

https://econtent.hogrefe.com/doi/pdf/10.1027/2698-1866/a000002 - Sunday, May 05, 2024 1:31:04 AM - IP Address:3.17.150.89

Editorial

Psychological Test Adaptation and Development – How Papers Are Structured and Why

Matthias Ziegler

Department of Psychology, Humboldt-Universität zu Berlin, Germany



Abstract. This article explains how papers should be structured to guide the preparation of papers to be submitted to *Psychological Test Adaptation and Development*. Each submission should adhere as strictly as possible to the following structure. If, for any reason, certain aspects cannot be provided, this should be explained and considered in the limitations and recommendations. The outline in Table 1 is followed by a detailed explanation for each section.

Keywords: PTAD, registered report, ABC of test construction, open science

The structure of a paper in *Psychological Test Adaptation and Development* (PTAD) generally follows the typical structure for scientific papers consisting of Introduction/Theory, Methods, Results, and Discussion. We encourage authors to also consider the opportunity to use online supplementary material to support any of these sections.

While the structure, in general, is pretty much the same as in most other psychological journals, the specific contents of each section are more closely predefined and follow the three lead questions from the “ABC of test construction” (Ziegler, 2014b):

- A. What is/are the construct(s) being measured?
- B. What are the intended uses?
- C. What is the intended target population?

Table 1. Content required in papers given by section

Section content required	Specifics for registered reports
Theory/Introduction	
(A) <i>What is the construct being measured?</i>	
• Define each construct measured • Define possible hierarchy • Elaborate on score intercorrelations • Elaborate on nomological net - Relations with manifest variables (items) - Relations with other constructs • Derive hypotheses regarding - Structural validity - Convergent and discriminant validity • Explain how items were constructed and how they reflect the theory defined above • Consider consequences of adaptations/translations	
(B) <i>What is the intended use?</i>	
• Define intended use(s)/rule out what it is not intended for - Elaborate on necessary data selection contexts	

(Continued on next page)

Table 1. (Continued)

Section content required	Specifics for registered reports
• Derive hypotheses regarding test criterion correlations - Pay attention to possible construct overlap, so not only bivariate correlations, use regression, for example, for facet scores • Delineate requirements for item difficulties • Delineate requirements for reliability estimates - Prognosis vs. status - Survey vs. individual assessment - Measurement precision (C) <i>What is the intended target population?</i> • Define target population(s) • Explain how this is adhered to during the studies • Delineate requirements for item difficulties and content Methods • Provide specifics about a possible translation or the adaptation/development undertaken • Data collection • Report all measures used in the entire study • Provide sample information - Descriptive statistics - Size - Composition (e.g., age, gender) • Justify sample size • Statistical analyses - Define alpha level and if a test is one- or two-tailed - Explain which methods match which hypotheses/assumptions and which result is indicative of supporting evidence - Structural validity • Clearly state whether exploratory (no evidence) or confirmatory • Exploratory factor analysis (EFA) vs. confirmatory factor analysis (CFA) • Cluster vs. factor mixture modeling • Define cutoffs • Define what to do in case of misfit - Define rules for item selection (e.g., based on loadings, difficulties, variance, content) - Report software (version) and packages - Provide code that allows a reproduction of analyses Results • Provide all information necessary to evaluate evidence with regard to the assumptions/hypotheses in ABC • <i>Report all exploratory analyses that were additionally conducted</i> Discussion • Evaluate evidence provided with regard to ABC • Elaborate on limitations • Make clear recommendations how the score(s) can be used	• Provide rules for stopping data collection • Not required upon initial submission

Introduction

The Introduction/Theory section should provide clear-cut answers to all three questions. The logical consequences of these answers will then be the foundation for the ensuing research program.

What Is/Are the Construct(s) Being Measured?

The first part of the theory section should be dedicated to this question. The aim is to clearly define each construct measured in the measurement tool featured. If only one unidimensional construct is assessed, one definition should be provided. However, if there are several test scores or hierarchical structures, this needs to be acknowledged as well. Moreover, in case several constructs are being assessed, the relations assumed between them need to be defined. Finally, a nomological net, meaning assumed relations with other constructs or observables, needs to be laid out whenever possible. If it is not possible to present a substantive nomological net, then this needs to be acknowledged as well.

