



FROM THE RESEARCH BENCH

Measurement invariance before comparing English and German scores

A translated scale is not automatically comparable: configural, metric, and scalar invariance license different cross-language claims.

Sinan Isoglu, MBA (Quantic)

Commercial growth leader, lecturer and doctoral researcher

PUBLISHED

1 September 2026

READING TIME

11 min

WORDS

2,514

<https://isoglu.com/insights/journal/measurement-invariance-before-comparing-english-and-german-scores/>

Management summary	2
Why does linguistic translation fail to establish cross-cultural survey comparability?	3
Which three distinct statistical hypotheses do configural, metric, and scalar invariance test?	3
Why must researchers declare the specific cross-national comparison before invariance testing?	4
Why does bilingual equivalence in one psychometric instrument not transfer to others?	5
How do exact, partial, and Bayesian approximate invariance alter permitted conclusions?	5
Why does survey score standardization fail to correct for underlying measurement non-invariance?	6
Why must measurement invariance be re-established whenever survey platforms change?	7
How should research teams state cross-national findings strictly at the verified invariance level?	8
Where are the psychometric limits of cross-linguistic measurement invariance?	8
References	9

MANAGEMENT SUMMARY

Translating an instrument and proving that its English and German scores are comparable are different jobs. Translation addresses wording and meaning; measurement invariance asks whether the same construct has the same measurement structure across groups. Configural invariance concerns the factor pattern, metric invariance comparable loadings and relations, and scalar invariance latent-mean comparisons under the specified model. Cieciuch and colleagues distinguish exact and approximate results, while Klopp and Klößner separate scaling restrictions from invariance conditions. Prem and colleagues provide a bounded English and German example for one four-subscale instrument with separate UK and German samples. This article builds a bilingual release gate. It does not validate a private instrument, generalize one study, or authorize a universal language claim.

Keywords: Measurement invariance · Configural invariance · Metric invariance · Scalar invariance · Back-translation

A questionnaire can be translated into English and German, administered to two groups, and produce two sets of scores. None of those steps proves that a difference between the scores reflects a difference in the construct.

The short answer is that translation completion is not measurement-release completion. Translation and back-translation address wording, meaning, and cultural fit. Measurement invariance tests whether the instrument represents and measures the same latent construct across groups. Configural, metric, and scalar invariance answer progressively stronger questions, and each level licenses a different comparison.

If a team wants to compare latent means, it needs evidence for the relevant scalar condition under its specified model. If it wants to compare relations such as covariances or regression coefficients, metric invariance may be the relevant boundary. A favorable translation review, a good overall fit, or an existing result in another study cannot silently authorize all three claims.

This is a methods article, not a release decision for a private instrument. The defensible evidence review explains why a claim should also have a stopping rule. RES-07 applies that discipline to the narrower problem of comparing English and German scores.

Why does linguistic translation fail to establish cross-cultural survey comparability?

Translation asks whether an item has been rendered in another language with sufficient semantic and pragmatic care. A good process can use bilingual experts, independent versions, reconciliation, cognitive interviews, and back-translation. These steps can expose wording drift, missing concepts, unnatural phrases, or culture-specific assumptions.

Measurement invariance asks a different question: after the translation has been made, does the measurement model operate comparably across groups? The object is not a sentence in isolation. It is the relation among items, latent factors, loadings, intercepts or thresholds, residual structure, and the populations that answered them.

The distinction prevents two opposite errors. One error treats a statistically acceptable model as proof that the translation is semantically excellent. The other treats a careful translation as proof that raw scores can be compared. Translation evidence and invariance evidence can support each other. They do not replace each other.

Cieciuch, Davidov, Schmidt, Algesheimer, and Schwartz state that measurement invariance is necessary for meaningful comparisons. Their wording is a release boundary, not a demand to reject every translated scale. It means that the intended comparison must be matched to the measurement evidence that the comparison needs.

Which three distinct statistical hypotheses do configural, metric, and scalar invariance test?

The familiar sequence is configural, metric, and scalar invariance. The names matter less than the restriction each name represents.

Configural invariance asks whether the same factor pattern is plausible across groups. The same items are associated with the same latent factors, in the same basic structure. This is a structural starting point. It does not mean that an item contributes equally or that the groups have the same latent mean.

Metric invariance adds equality constraints on factor loadings. If the loadings are comparable, a unit of the latent construct has a more comparable relation to the observed indicators across groups. This is the level that can support comparisons involving relations such as covariances or unstandardized regression coefficients in the specified model. It still does not, by itself, license a claim that one group has a higher latent mean.

