



BENUTZERHANDBUCH

OpenPLS Benutzerhandbuch

Vollständige Anleitung für Forschende und Praktiker

Version 0.3.28

2026-08-04

openpls.app · Modernes PLS-SEM für Forschende

Inhaltsverzeichnis

1	Was ist PLS-SEM?	4
1.1	Die Bausteine	4
1.2	Die Iteration	4
1.3	Was OpenPLS Ihnen bietet	5
2	Erste Schritte	6
2.1	Die Projekte-Seite	6
2.2	Ein Projekt anlegen	6
2.3	Ein wissenschaftliches Beispiel klonen	7
2.4	Der erste Blick in ein Projekt	8
3	Datensätze	8
3.1	Das leere Datensatz-Panel	9
3.2	Sobald der Datensatz bereit ist	10
3.3	Missing-Codes bearbeiten	10
3.4	Mehrere Datensätze verwalten	11
3.5	Einen Datensatz löschen	11
4	Ein Modell aufbauen	12
4.1	Der visuelle Editor	12
4.2	Die Leinwand	14
4.3	Die Formularansicht	15
4.4	Live-Compute	15
4.5	Versions	16
5	Berechnung und Lesen der Ergebnisse	17
5.1	Der Results-Header	17
5.2	Inner summary (Standard)	18
5.3	Paths	19
5.4	Inner model / outer model	20
5.5	Reliability	20
5.6	HTMT	20
5.7	VIF	21
5.8	Effects	21
5.9	Model fit	21
6	Fortgeschrittene Methoden	22
6.1	Bootstrap	22
6.2	Multi-Group-Analysis (MGA)	24
6.3	MICOM	24
6.4	Moderation	27

6.5	Nomologische Validität	27
6.6	PLSc	27
6.7	Gaussian Copula (Endogenität)	30
6.8	Common Method Bias (CMB, Lindell-Whitney)	31
6.9	IPMA	31
6.10	PLSpredict	34
6.11	CVPAT	34
6.12	FIMIX-PLS	34
6.13	Higher-Order Construct (HOC)	38
7	Exporte	40
7.1	PDF-Report	40
7.2	Excel-Arbeitsmappe (XLSX)	40
7.3	Report-JSON-Bundle	41
7.4	Modell-Export (Editor-Toolbar)	41
7.5	LaTeX-Snippets	41
7.6	Reproduzierbarkeit	41
8	Zusammenarbeit	42
8.1	Das Team-Panel	42
8.2	Rollen	42
8.3	Eine Einladung senden	43
8.4	Das Activity-Log	44
8.5	Ein Projekt verlassen	44
8.6	Was Mitwirkende sehen	45
9	Einstellungen und Sprache	45
9.1	Projekteinstellungen	45
9.2	Sprache	46
9.3	Kontoeinstellungen	47
9.4	Theme	49
10	Hilfe und weiterführende Literatur	49
10.1	In-App-Hilfe	49
10.2	Die OpenPLS-Engine (Open Source)	49
10.3	Support, Bugs und Feature-Wünsche	50
10.4	OpenPLS zitieren	50
11	Der AI Companion	51
11.1	Aktivieren	51
11.2	AI hints: der stets sichtbare Streifen	52
11.3	Chat: das schwebende Drawer	52
11.4	Ein Modell aus einer Forschungsfrage entwerfen	55

11.5 Einen Fragebogen entwerfen	55
11.6 Method-Selection-Assistant	58
11.7 Report-Drafting	61
11.7.1 Methods / Results / Discussion	61
11.7.2 Reviewer defense	62
11.8 Die Publikationsbibliothek	65
11.9 Kontingent und Datenschutz	65
11.10 Ausschalten	66
12 Öffentliche Umfragen	66
12.1 Den Fragebogen bauen	66
12.2 Veröffentlichen	68
12.2.1 Versions-Snapshots	68
12.3 Die öffentliche Respondentenansicht	70
12.4 Antworten beobachten	70
12.5 Antworten als Datensatz importieren	70
12.6 Rohe CSV herunterladen	73
12.7 Veröffentlichung zurücknehmen	73
12.8 Limits und Konventionen	73
13 Anhang	74
13.1 Glossar	74
13.2 Tastaturkürzel	76
13.3 Referenzen	76

1 Was ist PLS-SEM?

Partial Least Squares Structural Equation Modeling (PLS-SEM) ist eine Methode zur Schätzung von Modellen mit **latenten Konstrukten**: Dingen, die wir nicht direkt messen können, sondern nur über beobachtbare Indikatoren. Kundenzufriedenheit, Markenpersönlichkeit und organisationales Vertrauen sind allesamt latent: Niemand liefert Ihnen einen “Zufriedenheits”-Messwert; Sie leiten ihn aus Fragebogenitems ab, die ihn stellvertretend abbilden.

PLS-SEM beantwortet in einer einzigen Schätzung drei verknüpfte Fragen:

1. **Messung.** Wie gut repräsentieren die beobachtbaren Indikatoren (Fragebogenitems, Bewertungen, Sensormesswerte) jedes latente Konstrukt?
2. **Struktur.** Wie stark beeinflusst jedes Konstrukt die anderen? Treibt wahrgenommene Qualität die Zufriedenheit an? Vermittelt Vertrauen den Effekt von Nützlichkeit auf Adoption?
3. **Prognose.** Wie gut generalisiert das Modell? Wie würde das Ergebnis für eine neue Beobachtung aussehen?

Anders als kovarianzbasierte SEM (LISREL, Mplus, lavaan) verlangt PLS-SEM keine multivariante Normalverteilung der Daten, funktioniert mit kleinen Stichproben und verarbeitet **formative** Messungen (Indikatoren, die ein Konstrukt gemeinsam definieren, z. B. Einkommen + Bildung + Beruf → sozioökonomischer Status) problemlos. Dadurch ist es die Standardwahl in Wirtschaftsforschung, Marketing, Wirtschaftsinformatik und zunehmend auch in Gesundheits- und Sozialwissenschaften.

1.1 Die Bausteine

Ein PLS-SEM-Modell besteht aus vier Arten von Bestandteilen:

- **Latente Variablen (LVs):** die Konstrukte, um die es Ihnen geht. Als Kreise dargestellt.
- **Indikatoren / manifeste Variablen (MVs):** die beobachteten Spalten Ihres Datensatzes, die jede LV messen. In klassischen Pfaddiagrammen als Quadrate dargestellt; in OpenPLS liegen sie in der Seitenleiste und werden LVs zugewiesen.
- **Messmodell / äußeres Modell:** die Pfeile zwischen einer LV und ihren Indikatoren. Die Richtung unterscheidet **reflektive** (Mode A: LV verursacht die Indikatoren, z. B. Zufriedenheit → sat1..sat4) von **formativen** Messungen (Mode B: Indikatoren verursachen die LV, z. B. Einkommen + Bildung → SES).
- **Strukturmodell / inneres Modell:** die Pfeile zwischen LVs. Diese sind die zu prüfenden Hypothesen.

1.2 Die Iteration

PLS-SEM alterniert zwei Kleinst-Quadrate-Regressionen bis zur Konvergenz:

1. Schätzung jeder LV als gewichtetes Composite ihrer Indikatoren unter Verwendung der aktuellen Innenmodell-Gewichte.

2. Schätzung der Pfadkoeffizienten des Innenmodells durch Regression jeder endogenen LV auf ihre Vorgänger.

Die Konvergenz ist schnell (üblicherweise < 10 Iterationen). Standardfehler und p-Werte stammen aus dem **Bootstrapping**: mehrtausendfaches Ziehen mit Zurücklegen und Neuschätzung.

1.3 Was OpenPLS Ihnen bietet

OpenPLS ist ein modernes, offenes, cloud-natives Werkzeug für PLS-SEM. In einem einzigen Projekt können Sie:

- Einen Datensatz hochladen (CSV, XLSX, SPSS .sav, Stata .dta), oder **einen Fragebogen erstellen, ihn als öffentliche Umfrage veröffentlichen und die gesammelten Antworten mit einem Klick als Datensatz importieren** (Kapitel 13).
- Das Modell visuell zeichnen (per Drag von LV zu LV), in einem Formular eingeben, oder **mit dem AI Companion aus einer Forschungsfrage ein Startmodell generieren** (Kapitel 12).
- Die PLS-Schätzung berechnen, mit Live-Neuberechnung bei jeder Änderung.
- Die Messqualität beurteilen: Reliabilitäten, AVE, HTMT, VIF.
- Die strukturelle Qualität beurteilen: R^2 , adjustiertes R^2 , f^2 , Q^2 , SRMR, dULS, Fornell-Larcker.
- Pfade per Bootstrap für p-Werte und Konfidenzintervalle absichern.
- Die fortgeschrittenen Methoden ausführen: MGA, MICOM, Moderation, Mediation, IP-MA, PLSpredict, FIMIX, HOC, PLSc, Gaussian Copula, CVPAT, Nomologische Validität, Common-Method-Bias.
- **Kontextuelle KI-Hinweise auf jedem Tab erhalten, mit einem Assistenten chatten, der nur Lesezugriff auf Ihre Projektdaten hat, und veröffentlichungsreife Methoden- / Ergebnis- / Diskussionsabschnitte sowie antizipierte Verteidigungen gegenüber Reviewern verfassen** — alles verankert in einer kuratierten Bibliothek von Open-Access-PLS-SEM-Publikationen (Kapitel 12).
- Den gesamten Lauf als PDF, XLSX, LaTeX oder als in sich geschlossenes JSON-Bundle für Reproduzierbarkeit exportieren.
- Mitwirkende einladen, Aktivitäten nachverfolgen und einen Audit-Trail jeder Modellversion pflegen.

Der Rest dieses Leitfadens geht jeden dieser Punkte in der Reihenfolge durch, in der Sie sie tatsächlich verwenden würden.