The purpose of this exercise is at least twofold. From the perspective of the test users, it will help them interpret the test results with respect to the intended use (see below in “What Are the Intended Uses?”). At the same time, this follows the theoretical guidance by Cronbach and Meehl (1955) regarding construct validity. They stated that each latent variable (i.e., measured construct in most cases) needs to be defined with regard to how it manifests in observable variables and how it relates to other latent variables. In other words, the introduction should allow each reader to understand the nature of the construct(s) measured and the position within a nomological net. Based on this, direct consequences for the kind of validity evidence needed will emerge. The definition provides information about the internal structure of the measure and thus about the assumed structural validity (Loevinger, 1957). Moreover, the nomological net described directly informs possible hypotheses regarding the evidence needed to support convergent and discriminant validity (Campbell & Fiske, 1959; Wehner, Roemer, & Ziegler, 2018). Finally, as this journal specifically calls for submissions featuring test translations or adaptations, it is necessary to elaborate on potential consequences of this procedure and how these are planned to be tested for.

Structural Validity Evidence

Depending on the answer to question A, it is possible that the scores from the featured test reflect a specific theoretical structure. In that case, it will be relevant to test this

structure using confirmatory approaches such as structural equation modeling or item response theory. In other cases, a clear structure might not be known, which would justify the use of exploratory approaches. The same might apply when tests are translated or culturally adapted. However, it should be noted that exploratory approaches are only hypotheses generating, not hypotheses testing. Thus, without an additional confirmatory step, using additional data, evidence for structural validity will be limited, which must be acknowledged in the paper.

In cases where only one unidimensional score is expected, steps to provide evidence for this assumption are also necessary. We refer authors to the editorial by Ziegler and Hagemann (2015), who also called for confirmatory approaches. In the case of a misfit, the same paper called for modeling of the misfit in order to gauge its impact on the test score and the relations that are actually interpreted (Heene, Hilbert, Draxler, Ziegler, & Bühner, 2011). It needs to be stressed here that a potential misfit also impacts the choice of reliability estimate as the popular Cronbach's α is only suitable for unidimensional items (Cronbach, 1951). Thus, in case of misfit, other estimates such as omega are more suitable (Brunner & Süß, 2005; McNeish, 2018; Revelle & Zinbarg, 2009; Zinbarg, Revelle, Yovel, & Li, 2005).

Convergent Validity Evidence

Evidence for convergent validity, according to Campbell and Fiske (1959), can be obtained by showing that two test scores from different measures which are meant to capture the same construct are correlated more strongly than with scores from tests capturing other constructs. The more modern interpretation often is that the two scores from tests measuring the same construct should be strongly correlated. An even vaguer formulation, which has also found its way into the APA Standards, is to assume a strong correlation between scores from tests which measure the same *or similar* constructs. Schweizer (2012) pointed out that one crucial problem here is the often poor definition of psychological constructs (Michell, 1997, 2001). However, if a clear-cut answer to question A is provided, such a definition should not be found wanting. What remains is the question whether a correlation between scores reflecting similar or identical constructs should both be considered as evidence for convergent validity. Let us consider the example of two intelligence tests. Test A captures verbal reasoning ability, and Test B captures figural reasoning ability. Based on existing research, it can safely be assumed that the correlation between the verbal and the figural test scores will be around 0.5 (Schipolowski, Wilhelm, & Schroeders, 2014). Most of us would consider this to be sufficient evidence for convergent validity. However, when gauging the existing

evidence for the validity of a test score interpretation, it is also vital to consider the answer process and thus the actual psychological processes that occur when a respondent solves an item (AERA, APA, & NMCE, 2014; Bejar, 1983). Now, whereas the processes (for example, inductive and deductive reasoning) might be the same, the content (verbal vs. figural) clearly differs (Schneider & McGrew, 2018). And, importantly, the content is what is driving the difference between different reasoning components (Wilhelm, 2005). Further, when we construct Test A, aiming to measure verbal reasoning, we do not want to measure figural reasoning. Thus, the 0.5 correlation would not be evidence for convergent validity but should rather be interpreted as discriminant validity evidence (incidentally, a correlation with a test measuring the same construct is now needed and should be higher than .5). Therefore, when choosing the kind of evidence to obtain, the already defined nomological net should serve as a guiding framework. In PTAD, we will prefer to see convergent evidence in the stricter sense of Campbell and Fiske. If a more lenient approach, that is, correlation between scores from tests capturing similar constructs, is taken, a theoretical explanation will be required illustrating why the overlap should be considered due to psychological processes that are in the core of the target construct and not an overlapping construct.

