analysis—a single relationship between two constructs—researchers can incrementally build and validate more elaborate models, ensuring each step is firmly rooted in both theory and empirical evidence.

Establishing nomological validity (How). Establishing nomological validity involves developing a conceptual model grounded in a literature review, detailing the hypothesized relationships among constructs. This model is empirically tested using appropriate statistical techniques such as structural equation modeling (SEM).

Covariance-based SEM (CB-SEM) provides a suite of goodness-of-fit indices that offer a comprehensive assessment of model fit, contributing to the substantiation of nomological validity. Key goodness-of-fit indices include the goodness-of-fit index (GFI), adjusted goodness-of-fit index (AGFI), comparative fit index (CFI), normed-fit index (NFI) and Tucker-Lewis index (TLI), with a threshold of at least 0.90 indicating a good fit between the proposed model and the observed data (Hair *et al.*, 1998; Lim, 2015). Additionally, the root mean square error of approximation (RMSEA) should ideally be less than 0.08, and the relative chi-square (χ^2/df) should be less than 3 (Bollen, 1989; Kline, 1998; Lim, 2015) to suggest a good model fit. Whereas, partial least squares SEM (PLS-SEM) uses the standardized root mean square residual (SRMR) method, whereby an SRMR value of less than 0.08 suggests a good fit (Henseler *et al.*, 2016).

Other statistical methods that can contribute to nomological validity include the coefficient of determination (R^2) and predictive relevance (Q^2). The R^2 value for each endogenous (dependent) construct in the model indicates the proportion of variance explained by the exogenous (independent) constructs linked to it, providing a measure of the model's explanatory power. Higher R^2 values suggest that the model effectively captures the theoretical relationships among constructs, contributing to the model's nomological validity. Predictive relevance (Q^2) evaluates the model's ability to predict data points that were not used in the model estimation. A Q^2 value greater than 0 indicates that the model has predictive relevance for the endogenous construct (Chin, 1998; Chopra *et al.*, 2024), further supporting the nomological network's validity.

These metrics, alongside the goodness-of-fit indices, offer a comprehensive view of how well the conceptual model aligns with empirical data, bolstering the argument for nomological validity. Ensuring that the model not only fits the observed data well but also explains a significant proportion of variance in the constructs and possesses predictive relevance, researchers can confidently assert the theoretical coherence and empirical robustness of their study's nomological network. This comprehensive approach to establishing nomological validity ensures a rigorous validation of the model investigated in the study.

Timing for establishing nomological validity (When). Nomological validity assessment is strategically conducted after ensuring the constructs are well-defined and distinct through prior validity assessments and before exploring predictive validity. This sequence is deliberate, ensuring that the constructs are not only theoretically sound and empirically validated but also appropriately interrelated within the model before their predictive capabilities at the relationship level are examined. Situating nomological validity at this juncture, the study provides a solid empirical test of the structural model, affirming the constructs' roles and relationships within the conceptual model. This step is pivotal for theoretical contributions, validating the interconnectedness that forms the study's theoretical foundation. Subsequently, the study may proceed to assess predictive validity, shifting the focus from theoretical integration at the model level to the practical predictive utility of the constructs at the relationship level, thereby broadening the study's implications from theoretical insight to practical application.

3.6 Predictive validity

Following the rigorous validation of the model through nomological validity, the research process advances to the examination of predictive validity. This critical phase of validity assessment centers on evaluating *the ability of predictors to forecast outcomes*, underpinning the applicability and utility of the results regarding the established relationships. As an essential component of criterion validity, predictive validity emphasizes the functional relationship (criterion) between constructs (e.g. predictor and outcome), underscoring the constructs' predictive power in real-world or future scenarios.

What predictive validity is and is not. Predictive validity scrutinizes the ability of specific predictors to forecast outcomes. This is not merely about quantifying the strength of these relationships through predictive coefficients but also encompassing a comprehensive evaluation of their empirical robustness and practical relevance.

While *coefficients* are central to understanding the nature of predictive relationships—indicating the expected change in the outcome variable with a unit change in the predictor—they represent just one facet of predictive validity. The essence of predictive validity lies in its capacity to translate them into meaningful predictions that withstand empirical scrutiny and hold practical significance.

Statistical significance in the context of predictive validity acts as a crucial indicator that the observed predictive relationship is unlikely to be due to chance, underscoring the empirical evidence supporting the existence of a meaningful relationship between predictors and outcomes, beyond random fluctuations in the data. This statistical confirmation enhances the validity of the predictive relationship by providing empirical proof that the predictors under investigation genuinely possess predictive power for the specified outcomes. While statistical significance is essential, it is the interpretation of the predictive coefficients, in light of their significance, that truly enriches understanding of the validity of the predictive relationships. The focus remains on establishing that these relationships are not only statistically detectable but also theoretically meaningful and practically significant, thereby reinforcing the overall validity of the predictive assertions made by the research.

Effect size elevates the assessment of predictive validity by quantifying the magnitude of the predictor's impact on the outcome. This metric addresses the practical question of “how substantial is the effect?” This is crucial because a statistically significant relationship with a negligible effect size may hold limited value in practical applications. Conversely, a relationship with a substantial effect size, indicating a considerable impact of the predictor on the outcome, underscores the practical importance of the constructs in explaining real-world phenomena, even if statistical significance is modest.