Dieser Leitfaden bezieht sich auf OpenPLS Version 0.3.27, die unter `cloud.openpls.app` läuft. Bildschirmansichten können in späteren Releases leicht abweichen; die Abläufe bleiben identisch.

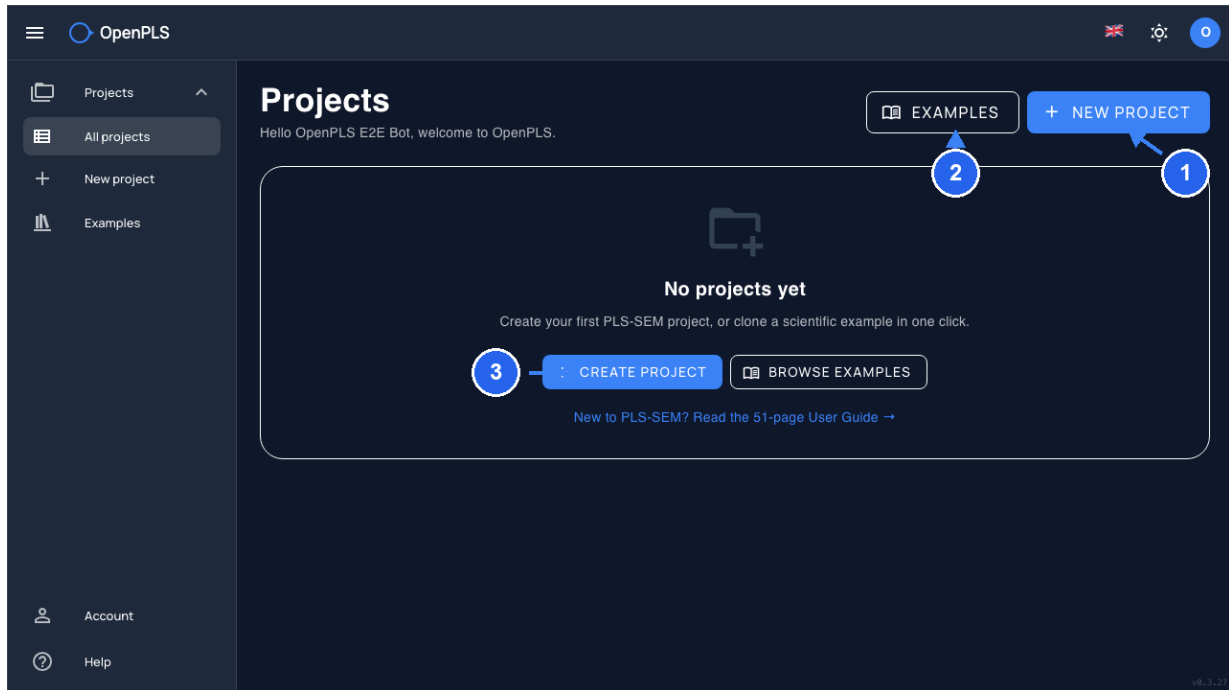


Abbildung 1 Projekte-Landingpage: der Header bietet **New project** und **Examples**; die Empty-State-Karte in der Mitte bietet dieselben beiden Aktionen.

2 Erste Schritte

Sie benötigen ein Konto, um Projekte zu speichern. Rufen Sie <https://cloud.openpls.app> auf, registrieren Sie sich mit einer E-Mail-Adresse, verifizieren Sie diese und melden Sie sich an. Sie landen auf der Seite **Projekte**, Ihrem Arbeitsbereich.

2.1 Die Projekte-Seite

Die Projekte-Seite (Abbildung 1) ist Ihr Startbildschirm. Sie listet jedes Projekt auf, auf das Sie Zugriff haben, und weist Sie im Header auf zwei Einstiegsaktionen hin, plus ein duplizierendes Paar innerhalb der Empty-State-Karte.

Drei Möglichkeiten zum Einstieg (nummeriert in Abbildung 1):

1. **New project** (oben rechts), ein leerer Arbeitsbereich. Sie laden Ihren eigenen Datensatz hoch und bauen das Modell von Grund auf neu.
2. **Examples**: die Beispielfall-Bibliothek. Ein Klick kloniert einen wissenschaftlichen PLS-SEM-Klassiker (ECSI, MOBI, Russett, mehrere synthetische Datensätze) in Ihren Arbeitsbereich, Datensatz + Modell vorverdrahtet.
3. **Create project** in der Empty-State-Karte, ein Shortcut zum selben Dialog wie (1). Verwenden Sie diesen oder den Header-Button, je nachdem, was näher an Ihrem Cursor ist.

2.2 Ein Projekt anlegen

Ein Klick auf eine der beiden **New project**-Aktionen öffnet den in Abbildung 2 gezeigten Dialog.

Abbildung 2 Der **New project**-Dialog mit Namensfeld, optionaler Beschreibung und der **Create**-Aktion unten rechts.

Ausfüllen (nummeriert in Abbildung 2):

1. **Projektname:** erforderlich, ≥ 2 Zeichen. Alles von “Kundenzufriedenheit 2026” bis “H1-Experiment, Länderkohorte”.
2. **Beschreibung:** optional, bis zu 500 Zeichen. Hier notieren Sie die Forschungsfrage oder die zu prüfende Hypothese. Nützlich, wenn Sie sechs Monate später wieder darauf zurückkommen.
3. Klicken Sie auf **Create**. Sie landen auf der frischen Projektdetailseite. Da noch kein Datensatz angehängt ist, öffnet die App direkt den Dataset-Tab, damit Sie sofort einen hochladen können.

2.3 Ein wissenschaftliches Beispiel klonen

Wenn Sie neu bei PLS-SEM sind, starten Sie mit einem bekannt guten Modell. Klicken Sie im Header auf **Examples**, um die Bibliothek zu öffnen (Abbildung 3).

1. Wählen Sie eine Karte, die zu Ihrer Domäne passt (Kundenzufriedenheit, Markenerleben, Technologieakzeptanz, Arbeitsengagement, Gesundheits-Apps, Politikwissenschaft...). Jede wird mit einem vollständig spezifizierten Datensatz + Modell ausgeliefert.
2. Klicken Sie auf **Clone into my workspace**. OpenPLS legt eine private Kopie unter Ihrem Konto an. Sie können Indikatoren beliebig löschen, LVs hinzufügen, Pfade ändern, das Original bleibt für andere Nutzer unverändert.

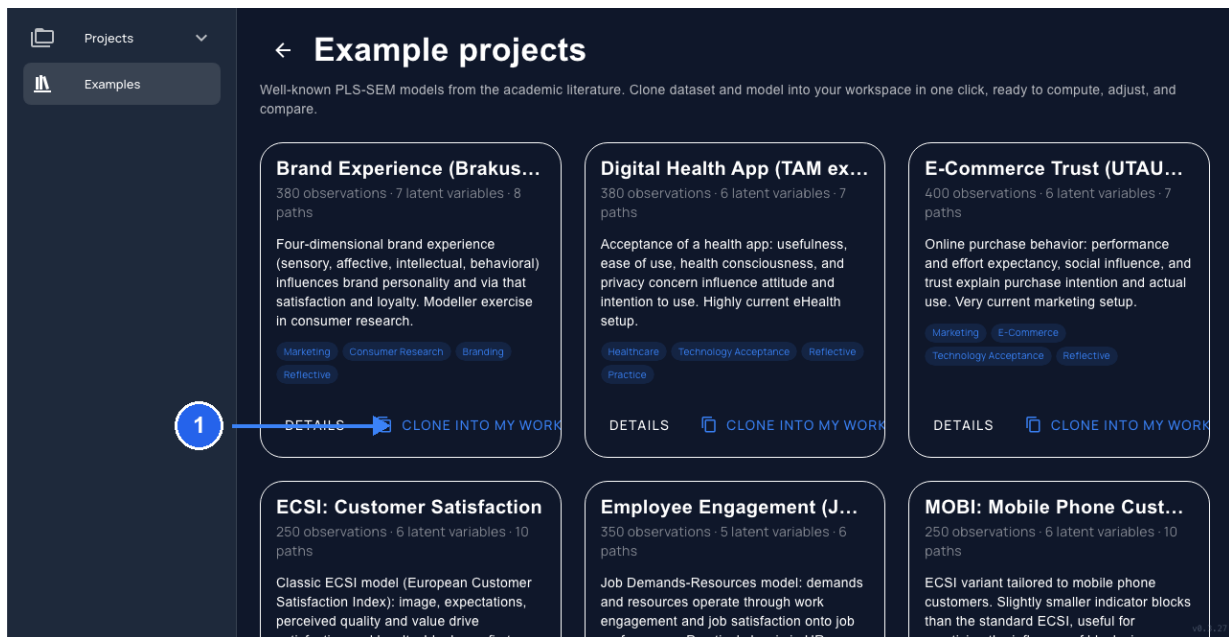


Abbildung 3 Die Examples-Bibliothek. Jede Karte zeigt die Datensatz-Kennzahlen und die Zitation; klicken Sie auf eine Karte, um sie zu öffnen, und dann auf **Clone into my workspace**, um eine bearbeitbare Kopie anzulegen.

Die Beispiele sind Open-Licence-Lehrfälle aus der PLS-SEM-Literatur (*Tenenhaus et al. 2005, Brakus & Schmitt 2009, Bakker & Demerouti 2017, Russett 1964, Venkatesh et al. 2012, Davis 1989*). Vollständige Zitationen finden Sie auf jeder Karte. Synthetische Datensätze sind klar gekennzeichnet und wurden ausschließlich zu Lehrzwecken gegen die publizierte Pfadstruktur simuliert.

