



AUS DER FORSCHUNG

Messinvarianz vor dem Vergleich englischer und deutscher Werte

Eine übersetzte Skala ist nicht automatisch vergleichbar: konfigurale, metrische und skalare Invarianz erlauben verschiedene Sprachvergleiche.

Sinan Isoglu, MBA (Quantic)

Führungskraft für Umsatzwachstum, Dozent und Doktorand

VERÖFFENTLICHT

1. September 2026

LESEZEIT

11 Min.

WÖRTER

2.377

<https://isoglu.com/de/insights/journal/messinvarianz-vor-dem-vergleich-englischer-und-deutscher-werte/>

Management Summary	2
Warum garantiert eine fehlerfreie Übersetzung noch keine interkulturelle Vergleichbarkeit?	3
Welche drei statistischen Hypothesen prüfen konfigurale, metrische und skalare Invarianz?	3
Warum muss der beabsichtigte Ländervergleich vor der Invarianzprüfung feststehen?	4
Warum gilt eine bilinguale Äquivalenzprüfung stets nur für ein konkretes Instrument?	5
Wie unterscheiden sich exakte, partielle und approximative Invarianz in den Befugnissen?	5
Warum kann eine Standardisierung von Befragungswerten fehlende Invarianz nicht heilen?	7
Warum erfordert jede Änderung der Erhebungsplattform eine erneute Invarianzprüfung?	7
Wie formulieren Forschungsteams Befunde exakt auf der empirisch nachgewiesenen Stufe?	8
Wo liegen die methodischen Grenzen psychometrischer Messinvarianzprüfungen?	9
Literatur	10

MANAGEMENT SUMMARY

Ein Instrument zu übersetzen und die Vergleichbarkeit englischer und deutscher Werte zu belegen, sind zwei Aufgaben. Übersetzung betrifft Formulierung und Bedeutung; Messinvarianz fragt nach derselben Messstruktur über Gruppen hinweg. Konfigurale Invarianz betrifft das Faktorenmuster, metrische Invarianz Ladungen und Beziehungen, skalare Invarianz unterstützt unter dem Modell Vergleiche latenter Mittelwerte. Cieciuch und Kollegen unterscheiden exakte und approximative Ergebnisse. Klopp und Klößner trennen Skalierungsrestriktionen von Invarianzbedingungen. Prem und Kollegen liefern ein begrenztes englisch-deutsches Beispiel mit vier Subskalen und getrennten britischen und deutschen Stichproben. Der Beitrag baut ein Freigabegate und verallgemeinert keine Einzelstudie zu einer universellen Sprachaussage.

Schlagwörter: Messinvarianz · Konfigurale Invarianz · Metrische Invarianz · Skalare Invarianz · Rückübersetzung

Ein Fragebogen kann ins Englische und Deutsche übersetzt, zwei Gruppen vorgelegt und mit zwei Wertesätzen ausgewertet werden. Keiner dieser Schritte beweist, dass ein Unterschied zwischen den Werten einen Unterschied im Konstrukt abbildet.

Die kurze Antwort lautet: Der Abschluss einer Übersetzung ist nicht der Abschluss der Messfreigabe. Übersetzung und Rückübersetzung betreffen Formulierung, Bedeutung und kulturelle Passung. Messinvarianz prüft, ob das Instrument dasselbe latente Konstrukt in verschiedenen Gruppen repräsentiert und misst. Konfigurale, metrische und skalare Invarianz beantworten schrittweise stärkere Fragen und erlauben jeweils andere Vergleiche.

Wenn ein Team latente Mittelwerte vergleichen will, braucht es unter dem spezifizierten Modell Evidenz für die relevante skalare Bedingung. Wenn es Beziehungen wie Kovarianzen oder Regressionskoeffizienten vergleichen will, kann metrische Invarianz die relevante Grenze sein. Eine gute Übersetzungsprüfung, ein guter Gesamt-Fit oder ein Ergebnis aus einer anderen Studie darf nicht stillschweigend alle drei Aussagen freigeben.