Integrating these three elements—deepening the interpretation of predictive coefficients, affirming the relationships’ validity through statistical significance and contextualizing their practical impact via effect size—forms the cornerstone of a rigorous assessment of predictive validity. This comprehensive approach ensures that predictive validity assessments are not only grounded in statistical rigor but are also imbued with practical relevance. Focusing on the relationship level, this validation process emphasizes the individual contributions of construct relationships to predictive outcomes, ensuring that the findings offer actionable insights and tangible implications for theory and practice.

While predictive validity shares an overarching focus with nomological validity on the relationships among constructs, it is essential to distinguish between the two. Predictive validity operates at the relationship level. This focused assessment ensures that each predictive link within the model is validated for its theoretical grounding and practical utility. In contrast, nomological validity is concerned with the model level, evaluating the coherence and logical consistency of the entire network of relationships within a model. The distinction lies in the unit of analysis: predictive validity zeroes in on individual relationships to assess their predictive capacity, while nomological validity takes a holistic view of the model to ensure its theoretical fit. Understanding this distinction is crucial for accurately applying and interpreting these forms of validity, thereby preventing the conflation of their roles and reinforcing the clarity and precision of validity assessments within empirical research.

Establishing predictive validity (How). The assessment of predictive validity typically involves statistical analyses such as regression modeling, where the outcome is regressed on the predictor(s). Key indicators in this assessment include:

- (1) *Coefficient* signifying the direction (+/−) and strength of the relationship between the predictor and the outcome. Coefficients are crucial for understanding how changes in the predictor are expected to influence the outcome, wherein a higher coefficient denotes a stronger influence of the predictor on the outcome.
- (2) *Statistical significance* of the relationships, commonly assessed through *p*-values, *t*-values and confidence intervals. A *p*-value of less than 0.05 typically indicates statistical significance, while a more stringent threshold of 0.01 offers stronger evidence of the relationship’s validity. Correspondingly, *t*-values that exceed $|t| > 1.96$ (two-tailed; no direction specified) or $t > 1.645$ (one-tailed; direction specified) for a 95% confidence level and $|t| > 2.576$ (two-tailed; no direction specified) or $t > 2.326$ (one-tailed; direction specified) for a 99% confidence level affirm the strength of the relationship with greater confidence. Similarly, *confidence intervals* provide additional depth by offering a range within which the true value of the coefficient is likely to lie, with narrower intervals indicating greater precision. It is crucial that the confidence interval does not include zero, as this would imply uncertainty about the direction or existence of the relationship, and thus, the exclusion of zero from the interval reinforces the certainty of the predictive relationship.
- (3) *Effect size* measures, such as Cohen’s f^2 for regression and Pearson’s r for correlations, provide a gauge of the practical significance of the relationships, helping to discern the extent of the predictor’s effect on the outcome. Standard benchmarks categorize effect sizes as small ($f^2 \geq 0.02$, $r \approx 0.1$ to 0.3), medium ($f^2 \geq 0.15$, $r \approx 0.3$ to 0.5), or large ($f^2 \geq 0.35$, $r > 0.5$), offering a context to the statistical significance (Cohen, 1988). Specifically, effect size illuminates the real-world significance of the predictive relationship, offering insight into the extent of the predictor’s influence on the outcome.

These comprehensive measures provide a robust framework for assessing predictive validity, ensuring that the predictive relationships under scrutiny are not only statistically substantiated but also of meaningful magnitude and practical significance.

Timing for establishing predictive validity (When). This assessment is ideally undertaken after the model's theoretical and structural integrity has been validated through previous validity assessments, including nomological validity. This ensures that the constructs are not only conceptually and empirically sound but also accurately reflect their proposed theoretical relationships before their predictive utility is tested. Positioning predictive validity at this stage allows for a meaningful transition from theoretical validation to practical application, highlighting the constructs' ability to provide actionable insights and predict outcomes within the specified model. Effectively establishing predictive validity, researchers underscore the real-world relevance of their model, enhancing their research's contribution to both academic knowledge and practical problem-solving.

3.7 Threats to validity

The integrity of research hinges not only on the establishment of various types of validity—including content, face, convergent, discriminant, nomological and predictive validity—but also on the robustness of the study's treatment of internal and external threats to validity. Internal threat to validity, affecting the *legitimacy of the study's insights, interpretations and implications* (i.e. the 3Is), and external threat to validity, concerning the study's *generalizability*, are foundational to the study's overall credibility and utility. Addressing potential threats to validity is crucial for ensuring that the study not only stands up to scrutiny within its immediate context but also resonates and remains relevant across wider contexts (Table 2).

3.7.1 Internal threat to validity. Internal threat to validity can threaten the degree to which a study successfully identifies genuine relationships between constructs while minimizing extraneous bias. A significant internal threat to validity in research involving content, face, convergent, discriminant, nomological and predictive validity is *common method bias*. This form of bias arises when the variations observed in the data stem more from the methodology (e.g. methods, measures) employed in data collection than from the constructs under investigation, potentially skewing the interpretation of the relationships among these constructs. To address this, researchers can employ procedural and statistical treatments.