Summing up, based on the answer to question A and the nomological net described, clearly stated hypotheses regarding evidence for convergent validity must be included. Of course, not every paper can include such evidence. However, the lack of it should then be noted as a limitation.

Discriminant Validity Evidence

Whereas convergent validity shows whether a score reflects the intended construct as much as other scores claiming to reflect the same construct, discriminant validity is necessary to show that no other, possibly overlapping, construct is being measured. The example above, using two reasoning tests, shows how strongly the two concepts are intertwined. Excellent examples for the importance of discriminant validity can be found in Mussel (2010) or Credé, Tynan, and Harms (2017) for the constructs of epistemic curiosity and grit, respectively. In both cases, the authors show that once theoretically overlapping constructs or constructs which might be close to each other in the nomological net are considered to be discriminant, the presumed core of the score under scrutiny dissolves into the allegedly discriminant constructs. In other words, when choosing measures to provide discriminant validity evidence, it is important to select measures reflecting close or overlapping constructs. Otherwise, the validity evidence will be very weak (Ziegler, Booth, & Bensch, 2013), and

problems of jingle-jangle fallacies (Kelley, 1927) might occur. Again, clear hypotheses need to be formulated.

Item Development

When the construct in question is defined and positioned in a nomological net, item development should be aligned to this information. The paper should report how this was achieved. Ideally, it should be possible to explain how differences in the construct evoke differences during the answering process (Borsboom, Mellenbergh, & Van Heerden, 2003). Moreover, each item should be placeable within the nomological net just defined. In the end, these thoughts ensure content validity.

Consequences of Adaptation

Tests are often translated or adapted to the needs of other populations (e.g., other cultures, age groups). During such processes, several things can happen which might potentially infringe the original purpose of the contained score(s). This is often discussed within the framework of measurement invariance (Borsboom, 2006; Chen, 2008; Fried et al., 2016; Sass, 2011). In fact, without providing evidence regarding measurement invariance with the original version, comparisons of scores are at least highly problematic (Ziegler & Bensch, 2013). Thus, each submission should consider in how far the adaptation or development might have distorted the measure from the original test. Ideally, invariance tests or similar means providing evidence that the adaptation has not resulted in a gross distortion of the relation between definition of the construct and score interpretation should be provided. However, PTAD is not as strict here and does allow submissions without such evidence, for example, when the original data are not available. Still, this must be highlighted as a limitation. Basically, in some cases, this could mean that the featured measure might only be valuable for research purposes, especially if other validity evidence is also limited.

What Are the Intended Uses?

Each scale development, translation, or adaptation should clearly state the intended uses for the resulting score(s). These can range from research purposes to specific applied uses such as personnel selection or clinical diagnosis. In any case, the intended uses have implications for the setting in which data are collected, criterion-related validity evidence, reliability estimators, and item content. At the same time, it can be advantageous to specifically rule out certain uses for the measure presented. For example, if it is stated that a score should not be used in high-stakes settings, the lack of empirical findings supporting such a

use cannot be criticized. In other words, by specifying certain uses and ruling out others, the test constructor(s) set the stage for possible criticism regarding (the lack of) criterion-related validity evidence.

Data Collection Setting

Psychological tests can be used in different settings. On the most general level, we can differentiate between high-stakes settings and low-stakes settings. While test takers may “gain” something for themselves (e.g., a job, a pension, a diagnosis) depending on the test score obtained in a high-stakes setting (Ziegler, Maaß, Griffith, & Gammon, 2015; Ziegler, MacCann, & Roberts, 2011), no such gains are apparent in a low-stakes setting (e.g., research). The consequences of this are mainly visible when it comes to self-reports and the influence of social desirability in low-stakes settings (Bäckström, 2007; Bäckström & Björklund, 2016; Bäckström, Björklund, & Larsson, 2009) or faking/malingering in high-stakes settings (Griffith, Chmielowski, & Yoshita, 2007). To summarize these findings, such response sets or styles often introduce an additional source of systematic variance, thereby potentially inflating internal consistency estimates and intercorrelations of scores obtained with the same method (Ziegler & Bühner, 2009). At the same time, item means and total scores could be distorted. Thus, there can be a strong influence on all kinds of evidence for a score’s psychometric quality. Not testing a measure in the intended setting could therefore limit the applicability of the measure. At the same time, it is not always possible to test a measure in the intended setting. This, however, should be clearly stated in the paper, and possible limitations, also regarding the uses finally recommended based on the provided evidence, must be discussed.