Dieser Beitrag ist ein Methodenartikel und keine Freigabeentscheidung für ein privates Instrument. Der Abbruchregel-Beitrag für belastbare Evidenzreviews erklärt, warum eine Aussage außerdem eine Abbruchregel braucht. RES-07 wendet diese Disziplin auf die engere Frage an, wann englische und deutsche Werte verglichen werden dürfen.

Warum garantiert eine fehlerfreie Übersetzung noch keine interkulturelle Vergleichbarkeit?

Übersetzung fragt, ob ein Item mit ausreichender semantischer und pragmatischer Sorgfalt in eine andere Sprache übertragen wurde. Ein guter Prozess kann bilinguale Fachpersonen, unabhängige Fassungen, Abstimmung, kognitive Interviews und Rückübersetzung nutzen. Diese Schritte können Bedeutungsverschiebung, fehlende Konzepte, unnatürliche Wendungen oder kulturspezifische Annahmen sichtbar machen.

Messinvarianz stellt eine andere Frage: Funktioniert das Messmodell nach der Übersetzung über Gruppen hinweg vergleichbar? Das Objekt ist nicht allein ein Satz. Es ist die Beziehung zwischen Items, latenten Faktoren, Ladungen, Intercepts oder Schwellen, Residualstruktur und den Populationen, die geantwortet haben.

Diese Unterscheidung verhindert zwei entgegengesetzte Fehler. Der erste behandelt ein statistisch akzeptables Modell als Beweis für eine semantisch hervorragende Übersetzung. Der zweite behandelt eine sorgfältige Übersetzung als Beweis dafür, dass Rohwerte verglichen werden dürfen. Übersetzungs- und Invarianzevidenz können sich ergänzen. Sie ersetzen einander nicht.

Cieciuch, Davidov, Schmidt, Algesheimer und Schwartz schreiben, dass Messinvarianz für sinnvolle Vergleiche notwendig ist. Diese Aussage ist eine Freigabegrenze und keine Aufforderung, jede übersetzte Skala abzulehnen. Sie bedeutet, dass der geplante Vergleich an die Messevidenz angepasst werden muss, die dieser Vergleich benötigt.

Welche drei statistischen Hypothesen prüfen konfigurale, metrische und skalare Invarianz?

Die bekannte Abfolge lautet konfigurale, metrische und skalare Invarianz. Entscheidend ist weniger der Name als die Restriktion, die jede Bezeichnung beschreibt.

Konfigurale Invarianz fragt, ob über Gruppen hinweg dasselbe Faktorenmuster plausibel ist. Dieselben Items sind mit denselben latenten Faktoren in derselben Grundstruktur verbunden. Das ist ein struktureller Ausgangspunkt. Es bedeutet nicht, dass ein Item gleich stark beiträgt oder dass Gruppen denselben latenten Mittelwert haben.

Metrische Invarianz ergänzt Gleichheitsrestriktionen für Faktorladungen. Sind die Ladungen vergleichbar, hat eine Einheit des latenten Konstrukts in den Gruppen eine vergleichbarere Beziehung zu den beobachteten Indikatoren. Diese Ebene kann im spezifizierten Modell Vergleiche von Beziehungen wie Kovarianzen oder unstandardisierten Regressionskoeffizienten unterstützen. Sie erlaubt für sich genommen noch nicht die Aussage, dass eine Gruppe einen höheren latenten Mittelwert hat.

Skalare Invarianz ergänzt Gleichheitsrestriktionen für Item-Intercepts oder die entsprechenden Schwellen, wenn das Modell dies verlangt. Die zusätzliche Restriktion betrifft systematische Ausgangsunterschiede zwischen Gruppen. Unter Modell und Annahmen ist skalare Invarianz die stärkere Bedingung für den Vergleich latenter Mittelwerte.