Scalar invariance adds equality constraints on indicator intercepts, or the corresponding thresholds when the model requires them. The additional restriction addresses systematic starting-point differences between groups. Under the model and assumptions, scalar invariance is the relevant stronger condition for latent-mean comparison.

These levels are not a quality ladder in which a lower result means the instrument is useless. They are comparison boundaries. A scale can support a relationship analysis while not supporting a latent-mean comparison. The article or decision should use the strongest language the evidence supports, not the strongest level the team hoped to reach.

Figure 1 The bilingual measurement-release gate

RELEASE QUESTION	EVIDENCE OR TEST	LEVEL OR METHOD	COMPARISON PERMITTED	COMPARISON NOT LICENSED
Construct, groups, languages, comparison?	Translation, loadings, intercepts, thresholds, sample, mode?	Configural, metric, scalar, partial or approximate?	What exact sentence can be stated from the result?	Which stronger claim remains outside the evidence, and why?

Five rows are blank reader inputs. The worksheet is a release gate, not a statistical result or universal translation standard.

Author's release worksheet grounded in Cieciuch et al. (2014), Klopp and Klößner (2023), and Prem et al. (2021). Blank fields are reader inputs; no instrument, participant, or score data is supplied.

Why must researchers declare the specific cross-national comparison before invariance testing?

Researchers often ask whether a scale is invariant as if the word referred to one yes-or-no property. The more useful question is invariant for what comparison, in which groups, under which model, and in which sample?

Suppose the intended claim is that English and German respondents relate two constructs in the same way. The relevant evidence may be metric invariance and a declared model for the relation. Suppose the claim is that German respondents have a higher average level of a latent construct. That claim needs the stronger scalar boundary, plus a defensible model for the latent mean. Suppose the claim is only that both versions represent the same conceptual dimensions. Configural evidence may be an important starting point, but the language should remain limited.

The distinction also applies to observed item means. Raw scores are not rescued by calling them simple. If an item has a different intercept or threshold across languages, the same observed response can encode a different location on the latent construct. A raw mean difference can then contain measurement noncomparability as well as a substantive difference.

The release record should therefore state the comparison before the fit indices are reviewed. Otherwise a team can discover a favorable result and only afterward choose a claim that the result does not support.

Why does bilingual equivalence in one psychometric instrument not transfer to others?

Prem and colleagues provide an especially relevant example because their instrument was developed and initially validated in German and English. The CODE scale has four subscales and was examined across three independent studies with an overall sample of 1,129. The authors used native speakers and a back-translation method. They then tested configural, metric, and scalar invariance and wrote that “the means of the subscales could be compared”.

Those facts support a precise sentence about the cited instrument and study. They do not support a claim about every English-German scale. The English study used 274 UK employees, while the German study used 303 employees in Germany. These are separate samples from different study settings, not a within-person language-switch experiment.

The study’s limitations also remain part of the evidence. Prem and colleagues note cross-sectional data, European samples, sample-specific validation, and the need to validate the instrument in other contexts . A result can be encouraging and still be bounded by instrument, sample, collection mode, and model.

The proper use of the example is therefore demonstrative. It shows how a translation process can be followed by a measurement comparison. It does not make the sequence optional for another scale, and it does not grant permission to compare scores from a new sample merely because the item wording was copied.

How do exact, partial, and Bayesian approximate invariance alter permitted conclusions?

Exact invariance imposes equality constraints as specified by the model. It is a strong and transparent test, but real instruments can contain small cross-group differences that make exact equality too rigid for the research question.

Partial invariance releases selected parameters while retaining enough equality constraints for the model and comparison. The release must be declared and justified. Partial invariance is not the same as saying that all parameters are approximately equal.

Approximate invariance permits small parameter differences through an explicit tolerance or prior distribution. Cieciuch and colleagues distinguish this approach from partial invariance and describe approximate equality as a different restriction choice. A favorable approximate result is not exact invariance with softer language. It is evidence under a different model.

Their eight-country human-values illustration makes the distinction concrete. In the prior analysis, exact scalar invariance was established for 10 of 19 values. Under the specified Bayesian approximate approach, approximate

scalar invariance was established for all 19 values. The contrast is a result for that instrument, sample, and prior specification. It is not a general count for translated instruments.