Criterion-Related Validity Evidence

The implications of the intended uses for criterion-related validity evidence are pretty straightforward. Each submission should include evidence supporting the use of the score(s) and interpretation(s) for the intended applications. Popular examples are correlations between test scores and criteria (e.g., IQ or openness score and school grades). It is also possible to look at mean differences between samples from specific populations. Some measures contain several scores or facets for the construct(s) in question. Here, multivariate analyses are necessary to exemplify the usefulness of each score (Sieglings, Petrides, & Martskvishvili, 2015; Ziegler & Bäckström, 2016). In most cases, this will mean providing evidence for incremental validity. For some scores, the intended use is individual assessment (e.g., personnel selection or clinical diagnoses). Here, criterion-related validity evidence should focus on sensitivity and specificity (Kemper, Trapp,

Kathmann, Samuel, & Ziegler, 2019). Importantly, clear hypotheses regarding the expected differences or relations must be stated a priori.

Reliability Estimators

Scores from psychological tests are often used for status assessment or prognosis. These different uses have different requirements regarding reliability evidence. For status assessment, it is important to showcase the internal consistency of a score. PTAD welcomes estimators that can take structural properties into account (Brunner & Süß, 2005; Revelle & Zinbarg, 2009). If a score is meant for prognoses, reliability evidence should support the stability of a score. Thus, test-retest reliability evidence is vital (Gnambs, 2014, 2015). Again, the use of a score for individual assessment comes along with specific requirements (Sijtsma, 2009; Sijtsma & Emons, 2011). Especially the relation between scale length and reliability must be considered. Here, PTAD strongly encourages the use of the concept of measurement precision (Kemper et al., 2019; Krueger, Emons, & Sijtsma, 2013a, 2013b, 2014). Thus, the choice of reliability evidence should be justified in accordance with the intended use(s).

Item Content

In some cases, the intended use also has implications for item content. For example, if a score is meant for personnel selection, legislation in some countries does not allow to ask about private matters (e.g., typical vacations). Such matters should be explained. Finally, the intended uses might also have implications for the item difficulties needed.

What Is the Intended Target Population?

The answer to this question has two main implications. First of all, it defines the population from which samples should be drawn for all stages of test construction, adaptation, or development. Second, the answer has implications for the content of the items.

Samples

All samples used in the paper must be drawn from the target population. If, for any reason, this is not the case, explanations regarding possible limitations should be provided. Implications of not using samples from the target population(s) could occur for item means or difficulties and thus item intercorrelations or test score intercorrelations. This should be considered when formulating hypotheses. For example, if a sample used is potentially restricted in variance, correlations between

scores could be diminished. This would affect convergent validity evidence negatively. At the same time, discriminant validity evidence could be incorrectly distorted to appear positive.

Item Content

Depending on different characteristics of the sample (e.g., cultural background, age, professional status), the constructs targeted might manifest differently. Consequently, using the same items for different populations can be problematic. Thus, when translating, adapting, or developing a measure, it is vital to demonstrate how the item content relates to characteristics of the target population. For example, pensioners who no longer have a regular job will have problems answering the following conscientiousness item: “I always appear on time at my place of work.” Thus, an adaptation of this item to an elderly target population is not straightforward and requires some substantial thinking and possible pretesting (Ziegler, Kemper, & Lenzner, 2015). The specification of the target population also informs the choice of item difficulties. It is not always wise to simply look for medium difficulties (Ziegler, 2014a).

In conclusion, these explanations highlight how answers to the three ABC questions inform and define the plan to evaluate a measure. Each paper should address all of these issues or discuss whether and how the lack of any such evidence decreases generalizability. For registered reports, the same holds true. Moreover, if necessary, papers should explain why certain evidence has not been pursued. At the end of the introduction, the reader should have a clear understanding of what is being measured, for what purposes, and in how far the ensuing information supports these claims.