Diese Ebenen sind keine Qualitätsleiter, auf der ein niedrigeres Ergebnis das Instrument nutzlos macht. Sie sind Vergleichsgrenzen. Eine Skala kann eine Beziehungsanalyse unterstützen, ohne einen Vergleich latenter Mittelwerte zu erlauben. Artikel oder Entscheidung sollten die stärkste Sprache verwenden, die die Evidenz trägt, nicht die stärkste Ebene, die das Team erreichen wollte.

Abbildung 1 Das zweisprachige Freibegate für Messung

FREIGABEFRAGE	EVIDENZ ODER TEST	EBENE ODER METHODE	VERGLEICH ERLAUBT	VERGLEICH NICHT FREIGEgeben
Konstrukt, Gruppen, Sprachen, Vergleich?	Übersetzung, Ladungen, Intercepts, Schwellen, Stichprobe, Modus?	Konfigural, metrisch, skalar, partiell oder approximativ?	Welcher genaue Satz kann aus dem Ergebnis abgeleitet werden?	Welche stärkere Aussage bleibt außerhalb der Evidenz und warum?

Fünf Zeilen sind leere Eingabefelder. Das Arbeitsblatt ist ein Freibegate, kein statistisches Ergebnis und kein universeller Übersetzungsstandard.

Freigabe-Arbeitsblatt des Autors, begründet mit Cieciuch et al. (2014), Klopp und Klößner (2023) sowie Prem et al. (2021). Leere Felder sind Eingaben der Lesenden; Instrument-, Teilnehmenden- und Wertedaten sind nicht enthalten.

Warum muss der beabsichtigte Ländervergleich vor der Invarianzprüfung feststehen?

Forschende fragen oft, ob eine Skala invariant ist, als wäre das Wort eine einzige Ja-oder-Nein- Eigenschaft. Die nützlichere Frage lautet: invariant für welchen Vergleich, in welchen Gruppen, unter welchem Modell und in welcher Stichprobe?

Angenommen, die geplante Aussage lautet, dass englische und deutsche Befragte zwei Konstrukte gleich miteinander verbinden. Dann kann metrische Invarianz für die Beziehung im erklärten Modell relevant sein. Angenommen, die Aussage lautet, deutsche Befragte hätten ein höheres durchschnittliches Niveau eines latenten Konstrukts. Dann braucht diese Aussage die stärkere skalare Grenze und ein tragfähiges Modell für den latenten Mittelwert. Angenommen, die Aussage lautet nur, dass beide Fassungen dieselben konzeptuellen Dimensionen darstellen. Dann kann konfigurale Evidenz ein wichtiger Anfang sein. Die Sprache muss aber begrenzt bleiben.

Die Unterscheidung gilt auch für beobachtete Itemmittelwerte. Rohwerte werden nicht dadurch gerettet, dass man sie einfach nennt. Wenn ein Item über Sprachen hinweg einen anderen Intercept oder eine andere Schwelle hat, kann dieselbe beobachtete Antwort eine andere Position auf dem latenten Konstrukt bedeuten. Eine Differenz der Rohmittelwerte kann dann Messnichtvergleichbarkeit und einen substantiellen Unterschied enthalten.

Das Freigabeprotokoll sollte den Vergleich deshalb festhalten, bevor Fit-Indizes geprüft werden. Sonst kann ein Team zuerst ein günstiges Ergebnis finden und erst danach eine Aussage auswählen, die das Ergebnis nicht trägt.

Warum gilt eine bilinguale Äquivalenzprüfung stets nur für ein konkretes Instrument?

Prem und Kollegen liefern ein besonders passendes Beispiel, weil ihr Instrument auf Deutsch und Englisch entwickelt und zunächst validiert wurde. Die CODE-Skala hat vier Subskalen und wurde in drei unabhängigen Studien mit insgesamt 1.129 Personen untersucht. Die Autoren verwendeten muttersprachliche Personen und eine Rückübersetzungsmethode. Anschließend prüften sie konfigurale, metrische und skalare Invarianz und schrieben: “the means of the subscales could be compared”. Eigene Übersetzung: Die Mittelwerte der Subskalen konnten verglichen werden.