The source also reports convenience student samples, mixed collection modes, unequal sample sizes, and limitations related to ordinal scores. These caveats do not invalidate the illustration. They define how far its conclusion can travel.

Figure 2 What each invariance result licenses

Result or method	What is constrained	Comparison it can support in the specified model	Stronger claim not licensed by the label alone
Configural invariance	Factor pattern and basic construct representation.	A claim that the same structural pattern is being examined across groups.	Equal item meaning, equal relations, or equal latent means.
Metric invariance	Factor loadings across groups.	Comparisons involving specified relations such as covariances or unstandardized regressions.	A higher or lower latent mean.
Scalar invariance	Loadings plus intercepts or thresholds as specified.	Latent-mean comparisons under the declared model and sample.	Permanence across samples, modes, translations, or instruments.
Partial invariance	Selected parameters are released while others remain equal.	The narrower comparison justified by the retained constraints and identification.	Treating every parameter as equal.
Approximate invariance	Small differences are allowed under an explicit tolerance or prior.	The comparison justified by that approximate model and its assumptions.	Calling the result exact or applying its tolerance universally.

Author's comparison framework grounded in Cieciuch et al. (2014). The rows state methodological boundaries, not results for a new instrument or sample.

Why does survey score standardization fail to correct for underlying measurement non-invariance?

Technical discussions can introduce another source of confusion: model scaling. Latent-variable models need scaling restrictions for identification. A scaling rule tells the model how a latent variable is anchored. It is not substantive evidence that an English and German version measures a construct equivalently.

Klopp and Klößner separate scaling restrictions from the invariance conditions being tested. Their paper focuses on metric measurement-invariance models. In the relevant setup, they write, “Apply the scaling restriction in one group only”. The instruction concerns correct model identification and comparison of fit quantities. It does not turn a referent indicator into proof of translation quality or scalar invariance.

The distinction is easy to lose in a software output. A model can converge because its scaling restrictions are adequate. That convergence does not answer whether loadings or intercepts are equal. The analyst needs to identify which constraints are being tested, which are imposed for identification, and which are released or tolerated.

Klopp and Klößner also show why a one-indicator partial metric model can be equivalent to the configural model in the relevant setup. The technical result is a warning against treating the presence of a restriction as proof that a meaningful extra comparison has been established. Degrees of freedom and identification choices are part of the evidence record, not decorative details.

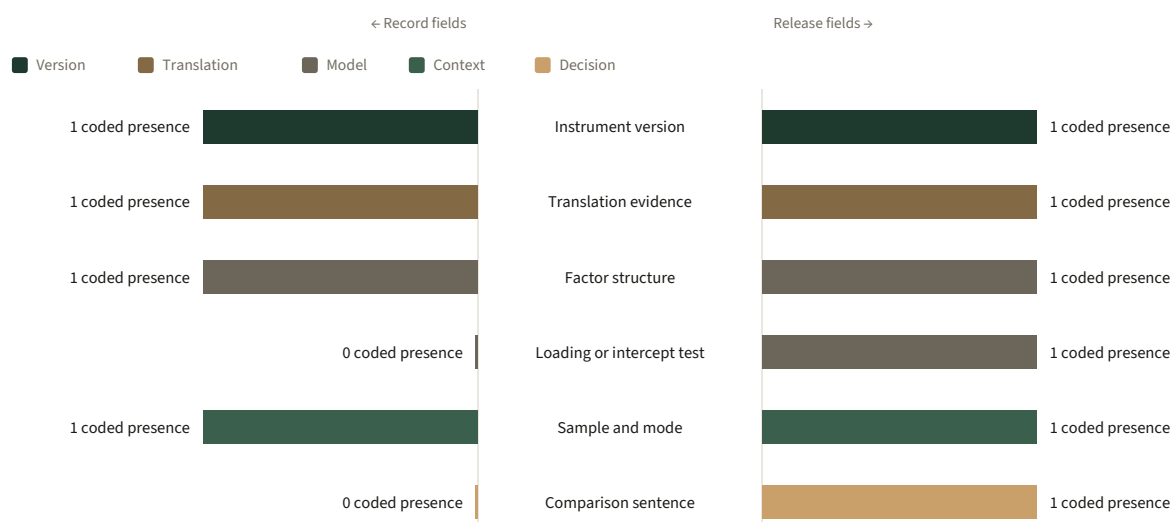
Why must measurement invariance be re-established whenever survey platforms change?