Procedurally, ensuring anonymity and reducing evaluation concerns (e.g. emphasizing the importance of honest feedback over “correct” answers) can minimize socially desirable responses (i.e. the tendency of respondents to answer questions in a manner that will be viewed favorably by others). Implementing temporal (e.g. collecting data at different points in time), psychological (e.g. framing questions in varied ways to tap into different cognitive processes), or methodological (e.g. using mixed methods that involve primary and secondary or qualitative and quantitative approaches) separation of measurement can also reduce this bias.

Statistically, addressing common method bias can be achieved through various advanced techniques, each providing a unique lens through which to assess the potential impact of methodological artifacts on the study's findings. Harman's (1976) single-factor test can be used to detect the presence of common method variance, where the criterion is that a single factor emerging from an unrotated factor analysis should not explain more than 50% of the variance among the constructs or variables (Kock, 2020). This threshold acts as a safeguard, ensuring that no single methodological factor disproportionately influences the data.

The full collinearity test offers another statistical avenue for assessing common method bias (Kock, 2015). Regressing a dependent variable, artificially constructed from random numbers, against all constructs in the model, this test scrutinizes the data for signs of collinearity that might indicate method bias. A key indicator in this test is the variance

Table 2.
Threats to validity

Threat to validity	Internal	External
Threat	Common method bias, stemming from the methodology employed in data collection than from the constructs	Generalizability
Threatened	The degree to which a study successfully identifies genuine relationships between constructs while minimizing extraneous bias	The degree to which research findings are applicable beyond the specific conditions under which the study was conducted
Treatment	Procedural • Ensure anonymity • Reduce evaluation concerns • Separation (temporal, psychological, methodological) Statistical • Harmon's single factor test, where an unrotated factor analysis should not explain more than 50% of the variance among the variables • Full collinearity test, where variance inflation factors (VIFs) are less than the critical threshold of 3.3 • Measured latent marker variable (MLMV) method, where introducing a construct unrelated to the study's main constructs as a marker variable should not significantly distort the relationships among the primary variables of interest, otherwise revisit measurement instruments to remove potential overlap (e.g. items with similar phrases) or impose statistical controls to partial out the effects of the marker variable	Methodological • Diversify samples (e.g. cluster or stratified sampling) • Engage in replication studies (e.g. industries, populations, time periods) • Triangulation (e.g. multiple data sources, methods, perspectives, studies)

Source(s): Author's own illustration

inflation factors (VIFs); values less than the critical threshold of 3.3 suggest that collinearity—and by extension, common method bias—is not unduly influencing the results.

The measured latent marker variable (MLMV) method is another approach to assess common method bias (Chin *et al.*, 2013). This method involves introducing a construct unrelated to the study's main constructs as a marker variable. The analysis then involves comparing path coefficients across two models: one that incorporates the marker variable and one that does not. The expectation is that the incorporation of this unrelated marker variable should not significantly distort the relationships among the primary constructs or variables of interest. Any significant changes in the path coefficients when the marker variable is included might indicate the presence of common method bias, warranting further investigation and potentially model adjustments. Such adjustments could include revisiting measurement instruments to remove potential overlap (e.g. items with similar phrases), or applying statistical controls to partial out the effects of the marker variable, thereby refining the model to more accurately reflect the true relationships among the constructs.

3.7.2 *External threat to validity.* External threat to validity affects *generalizability*, threatening the credibility and utility of research findings beyond the specific conditions under which the study was conducted. The representativeness of the study's sample and the

uniqueness of its setting pose challenges to validity, potentially restricting the findings' relevance to broader contexts and populations.

To counter these limitations and enhance the generalizability of research outcomes, adopting a multifaceted approach is beneficial. Replication studies, for instance, play a pivotal role by testing the study's hypotheses across varied contexts (e.g. industries, populations, time periods), thereby assessing the validity of findings under different conditions. Diverse sampling methods further contribute to this endeavor by encompassing a broad spectrum of participants, ensuring the sample's representativeness. Additionally, the technique of triangulation, which involves integrating multiple methodologies, perspectives and sources, enriches the research by providing varied perspectives on the phenomenon under investigation. This not only enhances the robustness of the findings but also their applicability to different contexts, thereby bolstering the study's validity.

In parallel, the pursuit of multiple evidences stands as a cornerstone in reinforcing the validity of research findings. Aggregating and comparing results from various studies, researchers can discern patterns and variances in findings, offering a grounded basis for generalizations. Systematic reviews, including bibliometric analyses and meta-analyses, serve as instrumental tools or references for comparison in this process. They allow researchers to collate and scrutinize a wide array of findings, identifying overarching trends and discrepancies. Such comparative analysis, especially when juxtaposed against observations from systematic reviews, provides a finer-grained understanding of the phenomena, ensuring the findings' applicability and relevance across diverse contexts.

4. Conclusion

This article has approached and reconciled the attributes and manifestations of validity, elucidating a typology that underscores its multifaceted nature. Commencing with a foundational understanding of validity, this article ventured through the peculiarities of content, face, convergent, discriminant, nomological and predictive validity, each serving as a critical pillar in the construction of robust research with integrity and impact.