Methods

The typical method section should contain information on how data were collected, all measures used, sample demographics, and descriptive statistics for the measures used, as well as a section on the statistical analyses (to be) conducted. Whereas most of these parts are pretty straightforward, the sample section and the analyses section come with specific challenges, especially when writing a registered report.

Sample

Whether the paper is a regular submission or a registered report, sample size needs to be justified. In many cases,

rules of thumb for manifest (Schönbrodt & Perugini, 2013) or latent correlations (Kretzschmar & Gignac, 2019) might suffice. However, there may be instances where more specific a priori power analyses might be required and could be a safeguard against replication failures (Open Science Collaboration, 2015). In such cases, authors often refer to simulations. Here, we just want to point out that, while this is certainly advisable in general, there can be serious problems when the simulations are based on inaccurate assumptions (Albers & Lakens, 2018). Thus, each paper, whether registered report or regular, needs to justify the sample size aimed for or actually acquired.

Statistical Analyses

This section, which is meant to inform the reader about the kind of analyses (to be) conducted and the software used, comes with specific challenges when we adapt or develop tests. In general, during such processes, many decisions need to be made. For example, items might be selected at several stages. Here, the paper needs to accurately inform in detail on decision criteria. In other cases, different reliability estimators could be available (e.g., Cronbach α , McDonald ω , or split half). Here, an explanation of why a certain estimator is used should be provided. All of this shows that the choices for the kind of analyses or estimator used must not only be stated but need to be justified! Importantly, this needs to be done in alignment with the answers provided to the three ABC questions. Thus, a kind of protocol matching the pursued evidence (e.g., criterion-related validity evidence) for psychometric quality to analyses (e.g., correlation, regression, or t -test) and, importantly, the required result(s) (e.g., expected effect size) needs to be defined. This should be especially fastidious when it comes to structural validity evidence. We will highlight this here using the example of structural equation modeling which is often used to provide such evidence. Here, authors should first clearly define their model of choice and, if possible, alternative models. In a next step, they should state how individual models are evaluated and how different models are being compared. Here, I explicitly refer all authors to literature which suggests a more differentiated approach to standard cutoffs for indices such as RMSEA or CFI (e.g., Greiff & Heene, 2017; Heene et al., 2011). Importantly, authors also need to explain what they plan to do (for a registered report) or justify what they did (for a regular submission) when the preferred model did not fit the data. Options might include to delete items, to add correlated residuals, to add crossloadings, or additional latent variables, to name just a few. In a registered report, these choices must be made clear. In all submissions, the consequence of such choices, for example,

regarding content validity or unidimensionality (Ziegler & Hagemann, 2015), need to be considered and stated.

Authors should not forget that this section is vital when it comes to showcasing how reliability and validity evidence was obtained and how trustworthy it is. Any deviations from the planned procedure must be explained in a registered report and possible limitations added. In regular submissions, the same holds true when there are deviations from the assumptions laid out in the introduction.

Results

This section should contain all information necessary to evaluate whether the score(s) from the measure presented in the paper actually measure the intended construct and are useful in the intended way. In registered reports, the kind of information that is planned to be reported can be portrayed. The option of using online supplementary material should be considered.

In general, we encourage that authors share raw data and the code or the converted score data underlying the main findings. Exceptions may be made (e.g., for reasons of data security or confidentiality) provided the reasons are set out during manuscript submission. The general premise here is that any researcher should be able to reproduce the central results of the study without contacting the original authors. This requires open code and open data. Furthermore, it should in principle be possible to replicate the study in an independent sample. This requires open material (i.e., items, instructions, study setup) or a reference to the source of the materials.

Discussion

Within this section, the evidence obtained should be evaluated with regard to the requirements formulated in the introduction. As a result, clear recommendations should be listed. This refers to whether the measured score(s) can be interpreted as intended.

These, in places detailed, elaborations should be understood as a kind of template. Submissions will be expected to follow the structure laid out here and to provide the information outlined (or explain the lack of it). There will be a formal check of each submission, and divergences from this template may result in the paper being sent back with a request for closer alignment with the template. Please keep in mind that in doing this we aim to help both readers and reviewers – and also our authors because the

chances of a fast peer review process and of the paper being cited later on should improve substantially!

At this point, I would once again like to highlight the option of submitting a registered report. Planning a test translation or adaptation will, we hope, be facilitated by following these guidelines – the template can be used in the sense of a checklist. Moreover, the opportunity to obtain timely feedback, before data are being collected, should allow us to weed out problems or even critical flaws that otherwise could not be undone later. This in turn should bolster the paper's final quality.