2.4 Der erste Blick in ein Projekt

Sobald Sie sich in einem Projekt befinden, sehen Sie oben die Tab-Leiste:

- **Dataset:** hochladen, inspizieren, Missing-Code-Konventionen bearbeiten
- **Model:** das Modell aufbauen (Editor / Form / Versions)
- **Results:** die Compute-Ausgabe lesen
- **Advanced:** die 13 fortgeschrittenen Methoden
- **Manage:** Team, Aktivitätslog, Projekteinstellungen

Wir gehen jeden Tab der Reihe nach durch. Als Nächstes: einen Datensatz hochladen.

3 Datensätze

Jedes Projekt benötigt genau einen Datensatz, bevor Sie ein Modell bauen können. OpenPLS liest **CSV**, **XLSX**, **SPSS .sav** und **Stata .dta**, bis zu 50 MB pro Datei.

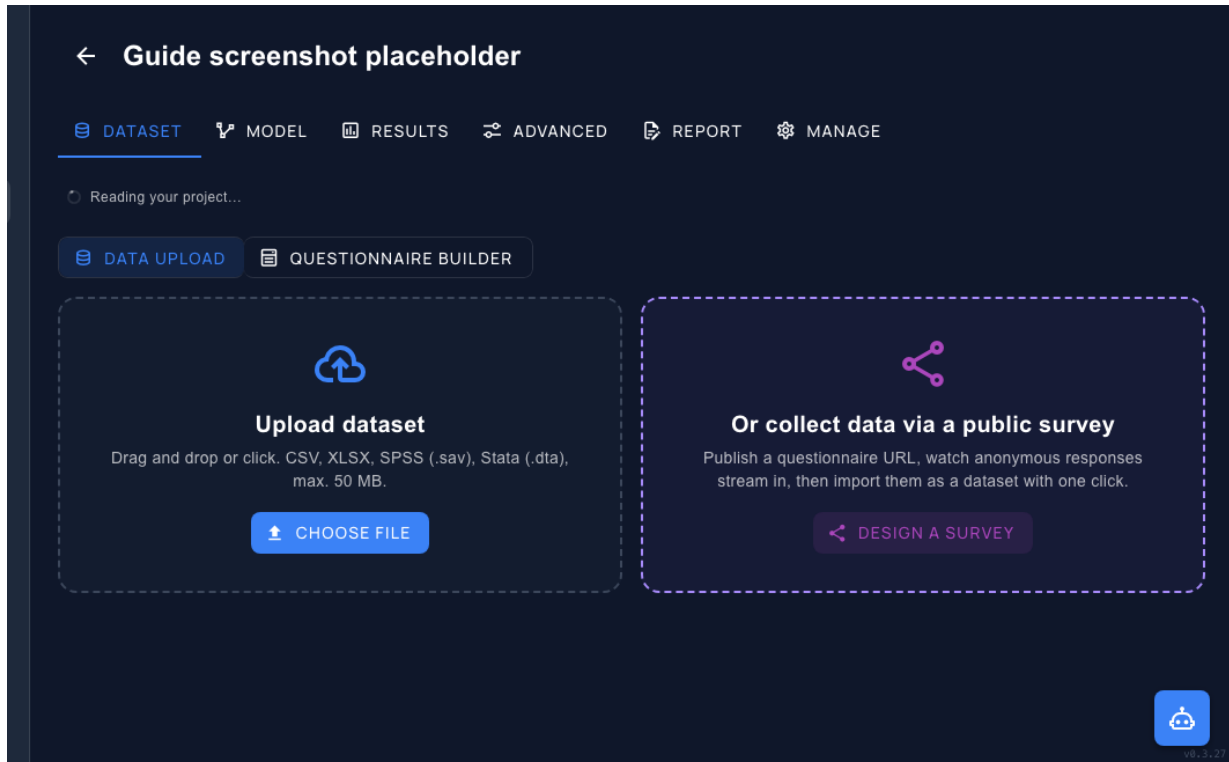


Abbildung 4 Leerer Dataset-Tab. Die Tab-Leiste am oberen Rand wechselt zwischen Dataset, Model, Results, Advanced, Report und Manage. Darunter ein Unter-Toggle, der zwischen dem Hochladen einer Datei und dem Erstellen einer öffentlichen Umfrage wählt.

3.1 Das leere Datensatz-Panel

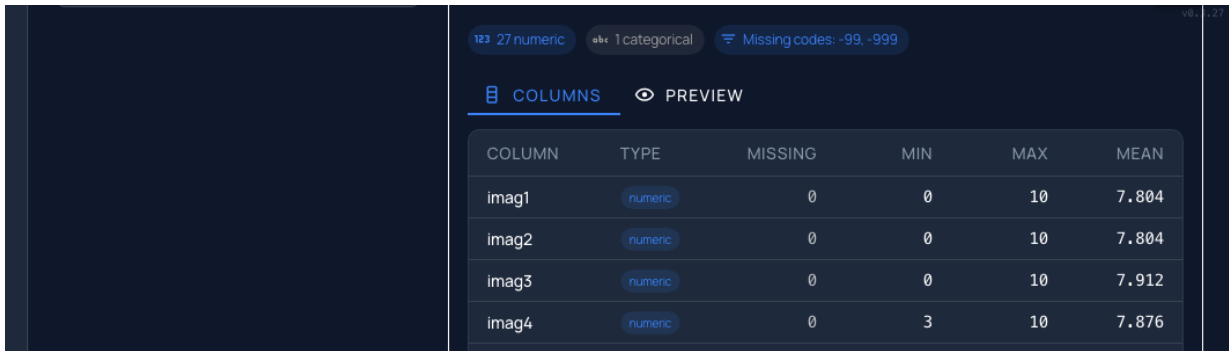
Ein frisches Projekt öffnet sich mit aktivem Dataset-Tab. Wenn noch keine Daten angehängt sind, sehen Sie die Upload-Zone in Abbildung 4.

Zwei Schritte, um Daten hineinzubekommen:

1. Bestätigen Sie, dass Sie auf dem **Dataset**-Tab sind. Er ist in jedem Projekt der erste Tab.
2. Ziehen Sie eine Datei auf die gestrichelte Drop-Zone oder klicken Sie darauf, um den Dateipicker Ihres Betriebssystems zu öffnen.

Für CSV-Uploads öffnet OpenPLS anschließend einen **CSV-separator**- Dialog. Er zeigt eine Vorschau der ersten Zeilen und schlägt den Separator automatisch vor (, , ; , Tab oder |); bestätigen Sie ihn, um fortzufahren.

CSVs aus europäischen Excel-Exporten verwenden üblicherweise ';' als Separator. OpenPLS erkennt dies automatisch, prüfen Sie jedoch die Vorschau, bevor Sie auf Upload klicken; ein falscher Separator ist der schnellste Weg, einen unsinnigen Datensatz zu erzeugen.



The screenshot shows the 'PREVIEW' tab of the dataset panel. At the top, there are three chips: '123 27 numeric', 'abs 1 categorical', and 'Missing codes: -99, -999'. Below these are two tabs: 'COLUMNS' and 'PREVIEW'. The 'PREVIEW' tab is active, displaying a table with the following data:

COLUMN	TYPE	MISSING	MIN	MAX	MEAN
imag1	numeric	0	0	10	7.804
imag2	numeric	0	0	10	7.804
imag3	numeric	0	0	10	7.912
imag4	numeric	0	3	10	7.876

Abbildung 5 Dataset-Panel nach dem Upload. Die Datensatzkarte zeigt den Ready-Chip und die aktuellen Missing-Code-Sentinels; das Hauptfeld rechts zeigt die Statistiken pro Spalte.

3.2 Sobald der Datensatz bereit ist

Das Parsen erfolgt serverseitig und dauert für übliche Dateien ein bis zwei Sekunden. Wenn es abgeschlossen ist, erscheint die Datensatzkarte in der linken Spalte und ihre Details werden rechts dargestellt (Abbildung 5). Zwei Chips sind wichtig:

1. **Status-Chip:** wird grün und zeigt "Ready" an, sobald das Parsen abgeschlossen ist. Bleibt er rot ("Error"), klicken Sie auf den Datensatz, um die Fehlermeldung zu sehen.
2. **Missing-codes-Chip:** zeigt, welche numerischen Werte OpenPLS als fehlend behandelt. Die Voreinstellung -99, -999 entspricht der SPSS-Konvention. Klicken Sie auf den Chip, um ihn zu bearbeiten.

Das rechte Feld listet jede Spalte mit ihrem Typ (numeric, string, mixed), der Anzahl fehlender Zellen und (für numerische Spalten) Min, Max und Mittelwert. Verwenden Sie dies als Plausibilitätsprüfung des Parsens: Eine Spalte, die Sie als numerisch erwartet haben und die als string zurückkommt, deutet meist auf eine versprengte nicht-numerische Zelle hin.

3.3 Missing-Codes bearbeiten

Manche Datensätze kodieren fehlende Werte als leere Zellen, andere verwenden -99 (SPSS-Standard), wieder andere 9999 oder -1. Dies falsch zu haben ist der **häufigste einzelne Grund** für einen "warum ist mein R^2 so niedrig?"-Bugreport, das Modell behandelt Sentinelwerte als echte Daten und ruiniert die Korrelationen.

Klicken Sie auf den Missing-Codes-Chip, um den in Abbildung 6 gezeigten Editor zu öffnen.