Content and face validity emerged as the precursors in the typology, laying the groundwork for a rigorous inquiry by ensuring the conceptual and operational alignment of measurement instruments with the constructs they aim to measure and the concepts they intend to represent. Convergent and discriminant validity, assessed within the confines of the measurement model, further refined understanding by affirming the coherence and distinctiveness of constructs. Nomological and predictive validity, which, situated within the structural model, attest to the theoretical and practical value of research in forecasting outcomes in a way that is theoretical coherent and explainable.

Moreover, the article illuminated the distinction between construct and criterion validity, underscoring the role of the former in concept-focused assessments and the latter in relationship-focused inquiries. This delineation not only enriches understanding of validity's dimensions but also enhances the scope of validity assessments in empirical research.

Yet, the passage of establishing the typology of validity is not devoid of challenges. The article highlighted the perils of common method bias and generalizability concerns, presenting them as internal and external threats to validity, respectively. Through the examination of these threats and the proposition of methodological and statistical countermeasures, the article underscored the imperative of safeguarding research against potential validity pitfalls to uphold its integrity.

As digital and online research methodologies become increasingly prevalent, the typology of validity presented in this article remains robust and relevant across both technological and non-technological contexts. The principles of content, face, convergent, discriminant, nomological and predictive validity are equally applicable regardless of the research setting

(e.g. brick-and-mortar stores, online stores, physical offices, virtual offices). For example, content validity ensures that survey questions in an online survey cover all relevant aspects of the construct being measured, just as they do in a traditional paper-and-pen survey. Face validity ensures that participants perceive the survey as appropriate and relevant in both technology-mediated (e.g. online surveys) and non-mediated (e.g. paper-and-pen surveys) data collection conditions. Convergent and discriminant validity remain essential in confirming that constructs are measured coherently and distinctly, regardless of the medium. Nomological and predictive validity are crucial in establishing that the model and relationships among constructs hold true in digital contexts, like e-commerce, as well as in traditional settings, such as shopping malls. Although challenges such as data integrity issues and sampling biases in online data collection require specific methodological adjustments (e.g. implementing data encryption to protect data integrity and using stratified sampling techniques to mitigate sampling bias), these adjustments fall under the broader scope of research methodology. Adhering to the rigorous validity standards outlined in this typology ensures that findings are robust and relevant, while methodological adjustments specifically address contextual challenges posed by technology. This approach allows researchers to produce findings that are both statistically and substantively significant, regardless of the context or method (technological or non-technological) in which the research is conducted.

In synthesizing these insights, this article contributes a comprehensive framework that guides researchers in the establishment of validity across different research stages. This contribution not only reconciles the discourse on validity in academic circles but also serves as a compass for practitioners, enabling them to navigate the complexities of research validation with clarity and confidence. The typology of validity presented herein is not only an academic concept or framework but also a pragmatic guide that underpins the pursuit of rigor and impact in research, and thus, researchers embracing this typology will be equipped to transcend the conventional boundaries of validity assessment, ultimately fostering a culture of integrity and innovation in the quest for knowledge and impact.

Notes

1. For a comprehensive understanding of the philosophy of science and research paradigms, see [Lim \(2023\)](#). An exemplar of stakeholders reflecting societal needs can be found in [Lim and Bowman \(2023\)](#).
2. *Conceptualization* refers to the process of defining concepts. This involves specifying what we mean by a particular concept in a theoretical context. This step is crucial because many concepts we use in research can have multiple meanings or interpretations. Through conceptualization, researchers develop a clear, precise definition that sets the boundaries and dimensions of the concept as it will be used in their study. This process transforms abstract ideas into something more specific (concept) by delineating their essential attributes and specifying the framework within which they are situated. *Operationalization*, on the other hand, is the process of translating these concepts into something measurable (constructs). This involves identifying specific, observable and measurable elements that can be assessed to represent the constructs in empirical research. Operationalization determines how a construct will be measured in the real world, which can involve choosing appropriate instruments (questionnaires) and measures (items, scales). Both conceptualization and operationalization are critical for ensuring the validity of research. They provide a systematic approach to moving from abstract ideas to concrete empirical evidence, enabling researchers to collect data that reflects the *concepts* (conceptualization—conceptual definition) and *constructs* (operationalization—operational measurement) of interest and supports meaningful analysis and interpretation of the phenomena under study.
3. The *theoretical model* (or framework) delineates the underlying theory(ies) guiding the study, articulating the theoretical assumptions and principles that inform the research. The *conceptual model* outlines the relationships among key concepts derived from or guided by the theoretical

model, serving as a blueprint for how these concepts are expected to interact. The *measurement model* details the operationalization of these concepts into measurable constructs, specifying how each construct is represented by its indicators (items) and ensuring their reliability and validity. The *structural model* presents the empirical testing of the relationships posited and their fit in the conceptual model, using statistical methods to examine the paths between constructs and validate the proposed theoretical relationships and overall fit.