Now all that remains for me to say is that I look forward to the start of the new journal – and to reading YOUR paper!

References

- AERA, APA, & NMCE. (2014). *Standards for educational & psychological testing*. Washington, DC: AERA Publications.
- Albers, C., & Lakens, D. (2018). When power analyses based on pilot data are biased: Inaccurate effect size estimators and follow-up bias. *Journal of Experimental Social Psychology*, 74, 187–195. <https://doi.org/10.1016/j.jesp.2017.09.004>
- Bäckström, M. (2007). Higher-order factors in a five-factor personality inventory and its relation to social desirability. *European Journal of Psychological Assessment*, 23, 63–70. <https://doi.org/10.1027/1015-5759.23.2.63>
- Bäckström, M., & Björklund, F. (2016). Is the general factor of personality based on evaluative responding? Experimental manipulation of item-popularity in personality inventories. *Personality and Individual Differences*, 96, 31–35. <https://doi.org/10.1016/j.paid.2016.02.058>
- Bäckström, M., Björklund, F., & Larsson, M. R. (2009). Five-factor inventories have a major general factor related to social desirability which can be reduced by framing items neutrally. *Journal of Research in Personality*, 43, 335–344. <https://doi.org/10.1016/j.jrp.2008.12.013>
- Bejar, I. I. (1983). *Achievement testing: Recent advances*. Beverly Hills, CA: Sage Publications.
- Borsboom, D. (2006). When does measurement invariance matter? *Medical Care*, 44, S176–S181. <https://doi.org/10.1097/01.mlr.0000245143.08679.cc>
- Borsboom, D., Mellenbergh, G., & Van Heerden, J. (2003). The theoretical status of latent variables. *Psychological Review*, 110, 203–218. <https://doi.org/10.1037/0033-295x.110.2.203>
- Brunner, M., & Süß, H. M. (2005). Analyzing the reliability of multidimensional measures: An example from intelligence research. *Educational and Psychological Measurement*, 65, 227–240. <https://doi.org/10.1177/0013164404268669>
- Campbell, D. T., & Fiske, D. W. (1959). Convergent and discriminant validation by the multitrait-multimethod matrix. *Psychological Bulletin*, 56, 81–105. <https://doi.org/10.1037/h0046016>
- Chen, F. F. (2008). What happens if we compare chopsticks with forks? The impact of making inappropriate comparisons in cross-cultural research. *Journal of Personality and Social Psychology*, 95, 1005–1018. <https://doi.org/10.1037/a0013193>
- Credé, M., Tynan, M. C., & Harms, P. D. (2017). Much ado about grit: A meta-analytic synthesis of the grit literature. *Journal of*