Die drei Elemente in Abbildung 6:

1. **Existierende Sentinels:** jeder ist ein Chip. Klicken Sie auf das × eines Chips, um ihn aus der Menge zu entfernen.
2. **Code hinzufügen:** geben Sie eine beliebige ganze Zahl ein (positiv oder negativ, gemäß der Konvention Ihres Datensatzes) und klicken Sie auf **Add**. Sie erscheint als neuer Chip.
3. **Apply and re-parse:** sendet die neue Menge an den Server, der die Datei mit den aktualisierten Sentinels erneut einliest und die Spaltenstatistiken aktualisiert. Warten Sie, bis der

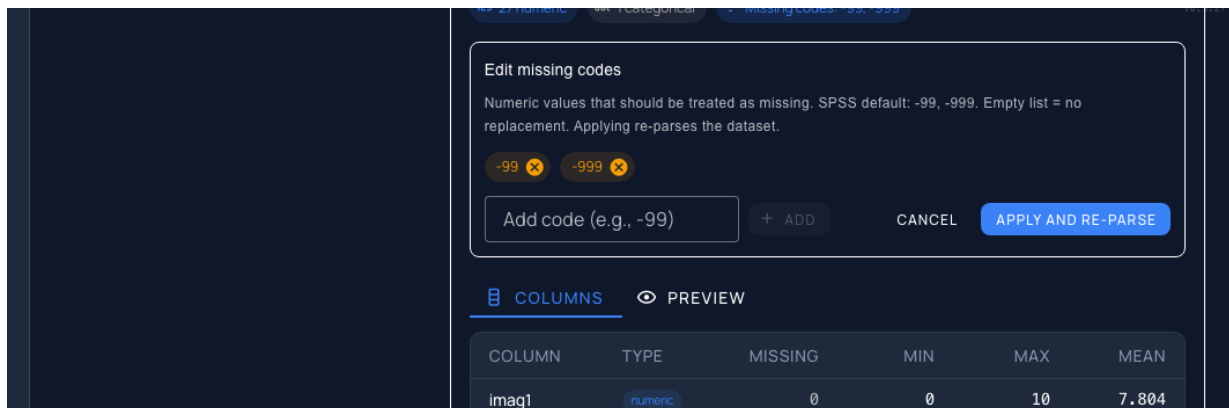


Abbildung 6 Missing-Codes-Editor. Draft-Chips sind über das × entfernbar; neue Sentinels werden über das Eingabefeld hinzugefügt; **Apply and re-parse** übernimmt die Änderung und liest die Datei erneut ein.

Spinner des Buttons aufhört, bevor Sie weitermachen.

Missing-Codes anzuwenden parst die Datei von Grund auf neu. Wenn Sie bereits ein Modell berechnet hatten, wird die Berechnung ungültig, starten Sie sie nach dem Neu-Parsen erneut.

3.4 Mehrere Datensätze verwalten

Ein einzelnes Projekt kann mehrere Datensätze parallel halten (z. B. einen pro Land oder eine Pilot- vs. Vollstichprobe). Fügen Sie weitere hinzu, indem Sie zusätzliche Dateien per Drag oder Datei-Picker auswählen. Jeder Datensatz erhält seine eigene Karte in der linken Spalte; klicken Sie auf eine Karte, um sie zu aktivieren. Modelle werden an einen bestimmten Datensatz gebunden, sodass ein Wechsel des aktiven Datensatzes über die Kartenliste nur ändert, was der Dataset-Tab anzeigt, das Modell verweist weiter auf den Datensatz, mit dem es erstellt wurde.

3.5 Einen Datensatz löschen

Um einen Datensatz zu löschen, bewegen Sie den Mauszeiger über seine Karte in der linken Spalte und klicken Sie auf das Papierkorb-Symbol. Die Löschung erfolgt sofort und kann nicht rückgängig gemacht werden. Wenn irgendein Modell im Projekt noch auf den gelöschten Datensatz verweist, wird seine Berechnung mit Fehler enden, bis Sie einen anderen Datensatz zuweisen.

Datensätze werden sofort gelöscht und können nicht wiederhergestellt werden. Jedes Modell im Projekt, das noch auf den gelöschten Datensatz verweist, wird bei der Berechnung Fehler werfen, verknüpfen Sie entweder einen anderen Datensatz im Model-Tab oder löschen Sie auch das Modell.

4 Ein Modell aufbauen

OpenPLS bietet Ihnen drei Wege, das Modell zu formen:

- **Editor:** die visuelle Drag-and-Drop-Leinwand. Dies ist der primäre Workflow und das Killer-Feature von OpenPLS.
- **Form:** ein einfaches Formular mit einer Tabelle für LVs und einer Tabelle für Pfade. Schneller, wenn Sie das Modell auswendig kennen oder aus einer Publikation übertragen.
- **Versions:** das Audit-Log jeder vergangenen Berechnung. Jede Berechnung erzeugt einen Snapshot von Modell + Ergebnis, sodass Sie jeden früheren Zustand wiederherstellen können.

Alle drei Unteransichten bearbeiten dasselbe zugrundeliegende Modell. Sie können frei zwischen ihnen wechseln.

4.1 Der visuelle Editor

Klicken Sie auf **Model** → **Editor**, um die in Abbildung 7 gezeigte Leinwand zu öffnen. Die Toolbar sitzt über der Leinwand; die Indikator-Seitenleiste rechts listet jede numerische Spalte des Datensatzes.

Von links nach rechts gelesen zeigt die Toolbar in Abbildung 7:

1. **Modellauswahl:** Ein Projekt kann mehrere Modelle für denselben Datensatz halten. Wählen Sie eines aus dem Dropdown.
2. **New model:** erstellt ein frisches, leeres Modell. Der Dialog fragt nach einem Namen; Sie können später in den Settings umbenennen.
3. **Dataset-Picker:** Jedes Modell wird an einen Datensatz gebunden. Wenn Sie mehr als eine CSV hochgeladen haben, wählen Sie hier aus.
4. **Undo / Redo:** Cmd+Z / Cmd+Shift+Z unter macOS, Ctrl+Z / Ctrl+Y unter Windows und Linux. Undo geht durch das Anlegen von LVs, Indikator-Zuweisungen und das Zeichnen von Pfaden zurück.
5. **Add LV:** setzt einen neuen Latente-Variable-Kreis an einer zufälligen Position auf der Leinwand ab und wählt ihn automatisch aus.
6. **Auto layout:** führt dagre aus, um Knoten in topologischer Reihenfolge von links nach rechts anzuordnen. Günstiger Weg, aus einem Wirrwarr überlappender Kreise etwas Veröffentlichungsreifes zu machen.
7. **Fit to window:** zentriert die Leinwand erneut auf alle sichtbaren Knoten.
8. **Export model:** lädt das Modell als `.openpls.json` herunter (vollständiges OpenPLS-Format, inklusive Positionen) oder als `.splsm` (SmartPLS-kompatibles XML).
9. **Live-Toggle:** aktivieren, und OpenPLS berechnet das Modell 500 ms nach jeder Änderung neu. Pfadkoeffizienten werden auf den Kanten dargestellt, R^2 erscheint innerhalb jeder endogenen LV. Schnelles Feedback beim Iterieren.
10. **Save:** schreibt das aktuelle Modell in die Datenbank. Es gibt kein Autosave; klicken Sie auf Save (oder Compute, das implizit speichert), um Ihre Änderungen zu persistieren.

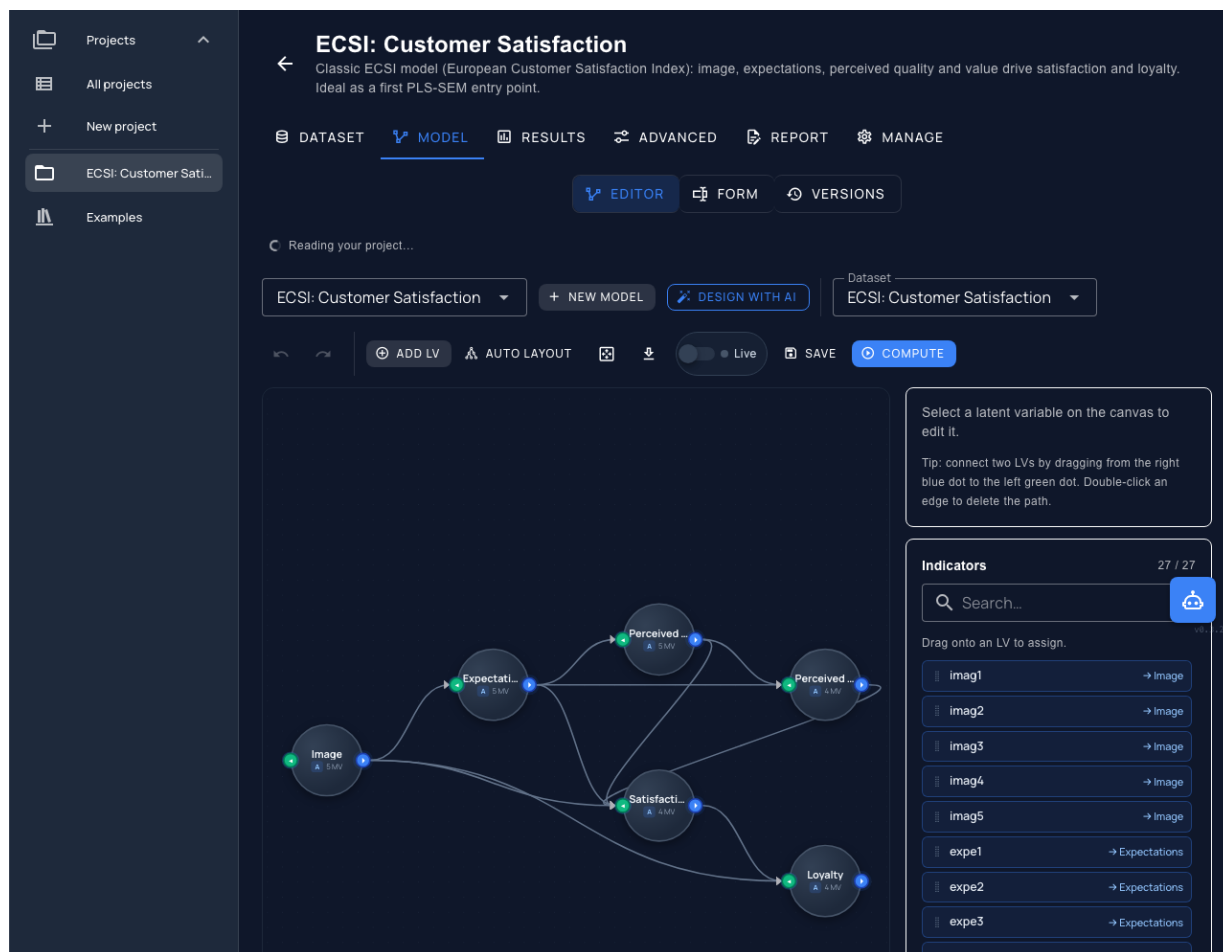


Abbildung 7 Der visuelle Editor mit der Toolbar (oben), der Vue-Flow-Leinwand (Mitte) und der Indikator-Seitenleiste (rechts).