4. Imagine a construct “health consciousness.” If indicators like “eating healthy food,” “frequency of exercising,” and “regular medical check-ups” are initially treated as reflective but show diverse associations with the construct, it might indicate these actions contribute to forming “health consciousness” rather than being manifestations (reflections) of it. In such a case, re-operationalizing “health consciousness” as a formative construct with these indicators would be more appropriate, as each action contributes uniquely to the overall concept of being health-conscious.
5. Single-item measurement is common and receives support from the literature (Diamantopoulos *et al.*, 2012; Sarstedt *et al.*, 2016).
6. Can be cross-checked against a correlation matrix (Hair *et al.*, 2021).
7. While discerning the maximum thresholds for HTMT values offers an understanding of construct relationships, a similar conceptual treatment for correlation thresholds proposes a maximum of 0.70 for conceptually distinct constructs, whereas a cautious consideration of up to 0.80 might be proposed for conceptually similar constructs, acknowledging the theoretical closeness of such constructs. However, this approach, informed by the range of 0.70–0.80 often used in the literature (Berry and Feldman, 1985; Dormann *et al.*, 2013), should be applied judiciously, with a clear theoretical justification and being mindful of the primary objective of ensuring discriminant validity—that constructs remain sufficiently distinct to uphold the integrity of their respective concepts. Such differentiation in thresholds emphasizes the importance of a contextually informed approach to validating construct relationships while maintaining a conservative stance to preserve the distinctiveness of each construct.
8. While nomological networks of concepts and relationships are typically guided by theories (as usually seen in empirical research), these networks can also be established through organizing frameworks like Paul and Benito’s (2018) antecedents, decisions and outcomes (ADO) framework or Luo *et al.*’s (2024) antecedents, mediators, moderators, outcomes and control variables (AMMO-CV) framework (prominently in review research; Kraus *et al.*, 2022; Lim *et al.*, 2022; Paul *et al.*, 2021).
9. The work of Khatri *et al.* (2024) exemplifies that within a seemingly simple nomological network of a single relationship between two constructs, a profound depth of theoretical and empirical richness can be embedded. Despite focusing on the singular relationship between student well-being and positive word-of-mouth, the authors delve into the multifaceted nature of student well-being, treating it as a higher-order construct composed of various dimensions including academic, financial, physical, psychological and relational well-being. This approach underscores that the complexity and depth of a nomological network are not solely contingent on the number of constructs or relationships examined, but also on the conceptual breadth and empirical substantiation of the constructs involved. This study demonstrates how a single, well-defined and theoretically grounded relationship can provide significant insights into complex phenomena, thereby affirming the model’s nomological validity.

References

- Becker, J.M., Ringle, C.M., Sarstedt, M. and Völckner, F. (2015), “How collinearity affects mixture regression results”, *Marketing Letters*, Vol. 26 No. 4, pp. 643-659, doi: [10.1007/s11002-014-9299-9](https://doi.org/10.1007/s11002-014-9299-9).
- Becker, J.M., Cheah, J.-H., Gholamzade, R., Ringle, C.M. and Sarstedt, M. (2023), “PLS-SEM’s most wanted guidance”, *International Journal of Contemporary Hospitality Management*, Vol. 35 No. 1, pp. 321-346, doi: [10.1108/ijchm-04-2022-0474](https://doi.org/10.1108/ijchm-04-2022-0474).
- Berry, W.D. and Feldman, S. (1985), *Multiple Regression in Practice: Quantitative Applications in the Social Sciences*, Sage Publications, Thousand Oaks, CA.