- Personality and Social Psychology*, 113, 492–511. <https://doi.org/10.1037/pspp0000102>
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297–334. <https://doi.org/10.1007/bf02310555>
- Cronbach, L. J., & Meehl, P. E. (1955). Construct validity in psychological tests. *Psychological Bulletin*, 52, 281–302. <https://doi.org/10.1037/h0040957>
- Fried, E. I., van Borkulo, C. D., Epskamp, S., Schoevers, R. A., Tuerlinckx, F., & Borsboom, D. (2016). Measuring depression over time ... Or not? Lack of unidimensionality and longitudinal measurement invariance in four common rating scales of depression. *Psychological Assessment*, 28, 1354–1367. <https://doi.org/10.1037/pas0000275>
- Gnambs, T. (2014). A meta-analysis of dependability coefficients (test-retest reliabilities) for measures of the Big Five. *Journal of Research in Personality*, 52, 20–28. <https://doi.org/10.1016/j.jrp.2014.06.003>
- Gnambs, T. (2015). Facets of measurement error for scores of the Big Five: Three reliability generalizations. *Personality and Individual Differences*, 84, 84–89. <https://doi.org/10.1016/j.paid.2014.08.019>
- Greiff, S., & Heene, M. (2017). Why psychological assessment needs to start worrying about model fit. *European Journal of Psychological Assessment*, 33, 313–317. <https://doi.org/10.1027/1015-5759/a000450>
- Griffith, R. L., Chmielowski, T., & Yoshita, Y. (2007). Do applicants fake? An examination of the frequency of applicant faking behavior. *Personnel Review*, 36, 341–357. <https://doi.org/10.1108/00483480710731310>
- Heene, M., Hilbert, S., Draxler, C., Ziegler, M., & Bühner, M. (2011). Masking misfit in confirmatory factor analysis by increasing unique variances: A cautionary note on the usefulness of cutoff values of fit indices. *Psychological Methods*, 16, 319–336. <https://doi.org/10.1037/a0024917>
- Kelley, T. L. (1927). *Interpretation of educational measurements*. Oxford, UK: World Book Co.
- Kemper, C. J., Trapp, S., Kathmann, N., Samuel, D. B., & Ziegler, M. (2019). Short versus long scales in clinical assessment: Exploring the trade-off between resources saved and psychometric quality lost using two measures of obsessive-compulsive symptoms. *Assessment*, 26, 767–782. <https://doi.org/10.1177/107319118810057>
- Kretschmar, A., & Gignac, G. E. (2019). At what sample size do latent variable correlations stabilize? *Journal of Research in Personality*, 80, 17–22. <https://doi.org/10.1016/j.jrp.2019.03.007>
- Kruyen, P. M., Emons, W. H. M., & Sijtsma, K. (2013a). On the shortcomings of shortened tests: A literature review. *International Journal of Testing*, 13, 223–248. <https://doi.org/10.1080/15305058.2012.703734>
- Kruyen, P. M., Emons, W. H. M., & Sijtsma, K. (2013b). Shortening the S-STAI: Consequences for research and clinical practice. *Journal of Psychosomatic Research*, 75, 167–172. <https://doi.org/10.1016/j.jpsychores.2013.03.013>
- Kruyen, P. M., Emons, W. H. M., & Sijtsma, K. (2014). Assessing individual change using short tests and questionnaires. *Applied Psychological Measurement*, 38, 201–216. <https://doi.org/10.1177/0146621613510061>
- Loevinger, J. (1957). Objective tests as instruments of psychological theory: Monograph supplement 9. *Psychological Reports*, 3, 635–694. <https://doi.org/10.2466/pr0.3.7.635-694>
- McNeish, D. (2018). Thanks coefficient alpha, we'll take it from here. *Psychological Methods*, 23, 412–433. <https://doi.org/10.1037/met0000144>
- Michell, J. (1997). Quantitative science and the definition of measurement in psychology. *British Journal of Psychology*, 88, 355–383. <https://doi.org/10.1111/j.2044-8295.1997.tb02641.x>
- Michell, J. (2001). Teaching and misteaching measurement in psychology. *Australian Psychologist*, 36, 211–218. <https://doi.org/10.1080/00050060108259657>
- Mussel, P. (2010). Epistemic curiosity and related constructs: Lacking evidence of discriminant validity. *Personality and Individual Differences*, 49, 506–510. <https://doi.org/10.1016/j.paid.2010.05.014>
- Open Science Collaboration. (2015). Estimating the reproducibility of psychological science. *Science*, 349, aac4716. <https://doi.org/10.1126/science.aac4716>
- Revelle, W., & Zinbarg, R. E. (2009). Coefficients alpha, beta, omega, and the glb: Comments on Sijtsma. *Psychometrika*, 74, 145–154. <https://doi.org/10.1007/s11336-008-9102-z>
- Sass, D. A. (2011). Testing measurement invariance and comparing latent factor means within a confirmatory factor analysis framework. *Journal of Psychoeducational Assessment*, 29, 347–363. <https://doi.org/10.1177/0734282911406661>
- Schipolowski, S., Wilhelm, O., & Schroeders, U. (2014). On the nature of crystallized intelligence: The relationship between verbal ability and factual knowledge. *Intelligence*, 46, 156–168. <https://doi.org/10.1016/j.intell.2014.05.014>
- Schneider, W. J., & McGrew, K. S. (2018). The Cattell-Horn-Carroll theory of cognitive abilities. In D. P. Flanagan & E. M. McDonough (Eds.), *Contemporary intellectual assessment: Theories, tests and issues* (Vol. 4, pp. 73–130). New York, NY: The Guilford Press.
- Schönbrodt, F. D., & Perugini, M. (2013). At what sample size do correlations stabilize? *Journal of Research in Personality*, 47, 609–612. <https://doi.org/10.1016/j.jrp.2013.05.009>
- Schweizer, K. (2012). On issues of validity and especially on the misery of convergent validity. *European Journal of Psychological Assessment*, 28, 249–254. <https://doi.org/10.1027/1015-5759/a000156>
- Siegl, A. B., Petrides, K. V., & Martskevich, K. (2015). An examination of a new psychometric method for optimizing multifaceted assessment instruments in the context of trait emotional intelligence. *European Journal of Personality*, 29, 42–54. <https://doi.org/10.1002/per.1976>
- Sijtsma, K. (2009). Correcting fallacies in validity, reliability, and classification. *International Journal of Testing*, 9, 167–194. <https://doi.org/10.1080/15305050903106883>
- Sijtsma, K., & Emons, W. H. (2011). Advice on total-score reliability issues in psychosomatic measurement. *Journal of Psychosomatic Research*, 70, 565–572. <https://doi.org/10.1016/j.jpsychores.2010.11.002>
- Wehner, C., Roemer, L., & Ziegler, M. (2018). Construct validity. In V. Zeigler-Hill & T. K. Shackelford (Eds.), *Encyclopedia of personality and individual differences* (pp. 1–3). Cham, Germany: Springer International Publishing.
- Wilhelm, O. (2005). Measuring reasoning ability. In O. Wilhelm & R. W. Engle (Eds.), *Handbook of understanding and measuring intelligence* (pp. 373–392). Thousand Oaks, CA: Sage.
- Ziegler, M. (2014a). Comments on item selection procedures. *European Journal of Psychological Assessment*, 30, 1–2. <https://doi.org/10.1027/1015-5759/a000196>
- Ziegler, M. (2014b). Stop and state your intentions!: Let's not forget the ABC of test construction. *European Journal of Psychological Assessment*, 30, 239–242. <https://doi.org/10.1027/1015-5759/a000228>
- Ziegler, M., & Bäckström, M. (2016). 50 facets of a trait 50 ways to mess up? *European Journal of Psychological Assessment*, 32, 105–110. <https://doi.org/10.1027/1015-5759/a000372>
- Ziegler, M., & Bensch, D. (2013). Lost in translation: Thoughts regarding the translation of existing psychological measures into other languages. *European Journal of Psychological Assessment*, 29, 81–83. <https://doi.org/10.1027/1015-5759/a000167>