Diese Fakten stützen einen präzisen Satz über das genannte Instrument und die zitierte Studie. Sie stützen keine Aussage über jede englisch-deutsche Skala. Die englische Studie verwendete 274 Beschäftigte aus dem Vereinigten Königreich, die deutsche Studie 303 Beschäftigte in Deutschland. Das sind getrennte Stichproben aus unterschiedlichen Studienkontexten, kein Sprachwechseltest mit denselben Personen.

Auch die Einschränkungen gehören zur Evidenz. Prem und Kollegen nennen querschnittliche Daten, europäische Stichproben, stichprobenspezifische Validierung und den Bedarf an Validierung in anderen Kontexten. Ein Ergebnis kann ermutigend sein und dennoch durch Instrument, Stichprobe, Erhebungsmodus und Modell begrenzt bleiben.

Das Beispiel sollte daher demonstrativ genutzt werden. Es zeigt, wie auf eine Übersetzung ein Messvergleich folgen kann. Es macht die Reihenfolge für eine andere Skala nicht optional und erlaubt keinen Vergleich einer neuen Stichprobe, nur weil der Itemtext kopiert wurde.

Wie unterscheiden sich exakte, partielle und approximative Invarianz in den Befugnissen?

Exakte Invarianz setzt die im Modell beschriebenen Gleichheitsrestriktionen. Das ist ein starker und transparenter Test. Reale Instrumente können jedoch kleine Gruppenunterschiede enthalten, die exakte Gleichheit für eine konkrete Forschungsfrage zu starr machen.

Partielle Invarianz gibt ausgewählte Parameter frei, während genügend Gleichheitsrestriktionen für Modell und Vergleich erhalten bleiben. Die Freigabe muss erklärt und begründet werden. Partielle Invarianz ist nicht dasselbe wie die Aussage, alle Parameter seien approximativ gleich.

Approximative Invarianz erlaubt kleine Parameterunterschiede innerhalb einer ausdrücklich gesetzten Toleranz oder Priorverteilung. Cieciuch und Kollegen unterscheiden diesen Ansatz von partieller Invarianz und beschreiben approximative Gleichheit als eine andere Restriktionsentscheidung. Ein günstiges approximatives Ergebnis ist nicht exakte Invarianz mit vorsichtigerer Sprache. Es ist Evidenz unter einem anderen Modell.

Die Illustration zu menschlichen Werten aus acht Ländern macht den Unterschied greifbar. In der früheren Analyse wurde für 10 von 19 Werten exakte skalare Invarianz festgestellt. Unter dem spezifizierten Bayes'schen approximativen Ansatz wurde für alle 19 Werte approximative skalare Invarianz festgestellt. Dieser Gegensatz gehört zu Instrument, Stichprobe und Prior-Spezifikation. Er ist keine allgemeine Zahl für übersetzte Instrumente.

Die Quelle berichtet außerdem über studentische Gelegenheitsstichproben, unterschiedliche Erhebungsmodi, ungleiche Stichprobengrößen und Einschränkungen im Umgang mit ordinalen Werten. Diese Hinweise entwerten die Illustration nicht. Sie bestimmen, wie weit ihre Schlussfolgerung reichen darf.