- Bollen, K.A. (1989), *Structural Equations with Latent Variables*, Wiley, New York, NY.
- Cheah, J.-H., Sarstedt, M., Ringle, C.M., Ramayah, T. and Ting, H. (2018), "Convergent validity assessment of formatively measured constructs in PLS-SEM: on using single-item versus multi-item measures in redundancy analyses", *International Journal of Contemporary Hospitality Management*, Vol. 30 No. 11, pp. 3192-3210, doi: [10.1108/IJCHM-10-2017-0649](https://doi.org/10.1108/IJCHM-10-2017-0649).
- Chin, W.W. (1998), "The partial least squares approach to structural equation modelling", *Modern Methods for Business Research*, Vol. 295 No. 2, pp. 295-336.
- Chin, W.W., Thatcher, J.B., Wright, R.T. and Steel, D. (2013), "Controlling for common method variance in PLS analysis: the measured latent marker variable approach", in *New Perspectives in Partial Least Squares and Related Methods*, Springer, New York, pp. 231-239.
- Chopra, I.P., Lim, W.M. and Jain, T. (2024), "Electronic word of mouth on social networking sites: what inspires travelers to engage in opinion seeking, opinion passing, and opinion giving?", *Tourism Recreation Research*, ahead-of-print doi: [10.1080/02508281.2022.2088007](https://doi.org/10.1080/02508281.2022.2088007).
- Cohen, J.E. (1988), *Statistical Power Analysis for the Behavioral Sciences*, Lawrence Erlbaum Associates, Hillsdale, NJ.
- Cronbach, L.J. and Meehl, P.E. (1955), "Construct validity in psychological tests", *Psychological Bulletin*, Vol. 52 No. 4, pp. 281-302, doi: [10.1037/h0040957](https://doi.org/10.1037/h0040957).
- Diamantopoulos, A., Sarstedt, M., Fuchs, C., Wilczynski, P. and Kaiser, S. (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*, Vol. 40 No. 3, pp. 434-449, doi: [10.1007/s11747-011-0300-3](https://doi.org/10.1007/s11747-011-0300-3).
- Dormann, C.F., Elith, J., Bacher, S., Buchmann, C., Carl, G., Carré, G., Marquéz, J.R.G., Gruber, B., Lafourcade, B., Leitão, P.J., Münkemüller, T., McClean, C., Osborne, P.E., Reineking, B., Schröder, B., Skidmore, A.K., Zurell, D. and Lautenbach, S. (2013), "Collinearity: a review of methods to deal with it and a simulation study evaluating their performance", *Ecography*, Vol. 36 No. 1, pp. 27-46, doi: [10.1111/j.1600-0587.2012.07348.x](https://doi.org/10.1111/j.1600-0587.2012.07348.x).
- Fornell, C. and Larcker, D.F. (1981), "Evaluating structural equation models with unobservable variables and measurement error", *Journal of Marketing Research*, Vol. 18 No. 1, pp. 39-50, doi: [10.2307/3151312](https://doi.org/10.2307/3151312).
- Green, J.P., Tonidandel, S. and Cortina, J.M. (2016), "Getting through the gate: statistical and methodological issues raised in the reviewing process", *Organizational Research Methods*, Vol. 19 No. 3, pp. 402-432, doi: [10.1177/1094428116631417](https://doi.org/10.1177/1094428116631417).
- Grey, A., Bolland, M., Gamble, G. and Avenell, A. (2019), "Quality of reports of investigations of research integrity by academic institutions", *Research Integrity and Peer Review*, Vol. 4 No. 1, 3, doi: [10.1186/s41073-019-0062-x](https://doi.org/10.1186/s41073-019-0062-x).
- Haertel, E. and Herman, J. (2005), *A Historical Perspective on Validity Arguments for Accountability Testing*, National Center for Research on Evaluation, Standards, and Student Testing (CRESST), Center for the Study of Evaluation (CSE), Graduate School of Education & Information Studies, University of California, Los Angeles, CA.
- Hair, J.F., Ronald, L.T., Anderson, R.E. and Black, W. (1998), *Multivariate Data Analysis*, Prentice-Hall International, London, UK.
- Hair, J.F., Black, W., Babin, B. and Anderson, R. (2014), *Multivariate Data Analysis*, 7th ed., Prentice Hall, New Jersey, NJ.
- Hair, J.F., Hult, G.T.M., Ringle, C.M., Sarstedt, M., Danks, N.P. and Ray, S. (2021), *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R: A Workbook*, Springer Nature, New York.
- Hair, J.F., Hult, G.T.M., Ringle, C.M. and Sarstedt, M. (2022), *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*, 3rd ed., Sage, Thousand Oaks, CA.
- Hardy, B. and Ford, L.R. (2014), "It's not me, it's you: miscomprehension in surveys", *Organizational Research Methods*, Vol. 17 No. 2, pp. 138-162, doi: [10.1177/1094428113520185](https://doi.org/10.1177/1094428113520185).