- Ziegler, M., Booth, T., & Bensch, D. (2013). Getting entangled in the nomological net. *European Journal of Psychological Assessment*, 29, 157–161. <https://doi.org/10.1027/1015-5759/a000173>
- Ziegler, M., & Bühner, M. (2009). Modeling socially desirable responding and its effects. *Educational and Psychological Measurement*, 69, 548–565. <https://doi.org/10.1177/0013164408324469>
- Ziegler, M., & Hagemann, D. (2015). Testing the unidimensionality of items: Pitfalls and loopholes. *European Journal of Psychological Assessment*, 31, 231–237. <https://doi.org/10.1027/1015-5759/a000309>
- Ziegler, M., Kemper, C. J., & Lenzner, T. (2015). The issue of fuzzy concepts in test construction and possible remedies. *European Journal of Psychological Assessment*, 31, 1–4. <https://doi.org/10.1027/1015-5759/a000255>
- Ziegler, M., Maaß, U., Griffith, R., & Gammon, A. (2015). What is the nature of faking? Modeling distinct response patterns and quantitative differences in faking at the same time. *Organizational Research Methods*, 18, 679–703. <https://doi.org/10.1177/1094428115574518>
- Ziegler, M., MacCann, C., & Roberts, R. D. (2011). Faking: Knowns, unknowns, and points of contention. In M. Ziegler, C. MacCann, & R. R. Roberts (Eds.), *New perspectives on faking in personality assessment* (pp. 3–16). New York, NY: Oxford University Press.
- Zinbarg, R. E., Revelle, W., Yovel, I., & Li, W. (2005). Cronbach's α , Revelle's β , and McDonald's ω H: Their relations with each other and two alternative conceptualizations of reliability. *Psychometrika*, 70, 123–133. <https://doi.org/10.1007/s11336-003-0974-7>

Published online June 25, 2020

Matthias Ziegler

Department of Psychology
Humboldt University Berlin
Rudower Chaussee 18
12489 Berlin
Germany
zieglema@hu-berlin.de