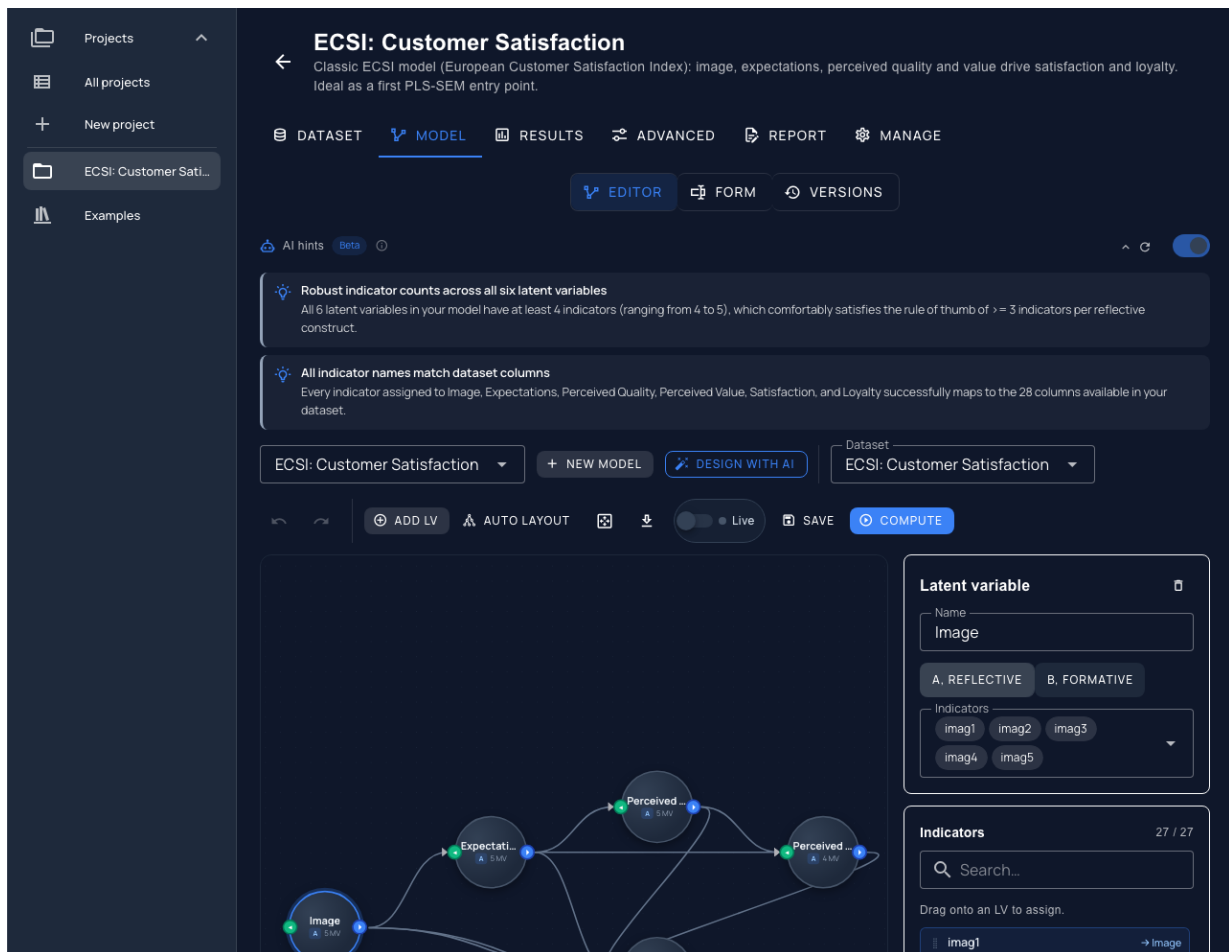


Abbildung 8 LV ausgewählt. Die Karte auf der rechten Seite ermöglicht Umbenennung, Änderung des Messmodus und Auswahl der Indikatoren.

11. **Compute:** führt die vollständige Schätzung serverseitig aus (Bootstrap + Metriken), schreibt ein Result-Dokument und wechselt zum Results-Tab.

4.2 Die Leinwand

Jede **latente Variable** wird als Kreis dargestellt mit:

- dem Namen der LV
- einer Pille, die den Messmodus anzeigt: **A** (reflektiv) oder **B** (formativ)
- einer Zählung zugewiesener Indikatoren, z. B. "5 MV"
- sobald Ergebnisse existieren, dem R^2 -Wert innerhalb des Kreises

Jeder **Pfad** wird als Pfeil zwischen zwei LVs dargestellt. Auf der Quell-LV ist der kleine blaue Punkt rechts das **Source-Handle**; auf der Ziel-LV ist der kleine grüne Punkt links das **Target-Handle**. Ziehen Sie von einem Quellpunkt zu einem Zielpunkt, um einen neuen Pfad zu zeichnen. Doppelklicken Sie einen bestehenden Pfad, um ihn zu löschen.

Ein Rechtsklick oder das Auswählen einer LV öffnet den Editor rechts (Abbildung 8).

1. **Name:** freier Text. Umbenennungen werden auf jeden Pfad, der auf diese LV verweist, propagiert.
2. **Mode A / Mode B:** reflektiv vs. formativ. Reflektiv bedeutet, dass die LV ihre Indikatoren verursacht (ändern Sie die Zufriedenheit, bewegen sich alle sat-Items gemeinsam); formativ bedeutet, dass die Indikatoren die LV definieren (Einkommen, Bildung, Beruf definieren gemeinsam SES). Wählen Sie falsch, und Reliabilitätsmetriken werden entweder bedeutungslos (Mode B mit reflektiver Interpretation) oder konvergieren nicht (Mode A mit einem formativen Konstrukt).
3. **Indicators:** eine Multi-Auswahl über jede numerische Spalte des Datensatzes. Ein hier hinzugefügter Indikator erscheint sofort als Pfeil innerhalb des LV-Kreises. Alternativ können Sie Chips aus der rechten **Indicators**-Seitenleiste direkt auf den LV-Kreis ziehen, OpenPLS legt den Indikator beim Loslassen ab.
4. **Delete LV:** das kleine Papierkorb-Symbol im Header. Entfernt die LV und alle Pfade, die auf sie zeigen oder von ihr ausgehen.

Die rechte **Indicators**-Seitenleiste listet jede verbleibende numerische Spalte auf, mit einem Badge, das anzeigt, welche LV (falls überhaupt) sie bereits verwendet. Ein einzelner Indikator kann in mehreren LVs auftauchen, aber das deutet meistens auf einen Modellierungsfehler hin, untersuchen Sie das, bevor Sie weitermachen.

4.3 Die Formularansicht

Nicht jeder Nutzer möchte Kreise klicken. Wechseln Sie zu **Form**, um den tabellenbasierten Editor zu sehen (Abbildung 9).

Das Formular ist ein einfaches HTML-Formular. Fügen Sie LVs mit dem Button unten in der LV-Tabelle hinzu; wählen Sie Indikatoren aus einer Multi-Auswahl in jeder Zeile. Fügen Sie Pfade mit dem Button unter der Pfadtabelle hinzu; Quelle und Ziel stammen aus dem LV-Dropdown. Alles, was Sie im Editor tun können, können Sie auch hier tun.

Wann Sie das Formular verwenden sollten: wenn Sie ein Modell aus einer Publikation per Copy-and-Paste übernehmen, als Tastaturnutzer oder bei automatisierter Modellkonstruktion.

4.4 Live-Compute

Der **Live**-Toggle in der Toolbar aktiviert inkrementelle Neuberechnung. In dem Moment, in dem Sie irgendetwas ändern (eine LV hinzufügen, einen Pfad verbinden, einen Indikator ankreuzen, umbenennen), entprellt OpenPLS für 500 ms und führt dann serverseitig eine vollständige PLS-Schätzung aus. Das Ergebnis fließt zurück und zeichnet:

- Pfadkoeffizienten auf jede Kante
- R^2 und R^2 adjustiert innerhalb jeder endogenen LV
- p-Werte auf den Pfaden (aus einem Bootstrap, der beim letzten vollständigen Compute