- Harman, H.H. (1976), *Modern Factor Analysis*, University of Chicago Press, Chicago, IL.
- Haws, K.L., Sample, K.L. and Hulland, J. (2023), "Scale use and abuse: towards best practices in the deployment of scales", *Journal of Consumer Psychology*, Vol. 33 No. 1, pp. 226-243, doi: [10.1002/jcpy.1320](https://doi.org/10.1002/jcpy.1320).
- Henseler, J., Ringle, C.M. and Sarstedt, M. (2015), "A new criterion for assessing discriminant validity in variance-based structural equation modeling", *Journal of the Academy of Marketing Science*, Vol. 43 No. 1, pp. 115-135, doi: [10.1007/s11747-014-0403-8](https://doi.org/10.1007/s11747-014-0403-8).
- Henseler, J., Hubona, G. and Ray, P.A. (2016), "Using PLS path modeling in new technology research: updated guidelines", *Industrial Management and Data Systems*, Vol. 116 No. 1, pp. 2-20, doi: [10.1108/imds-09-2015-0382](https://doi.org/10.1108/imds-09-2015-0382).
- Houston, M.B. (2004), "Assessing the validity of secondary data proxies for marketing constructs", *Journal of Business Research*, Vol. 57 No. 2, pp. 154-161, doi: [10.1016/s0148-2963\(01\)00299-5](https://doi.org/10.1016/s0148-2963(01)00299-5).
- Hulland, J., Baumgartner, H. and Smith, K.M. (2018), "Marketing survey research best practices: evidence and recommendations from a review of JAMS articles", *Journal of the Academy of Marketing Science*, Vol. 46 No. 1, pp. 92-108, doi: [10.1007/s11747-017-0532-y](https://doi.org/10.1007/s11747-017-0532-y).
- Khatrri, P., Duggal, H.K., Lim, W.M., Thomas, A. and Shiva, A. (2024), "Student well-being in higher education: scale development and validation with implications for management education", *International Journal of Management in Education*, Vol. 22 No. 1, 100933, doi: [10.1016/j.ijme.2024.100933](https://doi.org/10.1016/j.ijme.2024.100933).
- Kline, R.B. (1998), *Principles and Practice of Structural Equation Modeling*, Guilford Press, New York, NY.
- Kock, N. (2015), "Common method bias in PLS-SEM: a full collinearity assessment approach", *International Journal of E-Collaboration*, Vol. 11 No. 4, pp. 1-10, doi: [10.4018/ijec.2015100101](https://doi.org/10.4018/ijec.2015100101).
- Kock, N. (2020), "Harman's single factor test in PLS-SEM: checking for common method bias", *Data Analysis Perspectives Journal*, Vol. 2 No. 2, pp. 1-6.
- Kock, F., Berbekova, A., Assaf, A.G. and Josiassen, A. (2024), "Developing a scale is not enough: on the importance of nomological validity", *International Journal of Contemporary Hospitality Management*, ahead-of-print doi: [10.1108/IJCHM-07-2023-1078](https://doi.org/10.1108/IJCHM-07-2023-1078).
- Kraus, S., Breier, M., Lim, W.M., Dabić, M., Kumar, S., Kanbach, D., Mukherjee, D., Corvello, V., Piñeiro-Chousa, J., Liguori, E., Palacios-Marqués, D., Schiavone, F., Ferraris, A., Fernandes, C. and Ferreira, J.J. (2022), "Literature reviews as independent studies: guidelines for academic practice", *Review of Managerial Science*, Vol. 16 No. 8, pp. 2577-2595, doi: [10.1007/s11846-022-00588-8](https://doi.org/10.1007/s11846-022-00588-8).
- Lim, W.M. (2015), "Antecedents and consequences of e-shopping: an integrated model", *Internet Research*, Vol. 25 No. 2, pp. 184-217, doi: [10.1108/intr-11-2013-0247](https://doi.org/10.1108/intr-11-2013-0247).
- Lim, W.M. (2022), "The art of writing for premier journals", *Global Business and Organizational Excellence*, Vol. 41 No. 6, pp. 5-10, doi: [10.1002/joe.22178](https://doi.org/10.1002/joe.22178).
- Lim, W.M. (2023), "Philosophy of science and research paradigm for business research in the transformative age of automation, digitalization, hyperconnectivity, obligations, globalization and sustainability", *Journal of Trade Science*, Vol. 11 Nos 2/3, pp. 3-30, doi: [10.1108/jts-07-2023-0015](https://doi.org/10.1108/jts-07-2023-0015).
- Lim, W.M. and Bowman, C. (2023), "How to establish practical contributions and convey practical implications? Guidelines on locating practice gaps and making recommendations for practice", *Activities, Adaptation and Aging*, Vol. 47 No. 3, pp. 263-282, doi: [10.1080/01924788.2023.2232220](https://doi.org/10.1080/01924788.2023.2232220).
- Lim, W.M. and Koay, K.Y. (2024), "So you want to publish in a premier journal? An illustrative guide on how to develop and write a quantitative research paper for premier journals", *Global Business and Organizational Excellence*, Vol. 43 No. 3, pp. 5-19, doi: [10.1002/joe.22252](https://doi.org/10.1002/joe.22252).
- Lim, W.M., Kumar, S. and Ali, F. (2022), "Advancing knowledge through literature reviews: 'What', 'why', and 'how to contribute'", *Service Industries Journal*, Vol. 42 Nos 7-8, pp. 481-513, doi: [10.1080/02642069.2022.2047941](https://doi.org/10.1080/02642069.2022.2047941).