Abbildung 2 Was jedes Invarianzergebnis erlaubt

Ergebnis oder Methode	Was wird restringiert?	Welcher Vergleich ist im erklärten Modell möglich?	Welche stärkere Aussage erlaubt das Etikett allein nicht?
Konfigurale Invarianz	Faktorenmuster und grundlegende Konstruktstruktur.	Aussage, dass dieselbe Struktur über Gruppen hinweg geprüft wird.	Gleiche Itembedeutung, gleiche Beziehungen oder gleiche latente Mittelwerte.
Metrische Invarianz	Faktorladungen über Gruppen hinweg.	Vergleiche bestimmter Beziehungen wie Kovarianzen oder unstandardisierter Regressionen.	Ein höherer oder niedriger latenter Mittelwert.
Skalare Invarianz	Ladungen sowie spezifizierte Intercepts oder Schwellen.	Vergleiche latenter Mittelwerte unter dem erklärten Modell und in der erklärten Stichprobe.	Gültigkeit über Stichproben, Modi, Übersetzungen oder Instrumente hinweg.
Partielle Invarianz	Ausgewählte Parameter werden freigegeben, andere bleiben gleich.	Der engere Vergleich, den die verbleibenden Restriktionen und Identifikation rechtfertigen.	Die Gleichheit aller Parameter.
Approximative Invarianz	Kleine Unterschiede werden unter einer Toleranz oder Priorverteilung zugelassen.	Der Vergleich, den das approximative Modell und seine Annahmen rechtfertigen.	Das Ergebnis als exakt zu bezeichnen oder die Toleranz universell anzuwenden.

Vergleichsrahmen des Autors, begründet mit Cieciuch et al. (2014). Die Zeilen beschreiben methodische Grenzen und keine Ergebnisse für ein neues Instrument oder eine neue Stichprobe.

Warum kann eine Standardisierung von Befragungswerten fehlende Invarianz nicht heilen?

Technische Diskussionen führen oft eine weitere Verwechslung ein: die Skalierung des Modells. Modelle mit latenten Variablen brauchen Skalierungsrestriktionen zur Identifikation. Eine Skalierungsregel zeigt, wie eine latente Variable verankert wird. Sie ist keine inhaltliche Evidenz dafür, dass eine englische und eine deutsche Fassung ein Konstrukt gleichwertig messen.

Klopp und Klößner trennen Skalierungsrestriktionen von den geprüften Invarianzbedingungen. Ihre Arbeit konzentriert sich auf Modelle metrischer Messinvarianz. In der relevanten Konstellation schreiben sie: “Apply the scaling restriction in one group only”. Eigene Übersetzung: Wenden Sie die Skalierungsrestriktion nur in einer Gruppe an. Diese Anweisung betrifft korrekte Modellidentifikation und den Vergleich von Fitgrößen. Sie macht einen Referenzindikator nicht zum Beweis für Übersetzungsqualität oder skalare Invarianz.

Die Unterscheidung geht in einem Softwareoutput leicht verloren. Ein Modell kann konvergieren, weil seine Skalierungsrestriktionen ausreichend sind. Die Konvergenz beantwortet nicht, ob Ladungen oder Intercepts gleich sind. Die Analyse muss deshalb unterscheiden, welche Restriktionen geprüft, welche zur Identifikation gesetzt und welche freigegeben oder toleriert werden.

Klopp und Klößner zeigen außerdem, warum ein partielles metrisches Modell mit nur einem als invariant angenommenen Indikator je latenter Variable in der relevanten Konstellation dem konfiguralen Modell entsprechen kann. Das ist eine Warnung davor, das Vorhandensein einer Restriktion mit dem Nachweis eines sinnvollen zusätzlichen Vergleichs gleichzusetzen. Freiheitsgrade und Identifikationsentscheidungen sind Teil der Evidenz, keine dekorativen Details.