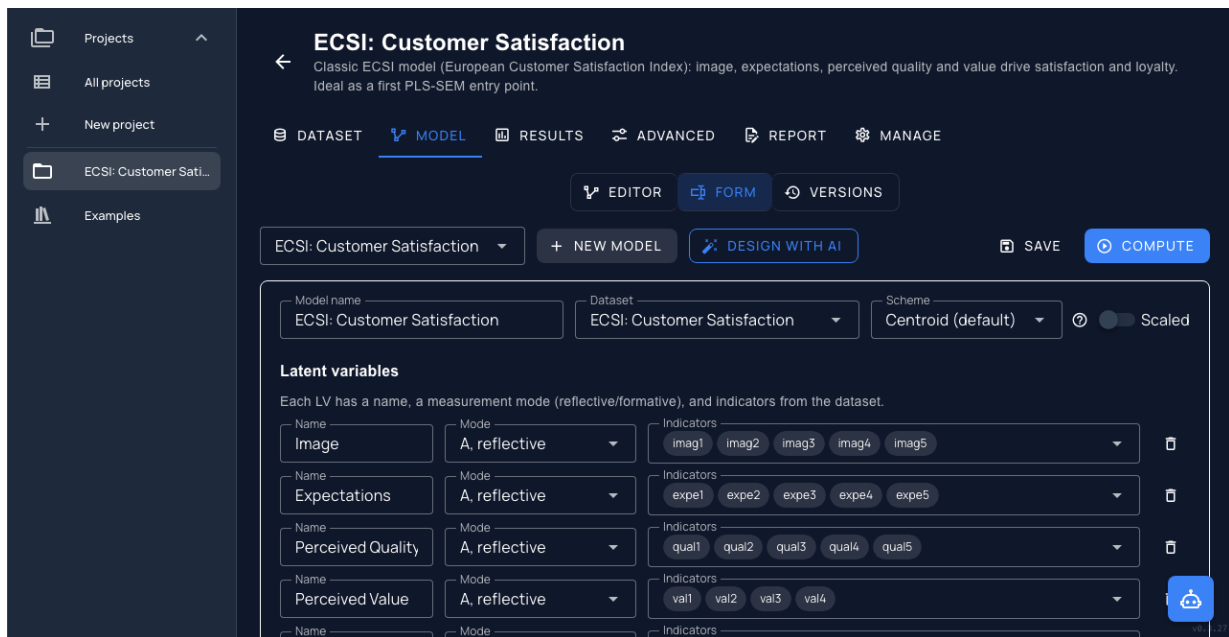


Abbildung 9 Die Form-Unteransicht: LV-Tabelle oben (Name, Modus, Indikatoren), Pfadtabelle darunter, Save + Compute oben rechts.

läuft, nicht bei jedem Tastendruck neu gebootstrapped)

Der Live-Chip in der Toolbar wechselt zwischen **computing...**, **waiting...** und **error**, damit Sie den Zustand des Hintergrundjobs kennen.

Live-Compute ist der schnellste Weg zu erkunden, was passiert, wenn Sie einen neuen Pfad hinzufügen oder einen Indikator neu zuweisen. Schalten Sie es bei großen Modellen aus, wenn Sie viele Änderungen machen möchten, bevor Sie manuell auf Compute drücken.

4.5 Versions

Jedes Mal, wenn Sie auf **Compute** klicken, erstellt OpenPLS einen Snapshot des gesamten Modells + Ergebnisses als versionierten Eintrag. Wechseln Sie zur **Versions**-Unteransicht (Abbildung 10), um sie zu durchsuchen.

Jede Versionskarte zeigt:

- **v#**: fortlaufende Versionsnummer
- das Modell, zu dem sie gehört (ein Chip)
- **GoF**: die Goodness-of-Fit für das Ergebnis dieser Version, falls erfolgreich
- den Avatar des Erstellers und eine Statistikzeile (LV-Anzahl, Pfadanzahl)
- **Rename** (Bleistift), Beschriften Sie die Version zur späteren Bezugnahme, z. B. "vor Invarianztest" oder "finale Einreichung"
- **Delete**: entfernt die Version. Die zugrundeliegenden Modell- + Result-Dokumente werden ebenfalls gelöscht.

Eine Version wiederherstellen ersetzt den aktuellen Live-Zustand des Modells durch den Snapshot.

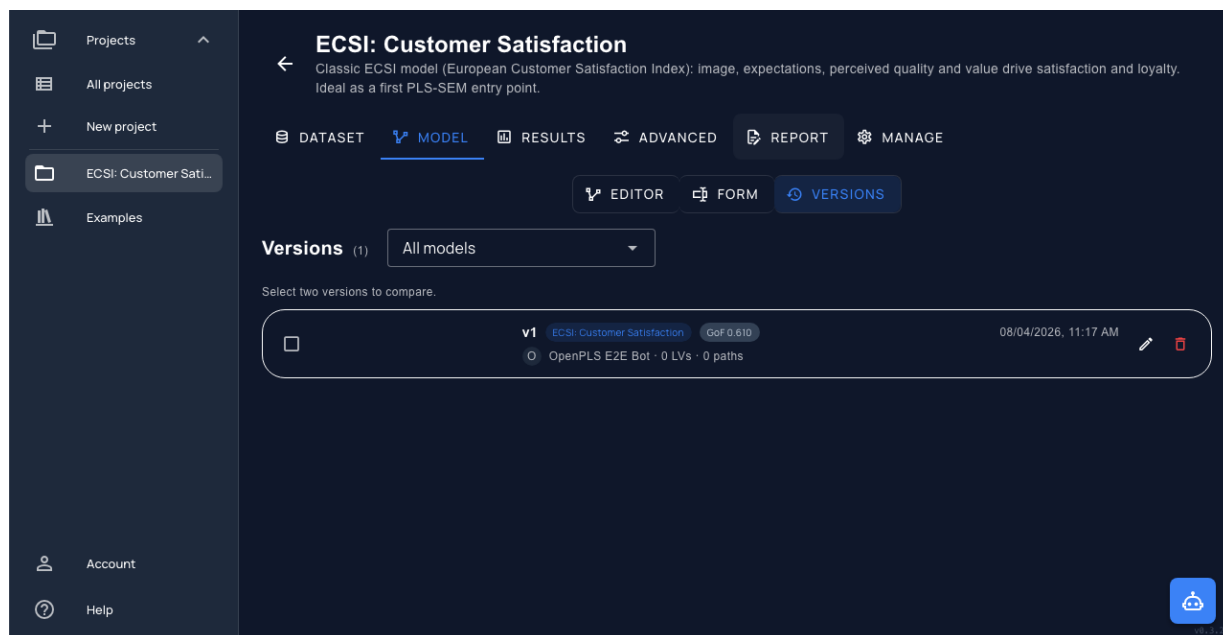


Abbildung 10 Die Versions-Liste. Jede vorherige Berechnung ist eine Karte mit LV-/Pfadzählung, GoF und Ersteller-Avatar.

Pfadpositionen, Indikatoren, Modus, alles.

Zwei Versionen vergleichen: Aktivieren Sie das Kontrollkästchen auf zwei Karten. Rechts öffnet sich ein Side-by-Side-Diff, das anzeigt, welche Pfade zwischen den beiden hinzugefügt, entfernt oder geändert wurden.

Das Wiederherstellen einer Version überschreibt Ihren Live-Editor-Zustand. Wenn Sie ungespeicherte Änderungen haben, klicken Sie zuerst auf Save (oder erstellen Sie eine frische Berechnung), damit ein Snapshot existiert, auf den Sie zurückgreifen können.

5 Berechnung und Lesen der Ergebnisse

Sobald das Modell mindestens eine endogene LV hat und jede LV ≥ 1 Indikator besitzt, wird der **Compute**-Button in der Editor-/Form-Toolbar solide blau. Klicken Sie ihn. Die Berechnung läuft serverseitig; typische Läufe auf ECSI-Größenordnung sind in 3-6 Sekunden abgeschlossen. Bei Erfolg navigiert OpenPLS zum Results-Tab.

5.1 Der Results-Header

Der Results-Tab (Abbildung 11) öffnet sich mit vier KPI-Karten oben, drei Export-Buttons rechts und einer Tab-Leiste darunter für die detaillierten Tabellen.

Der Header zeigt vier KPI-Karten:

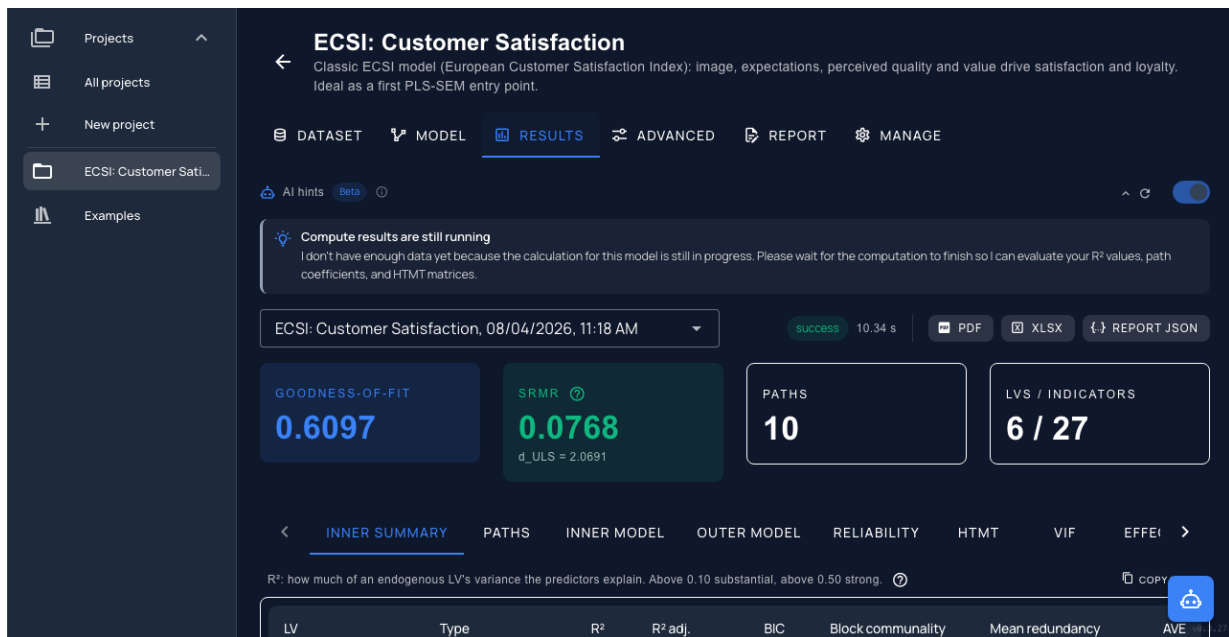


Abbildung 11 Results-Tab: KPI-Karten für Goodness-of-Fit, SRMR, Pfadanzahl, LV/MV-Anzahl. Darunter die tabbed Detailansicht: Inner Summary, Paths, Inner Model, Outer Model, Reliability, HTMT, VIF, Effect Sizes, IPMA.