- Luo, X., Lim, W.M., Cheah, J.-H., Lim, X.J. and Dwivedi, Y.K. (2024), "Live streaming commerce: a review and research agenda", *Journal of Computer Information Systems*, ahead-of-print doi: [10.1080/08874417.2023.2290574](https://doi.org/10.1080/08874417.2023.2290574).
- Memon, M.A., Ting, H., Ramayah, T., Chuah, F. and Cheah, J.-H. (2017), "A review of the methodological misconceptions and guidelines related to the application of structural equation modelling: a Malaysian scenario", *Journal of Applied Structural Equation Modeling*, Vol. 1 No. 1, pp. i-xiii.
- Memon, M.A., Thursamy, R., Cheah, J.-H., Ting, H., Chuah, F. and Cham, T.H. (2023), "Addressing common method bias, operationalization, sampling, and data collection issues in quantitative research: review and recommendations", *Journal of Applied Structural Equation Modeling*, Vol. 7 No. 2, pp. i-xiv, doi: [10.47263/jasem.7\(2\)01](https://doi.org/10.47263/jasem.7(2)01).
- Messick, S. (1989), "Validity", in Linn, R.L. (Ed.), *Educational Measurement*, 3rd ed., American Council on Education and Macmillan, New York, NY, pp. 13-103.
- Messick, S. (1995), "Standards of validity and the validity of standards in performance assessment", *Educational Measurement: Issues and Practice*, Vol. 14 No. 4, pp. 5-8, doi: [10.1111/j.1745-3992.1995.tb00881.x](https://doi.org/10.1111/j.1745-3992.1995.tb00881.x).
- Moher, D., Bouter, L., Kleinert, S., Glasziou, P., Sham, M.H., Barbour, V., Coriat, A.M., Foeger, N. and Dirnagl, U. (2020), "The Hong Kong Principles for assessing researchers: fostering research integrity", *PLoS Biology*, Vol. 18 No. 7, e3000737, doi: [10.1371/journal.pbio.3000737](https://doi.org/10.1371/journal.pbio.3000737).
- Moore, G.C. and Benbasat, I. (1991), "Development of an instrument to measure the perceptions of adopting an information technology innovation", *Information Systems Research*, Vol. 2 No. 3, pp. 192-222, doi: [10.1287/isre.2.3.192](https://doi.org/10.1287/isre.2.3.192).
- Paul, J. and Benito, G.R. (2018), "A review of research on outward foreign direct investment from emerging countries, including China: what do we know, how do we know and where should we be heading?", *Asia Pacific Business Review*, Vol. 24 No. 1, pp. 90-115, doi: [10.1080/13602381.2017.1357316](https://doi.org/10.1080/13602381.2017.1357316).
- Paul, J., Lim, W.M., O'Cass, A., Hao, A.W. and Bresciani, S. (2021), "Scientific procedures and rationales for systematic literature reviews (SPAR-4-SLR)", *International Journal of Consumer Studies*, Vol. 45 No. 4, pp. O1-O16, doi: [10.1111/ijcs.12695](https://doi.org/10.1111/ijcs.12695).
- Podsakoff, P.M., Podsakoff, N.P., Williams, L.J., Huang, C. and Yang, J. (2024), "Common method bias: it's bad, it's complex, it's widespread, and it's not easy to fix", *Annual Review of Organizational Psychology and Organizational Behavior*, Vol. 11 No. 1, pp. 17-61, doi: [10.1146/annurev-orgpsych-110721-040030](https://doi.org/10.1146/annurev-orgpsych-110721-040030).
- Preckel, F. and Brunner, M. (2017), "Nomological nets", in Zeigler-Hill, V. and Shackelford, T. (Eds), *Encyclopedia of Personality and Individual Differences*, Springer, Cham, Switzerland.
- Rao, P., Kumar, S., Lim, W.M. and Rao, A.A. (2024), "The ecosystem of research tools for scholarly communication", *Library Hi Tech*, Ahead-of-print doi: [10.1108/LHT-05-2022-0259](https://doi.org/10.1108/LHT-05-2022-0259).
- Robinson, M.A. (2018), "Using multi-item psychometric scales for research and practice in human resource management", *Human Resource Management*, Vol. 57 No. 3, pp. 739-750, doi: [10.1002/hrm.21852](https://doi.org/10.1002/hrm.21852).
- Sarstedt, M., Diamantopoulos, A., Salzberger, T. and Baumgartner, P. (2016), "Selecting single items to measure doubly concrete constructs: a cautionary tale", *Journal of Business Research*, Vol. 69 No. 8, pp. 3159-3167, doi: [10.1016/j.jbusres.2015.12.004](https://doi.org/10.1016/j.jbusres.2015.12.004).
- Terman, L.M. (1919), *The Intelligence of School Children*, Houghton Mifflin, Boston, MA.
- Thorndike, E.L. (1910), "The contribution of psychology to education", *Journal of Educational Psychology*, Vol. 1, pp. 5-12, doi: [10.1037/h0070113](https://doi.org/10.1037/h0070113).
- Thorndike, E.L. (1918), "The nature, purposes, and general methods of measurements of educational products", in Whipple, G.M. (Ed.), *The Measurement of Educational Products (17th Yearbook of the National Society for the Study of Education, Part II)*, Public School Publishing Company, Bloomington, IL, pp. 16-24.

Vaithilingam, S., Ong, C.S., Moisescu, O.I. and Nair, M.S. (2024), "Robustness checks in PLS-SEM: a review of recent practices and recommendations for future applications in business research", *Journal of Business Research*, Vol. 173, 114465, doi: [10.1016/j.jbusres.2023.114465](https://doi.org/10.1016/j.jbusres.2023.114465).

Journal of Trade
Science

179

About the author

Weng Marc Lim is a Distinguished Professor and the Dean of Sunway Business School at Sunway University in Malaysia and an Adjunct Professor at the Swinburne University of Technology's home campus in Melbourne, Australia and international branch campus in Sarawak, Malaysia. He has authored more than 100 manuscripts in journals ranked "A*" and "A" such as *Australasian Marketing Journal*, *European Journal of Marketing*, *Industrial Marketing Management*, *Journal of Advertising*, *Journal of Advertising Research*, *Journal of Business Research*, *Journal of Business and Industrial Marketing*, *Journal of Consumer Behaviour*, *Journal of Consumer Marketing*, *International Journal of Consumer Studies*, *Journal of Brand Management*, *Journal of Product and Brand Management*, *Journal of Retailing and Consumer Services*, *Journal of International Marketing*, *Journal of Strategic Marketing*, *Marketing Theory*, *Marketing Intelligence and Planning* and *Psychology and Marketing*, among others. He has also presented his work and led high-level policy discussions at the United Nations Educational, Scientific and Cultural Organization and the World Economic Forum. Contact: @limwengmarc on Instagram and Twitter (X), LinkedIn or his personal homepage at <https://www.wengmarc.com>. Weng Marc Lim can be contacted at: lim@wengmarc.com, marcl@sunway.edu.my, marclim@swin.edu.au, wlim@swinburne.edu.my

For instructions on how to order reprints of this article, please visit our website:

www.emeraldgroupublishing.com/licensing/reprints.htm

Or contact us for further details: permissions@emeraldinsight.com