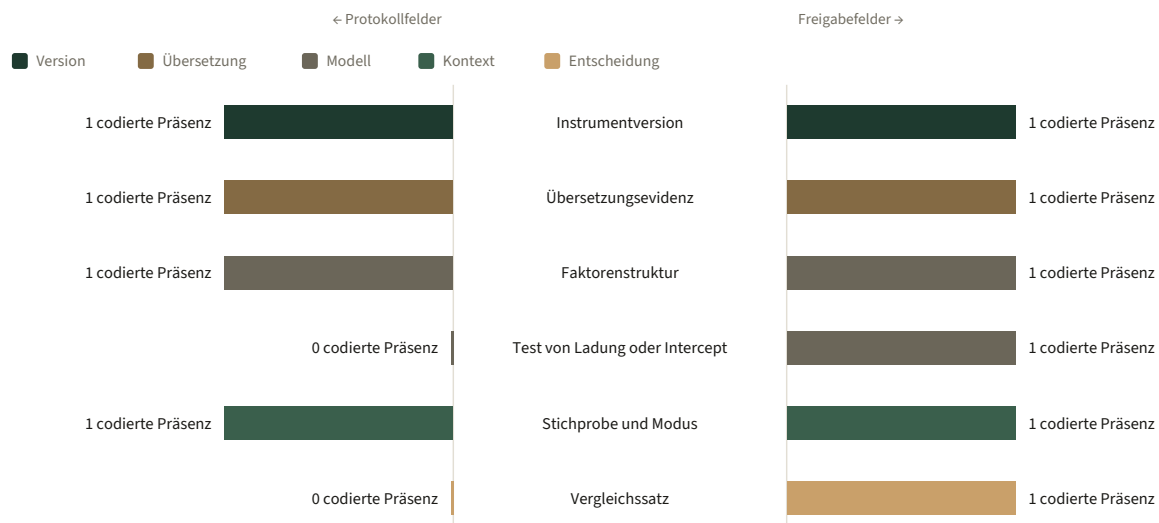
Warum erfordert jede Änderung der Erhebungsplattform eine erneute Invarianzprüfung?

Messinvarianz ist kein dauerhaftes Zertifikat, das an einem Fragebogen hängt. Cieciuch und Kollegen schreiben: “Establishing measurement invariance in one study does not signify that a questionnaire is always measurement invariant.” Eigene Übersetzung: Der Nachweis von Messinvarianz in einer Studie bedeutet nicht, dass ein Fragebogen immer messinvariant ist.

Das Ergebnis gehört zu einer Skala, Übersetzung, Stichprobe, Erhebungsart, einem Faktorenmodell und einem Vergleich. Eine neue Sprachfassung, eine Itemrevision, ein anderer Antwortmodus, eine neue Population oder ein anderer Anwendungskontext kann das Gate erneut öffnen. Ein Webformular ist nicht zwingend derselbe Messkontext wie eine beaufsichtigte Papierbefragung. Eine Beschäftigtenstichprobe ist nicht zwingend dieselbe Population wie Kundinnen, Kunden oder Studierende. Ein verändertes Itembeispiel kann die kognitive Anforderung ändern, obwohl der Konstruktnamen gleich bleibt.

Das praktische Protokoll sollte Instrument und Test versionieren. Halten Sie Itemtext, Übersetzungsentscheidung, Rückübersetzungsnotiz, Stichprobendefinition, Erhebungsmodus, Modellspezifikation, Restriktionen, Fit-Evidenz, freigegebenen Vergleich und offene Einschränkung zusammen. Eine spätere Leserin oder ein späterer Leser sollte erkennen können, zu welcher Version ein Ergebnis gehört.

Abbildung 3 Von der Übersetzungsversion zum Vergleichssatz



Schematischer Rahmen des Autors, begründet mit Cieciuch et al. (2014), Klopp und Klößner (2023) sowie Prem et al. (2021). Die Abfolge ist ein Governance-Modell und behauptet nicht, dass jede Studie genau diese Schritte verwendet.

Wie formulieren Forschungsteams Befunde exakt auf der empirisch nachgewiesenen Stufe?

Das letzte Gate ist sprachlich. Schreiben Sie zuerst, was das Ergebnis erlaubt, und danach, was es nicht erlaubt.