1. **Goodness-of-Fit (GoF)**: $\sqrt{(\text{mean}(\text{communality}) \times \text{mean}(R^2))}$. Grober Gesamt-Fit; unter 0.10 schwach, 0.25 mittel, 0.36+ stark.
2. **SRMR**: Standardized Root Mean Square Residual. Diskrepanz zwischen beobachteten und modellimplizierten Korrelationen. Unter 0.08 gilt als konventioneller Cutoff für guten Fit; unter 0.10 akzeptabel. Der dULS-Wert im Untertitel ist ein weiteres Diskrepanzmaß.
3. **Paths**: die Anzahl der Innenmodell-Pfade im geschätzten Modell.
4. **LVs / Indicators**: die Anzahl der latenten Variablen und Indikatoren gesamt (MVs). Nützlich als Plausibilitätsprüfung beim Vergleich von Läufen.

Über den KPIs sitzt ein Result-Selector-Dropdown, das jede vergangene Berechnung für das aktuelle Modell auflistet, mit Zeitstempel. Wechseln Sie zwischen historischen Berechnungen, ohne den Tab zu verlassen.

Rechts vom Selector befinden sich die drei **Export**-Buttons:

- **PDF**: ein druckfertiger Report (Pfaddiagramm, Tabellen, Literaturhinweise).
- **XLSX**: jede Ergebnistabelle in einer Arbeitsmappe, ein Sheet pro Metrik (Paths, Inner summary, Outer loadings, Reliability, HTMT, Fornell-Larcker, VIF, Effects, ...).
- **Report JSON**: ein in sich geschlossenes JSON-Bundle mit dem Modell, dem Datensatz-Schema und jeder Ergebnistabelle. Ideal als Zusatzmaterial in einer Publikation: dieselbe Berechnung kann später aus dem Bundle exakt reproduziert werden.

5.2 Inner summary (Standard)

Die **Inner Summary**-Tabelle listet jede latente Variable mit:

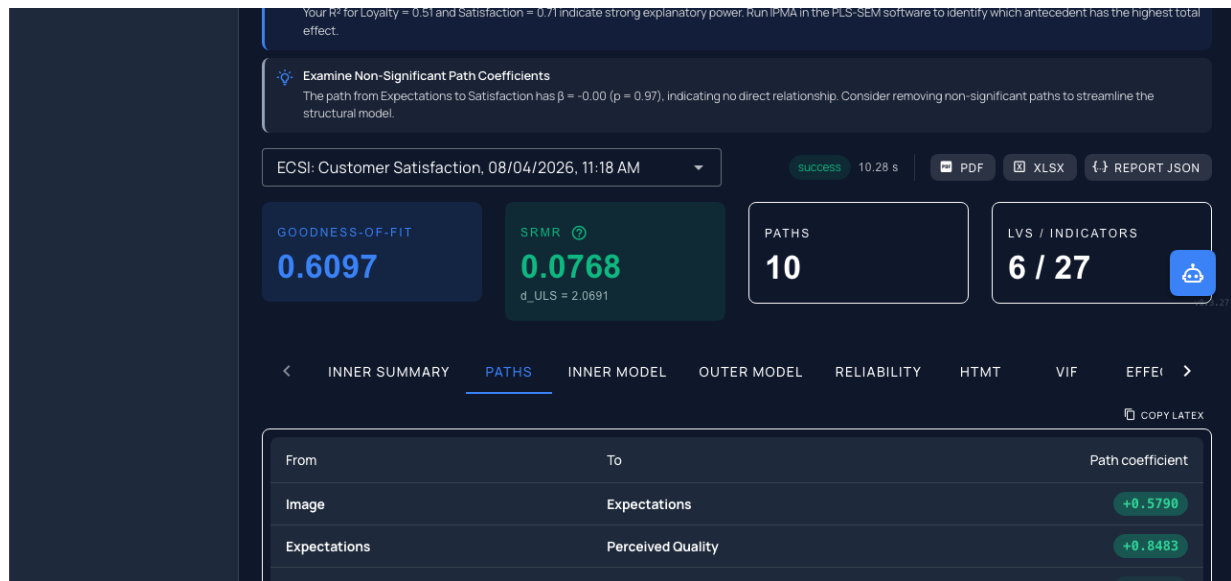


Abbildung 12 Der Paths-Untertab: jeder Innenmodell-Pfad mit Schätzwert, Standardfehler, t-Statistik, p-Wert und 95%-KI.

- **R²**: durch ihre Vorgänger erklärte Varianz der LV. 0.10 schwach, 0.25 mittel, 0.50 stark. Exogene LVs (keine Vorgänger) zeigen 0.
- **R² adj.**: adjustiert für die Anzahl der Prädiktoren. Beim Vergleich zwischen Modellen zu bevorzugen.
- **BIC**: Bayesian Information Criterion für das äußere Modell der LV. Niedriger ist besser; nur zwischen genesteten Spezifikationen sinnvoll.
- **Block communality**: durchschnittliche quadrierte Ladung der Indikatoren der LV. Höher = stärkere Messung.
- **Mean redundancy**: wie viel Varianz im Block "redundant" mit den Prädiktoren der LV ist.
- **AVE**: Average Variance Extracted. Über 0.5 bedeutet, dass die LV mehr von der Varianz ihrer Indikatoren einfängt als der Fehler.

5.3 Paths

Wechseln Sie zum **Paths**-Untertab (Abbildung 12) für die Tabelle der Strukturkoeffizienten.

Für jeden Pfad (Quelle → Ziel) listet die Tabelle:

- **Estimate**: der standardisierte Pfadkoeffizient. Interpretation wie ein Beta aus einer OLS-Regression.
- **Std. err.**: aus dem Bootstrap.
- **t**: estimate / std. err. Über 1.96 ist bei $\alpha = 0.05$ in einem zweiseitigen Test signifikant.
- **p-value**: Bootstrap-p, entsprechend dem t.
- **95% CI (lower, upper)**: Perzentil-Bootstrap-KI. Schließt es null aus, ist der Pfad signifikant.

Kopieren Sie die LaTeX-fertige Version über das kleine Copy-Symbol oben rechts in der Tabelle.

5.4 Inner model / outer model

Der **Inner model**-Tab stellt die Pfade in einer breiteren Tabelle mit den LV-Paaren in den Zeilen dar. **Outer model** dreht dies um: eine Zeile pro Indikator mit seiner LV, Ladung (Mode A) oder Gewicht (Mode B) und t-Statistik.

Interpretation der Ladungen: Über 0.708 ($\sqrt{0.5}$) bedeutet, dass die LV $\geq 50\%$ der Varianz des Indikators erklärt. Ladungen unter 0.4 sind Kandidaten zur Entfernung; zwischen 0.4 und 0.7 nur behalten, wenn ihre Entfernung das AVE ruiniert.

5.5 Reliability

Der Reliability-Tab zeigt zwei Indizes pro LV plus zwei blockbezogene Diagnostiken:

- **Cronbach's α** : der klassische Index. Setzt tau-äquivalente Ladungen voraus; konservativer.
- **Composite reliability (Dillon-Goldstein ρ)**: der moderne PLS-SEM-Standard. 0.7-0.9 akzeptabel, > 0.95 deutet auf redundante Indikatoren hin.
- **1st / 2nd eigenvalue**: der Korrelationsmatrix des Blocks. Eine klare Lücke zwischen ihnen (1. deutlich > 1 , 2. ≈ 1) unterstützt die Unidimensionalität eines reflektiven Blocks.

Zielen Sie auf $\alpha > 0.7$ und $\rho > 0.7$ bei reflektiven Konstrukten. Bei formativen Konstrukten ist Reliabilität nicht sinnvoll, prüfen Sie stattdessen den VIF.

Die konsistente Reliabilität nach Dijkstra & Henseler (ρ_A) wird nur ausgewiesen, wenn Sie PLSc ausführen (siehe Abschnitt 6.6).

5.6 HTMT

Der **HTMT**-Untertab (Abbildung 13) zeigt das Heterotrait-Monotrait-Verhältnis für jedes Konstruktpaar, wobei rote Zellen Werte oberhalb der Schwelle für Diskriminanzvalidität hervorheben.

HTMT ist der moderne Test für Diskriminanzvalidität. Für jedes Konstruktpaar berechnet er das Verhältnis der "Heterotrait"-Korrelationen (konstruktübergreifend) zu den "Monotrait"-Korrelationen (konstruktinnerhalb). Nähert sich dieses Verhältnis 1, sind die beiden Konstrukte nicht unterscheidbar.

Schwellenwerte:

- **< 0.85** (Henseler, Ringle & Sarstedt 2015, konservativ), sicher
- **< 0.90** (Kline 2011, liberal), akzeptabel
- **≥ 0.90** : Konstrukte zusammenführen, eines fallen lassen oder die Indikatoren erneut prüfen

OpenPLS führt außerdem eine Bootstrap-Version (HTMT-BST) aus, die pro Paar ein Konfidenzintervall zurückgibt. Schließt das obere KI 1.0 aus, sind die Konstrukte nachweisbar diskriminant.