Measurement invariance is not a permanent certificate attached to a questionnaire. Cieciuch and colleagues write: “Establishing measurement invariance in one study does not signify that a questionnaire is always measurement invariant.” The result belongs to a scale, translation, sample, collection mode, factor model, and comparison.

That means a new language edition, item revision, response mode, population, or administration context can reopen the gate. A web form is not necessarily the same measurement context as a supervised paper survey. An employee sample is not necessarily the same population as customers or students. A changed example in an item can change its cognitive demand even when the construct label remains the same.

The practical record should version the instrument and the test. Keep the item text, translation decision, back-translation note, sample definition, collection mode, model specification, constraints, fit evidence, released comparison, and unresolved limitation together. A later reader should be able to see which result belongs to which version.

Figure 3 From translation version to comparison sentence



Author's schematic framework grounded in Cieciuch et al. (2014), Klopp and Klößner (2023), and Prem et al. (2021). The sequence is a governance model, not a claim that every study follows these exact steps.

How should research teams state cross-national findings strictly at the verified invariance level?

The final gate is linguistic. State what the result allows, then state what it does not allow.

For configural evidence, a careful sentence might say that the same factor pattern was examined across the English and German groups under the specified model. For metric evidence, the sentence can address the relevant relations if the model and sample support that use. For scalar evidence, it can state that a latent-mean comparison was supported for the specified instrument, groups, and model. The sentence should not become English and German respondents are equivalent because that wording is broader than the test.

Partial and approximate results need even more precision. Name which parameters were released or what tolerance was used. Report the comparison that remains defensible and the stronger comparison that remains open. A smaller claim with a visible boundary is more useful than a universal statement that cannot be reproduced.

The release gate should also have a stop rule. Stop the comparison when the intended construct is not stable, the translation revision is not recorded, the sample or mode changed without a new test, the relevant invariance level is not addressed, or the model output cannot distinguish identification from substantive equality. Not yet comparable is a valid methods result.

Where are the psychometric limits of cross-linguistic measurement invariance?

Cieciuch and colleagues explain the levels and the exact-versus-approximate distinction in a bounded values-study illustration. Klopp and Klößner clarify technical scaling issues for metric invariance models. Prem and colleagues provide one four-subscale English and German example with its sample and instrument limits. Together, the sources do not validate every translation, instrument, population, website, survey, or language comparison.

The localization article addresses language as a market-entry decision. RES-07 addresses a narrower measurement question: whether a score comparison is licensed by the evidence. The evidence-review stopping rule provides a related discipline for deciding when a claim should pause.

The reliable conclusion is modest: translate carefully, test the measurement model, name the intended comparison, and write only the sentence that the invariance result supports. A favorable result belongs to the instrument, version, sample, mode, and model that produced it. When any of those change, the comparison gate opens again.

- Cieciuch, J., Davidov, E., Schmidt, P., Algesheimer, R., & Schwartz, S. H. (2014). Comparing results of an exact vs. an approximate (Bayesian) measurement invariance test: A cross-country illustration with a scale to measure 19 human values. *Frontiers in Psychology*, 5, 982. <https://doi.org/10.3389/fpsyg.2014.00982>
- Klopp, E., & Klößner, S. (2023). Scaling metric measurement invariance models. *Methodology*, 19(3), 192-227. <https://doi.org/10.5964/meth.10177>
- Prem, R., Kubicek, B., Uhlig, L., Baumgartner, V., & Korunka, C. (2021). Development and initial validation of a scale to measure cognitive demands of flexible work. *Frontiers in Psychology*, 12, 679471. <https://doi.org/10.3389/fpsyg.2021.679471>

END OF ESSAY

HOW TO CITE THIS ESSAY

Isoglu, S. (2026, September 1). Measurement invariance before comparing English and German scores. Sinan Isoglu.
<https://isoglu.com/insights/journal/measurement-invariance-before-comparing-english-and-german-scores/>

KEEP READING

From the research bench

What is selection bias? The sample can change the answer

<https://isoglu.com/insights/journal/what-is-selection-bias/>

From the research bench

Staggered difference-in-differences needs a cohort-time estimand

<https://isoglu.com/insights/journal/staggered-difference-in-differences-needs-a-cohort-time-estimand/>

From the research bench

What are parallel trends? The assumption behind difference-in-differences

<https://isoglu.com/insights/journal/what-are-parallel-trends/>



Sinan Isoglu, MBA (Quantic)
Commercial growth leader, lecturer and doctoral researcher