Für konfigurale Evidenz könnte ein vorsichtiger Satz lauten, dass dasselbe Faktorenmuster in den englischen und deutschen Gruppen unter dem spezifizierten Modell geprüft wurde. Bei metrischer Evidenz kann der Satz die relevanten Beziehungen betreffen, wenn Modell und Stichprobe diese Nutzung tragen. Bei skalarer Evidenz kann er sagen, dass für das spezifizierte Instrument, die erklärten Gruppen und das Modell ein Vergleich latenter Mittelwerte unterstützt wurde. Der Satz sollte nicht zu englische und deutsche Befragte sind gleichwertig werden. Diese Formulierung ist breiter als der Test.

Partielle und approximative Ergebnisse brauchen noch mehr Präzision. Nennen Sie, welche Parameter freigegeben oder welche Toleranzen verwendet wurden. Berichten Sie den weiterhin vertretbaren Vergleich und den stärkeren Vergleich, der offen bleibt. Eine kleinere Aussage mit sichtbarer Grenze ist nützlicher als eine universelle Aussage, die nicht reproduziert werden kann.

Das Freigabegate sollte außerdem eine Abbruchregel haben. Stoppen Sie den Vergleich, wenn das geplante Konstrukt nicht stabil ist, die Übersetzungsrevision nicht dokumentiert wurde, Stichprobe oder Modus ohne neuen Test geändert wurden, das relevante Invarianzniveau nicht angesprochen ist oder der Modelloutput Identifikation und inhaltliche Gleichheit nicht auseinanderhält. Noch nicht vergleichbar ist ein gültiges methodisches Ergebnis.

Wo liegen die methodischen Grenzen psychometrischer Messinvarianzprüfungen?

Cieciuch und Kollegen erklären Ebenen sowie die Unterscheidung zwischen exakter und approximativer Invarianz in einer begrenzten Werte-Studie. Klopp und Klößner klären technische Skalierungsfragen in Modellen metrischer Invarianz. Prem und Kollegen liefern ein englisch-deutsches Beispiel mit vier Subskalen und benannten Grenzen von Stichprobe und Instrument. Zusammen validieren die Quellen nicht jede Übersetzung, jedes Instrument, jede Population, jede Website, jede Befragung oder jeden Sprachvergleich.

Der Beitrag zur Lokalisierung behandelt Sprache als Markteintrittsentscheidung. RES-07 behandelt die engere Messfrage, ob ein Wertvergleich durch die Evidenz freigegeben ist. Der Beitrag zur Abbruchregel eines Evidenzreviews liefert eine verwandte Disziplin für die Entscheidung, wann eine Aussage pausieren sollte.

Die belastbare Schlussfolgerung ist bescheiden: Übersetzen Sie sorgfältig, prüfen Sie das Messmodell, benennen Sie den geplanten Vergleich und schreiben Sie nur den Satz, den das Invarianzergebnis trägt. Ein günstiges Ergebnis gehört zu dem Instrument, der Version, der Stichprobe, dem Modus und dem Modell, die es erzeugt haben. Ändert sich eines davon, öffnet sich das Vergleichsgate erneut.

- 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>

ZITIERWEISE

Isoglu, S. (2026, 1. September). Messinvarianz vor dem Vergleich englischer und deutscher Werte. Sinan Isoglu.
<https://isoglu.com/de/insights/journal/messinvarianz-vor-dem-vergleich-englischer-und-deutscher-werte/>

WEITERLESEN

Aus der Forschung

Was ist Selektionsbias? Die Stichprobe kann die Antwort verändern

<https://isoglu.com/de/insights/journal/was-ist-selektionsbias/>

Aus der Forschung

Hybride Intelligenz braucht eine Grenze zwischen Fähigkeit und Ergebnis

<https://isoglu.com/de/insights/journal/hybride-intelligenz-braucht-eine-grenze-zwischen-faehigkeit-und-ergebnis/>

Aus der Forschung

Ein System für immaterielle Werte, zwei nützliche Sichten

<https://isoglu.com/de/insights/journal/ein-system-fuer-immaterielle-werte-zwei-nuetzliche-sichten/>



Sinan Isoglu, MBA (Quantic)

Führungskraft für Umsatzwachstum, Dozent und Doktorand