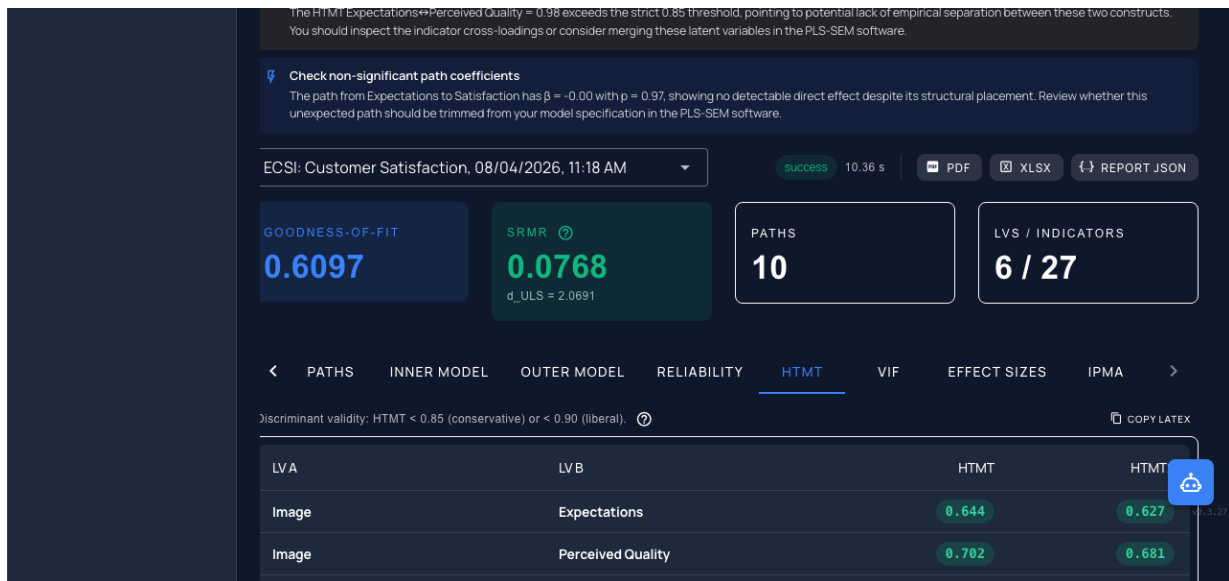


Abbildung 13 Der HTMT-Untertab: Heterotrait-Monotrait-Verhältnisse in einer Dreiecksmatrix. Zellen oberhalb der Schwelle (0.85 konservativ, 0.90 liberal) sind hervorgehoben.

5.7 VIF

Variance Inflation Factor. Zwei Ausprägungen:

- **Outer VIF:** für Indikatoren innerhalb einer formativen LV. Erkennt Redundanz zwischen formativen Items; unter 5 halten (unter 3 strenger).
- **Inner VIF:** für Prädiktor-LVs in jedem endogenen Block. Erkennt Kollinearität zwischen Prädiktoren; dieselben Schwellen.

5.8 Effects

Direkte, indirekte und totale Effekte für jedes Quelle-→-Ziel-Paar. Nützlich beim Berichten von Mediationen: der direkte Effekt ist der gezeichnete Pfeil, der indirekte Effekt ist die Summe aller Pfade durch zwischengeschaltete LVs, und der totale ist deren Summe.

5.9 Model fit

Unterhalb des KPI-Bandes zeigt OpenPLS zusätzlich:

- **SRMR** und **dULS:** bereits oben behandelt.
- **Fornell-Larcker-Matrix:** die klassische Alternative zu HTMT. Die außerdiagonalen quadrierten Korrelationen müssen kleiner sein als das diagonale $\sqrt{AVE}$, um Diskriminanzvalidität zu gewährleisten.

Zwischen HTMT und Fornell-Larcker ist HTMT die neuere und sensitivere Diagnostik. Berichten Sie beide, wenn Ihre Reviewer eines der beiden erwarten.

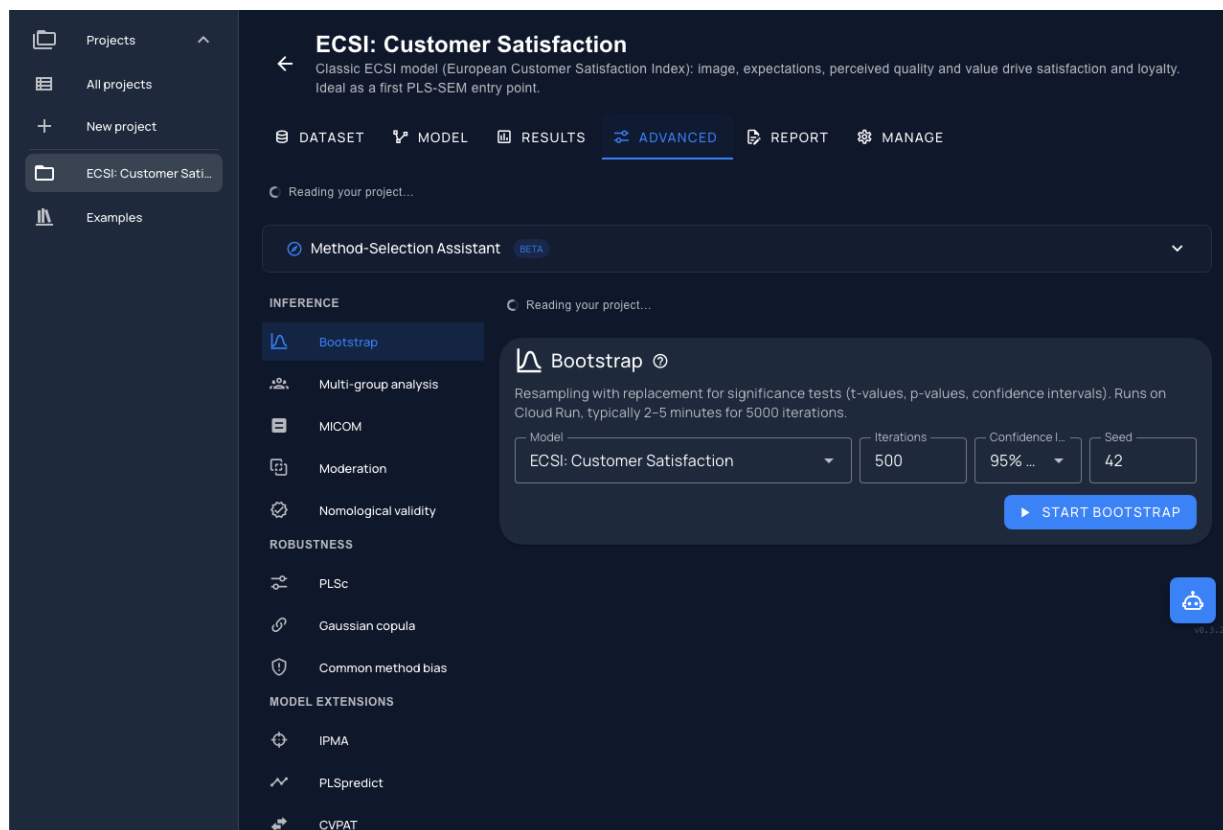


Abbildung 14 Der Advanced-Tab: 13 Methoden in der Seitennavigation, aktuelles Panel rechts.

6 Fortgeschrittene Methoden

Der **Advanced**-Tab ist der Ort, an dem OpenPLS über einen Basis-PLS-Lauf hinausgeht. Dreizehn Methoden sind in drei Gruppen in der Seitennavigation untergebracht:

- **Inference:** Bootstrap, MGA, MICOM, Moderation, Nomologische Validität
- **Robustness:** PLSc, Gaussian Copula, Common Method Bias
- **Model extensions:** IPMA, PLSpredict, CVPAT, FIMIX-PLS, Higher-Order-Model

Jede Methode folgt derselben Struktur: Konfigurieren Sie die Eingaben oben, klicken Sie auf Run, warten Sie, bis der Status-Chip grün wird, und inspizieren Sie die Ergebnistabellen darunter. Jeder vergangene Lauf bleibt erhalten und kann über das “Past runs”-Dropdown unten ausgewählt werden (Abbildung 14).

6.1 Bootstrap

Zweck. Berechnung von Standardfehlern, t-Statistiken, p-Werten und Konfidenzintervallen für jeden Pfadkoeffizienten und jede äußere Ladung. Berechnet außerdem bias-korrigierte akzelebrierte (BCa) Intervalle für Mediationseffekte (Abbildung 15).

Eingaben.

1. **Iterations:** Standard 500, publikationsreif ≥ 5000 . Jede Iteration zieht mit Zurücklegen und

AI hints Beta

Run IPMA for Loyalty

Your R^2 for Loyalty = 0.51, making it a strong candidate for importance-performance matrix analysis. Run IPMA in the PLS-SEM software to contrast total effects against latent variable index values.

Evaluate PLSpredict for key endogenous targets

Your R^2 for Satisfaction = 0.71 and Loyalty = 0.51 exceed the 0.19 threshold. Run PLSpredict in the PLS-SEM software to generate case-level predictions and evaluate out-of-sample performance.

Check HTMT between Expectations and Perceived Quality

Your HTMT Expectations↔Perceived Quality = 0.98, signaling potential discriminant validity concerns. Consider treating these constructs as a higher-order construct in the PLS-SEM software.

Method-Selection Assistant BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

AI hints Beta

Evaluate Discriminant Validity Against HTMT Thresholds

Your HTMT between Expectations and Perceived Quality reaches 0.98, exceeding the conservative 0.85 threshold. Review indicator overlap for these latent variables before finalizing model inferences.

Check Fornell-Larcker Criterion Violations

Fornell-Larcker evaluation shows some latent variables share higher correlations with others than their own square root of AVE (e.g. Expectations and Perceived Quality). Inspect cross-loadings to ensure unidimensionality.

Assess Structural Path Significance and Effect Sizes

Your R^2 for Loyalty = 0.51 and Satisfaction = 0.71 show strong explanatory power, but paths like Expectations -> Satisfaction ($p = 0.97$) are non-significant. Consider trimming unsupported paths to streamline the structural model.

Bootstrap

Resampling with replacement for significance tests (t-values, p-values, confidence intervals). Runs on Cloud Run, typically 2-5 minutes for 5000 iterations.

Model

Iterations

Confidence I...

Seed

ECSI: Customer Satisfaction

500

95% ...

42

START BOOTSTRAP

Bootstrap runs

Run

500 iter., success, 2 min ago

success 500 iterations · $\alpha = 0.05$ · Seed 42 · 144.5s

Abbildung 15 Bootstrap-Panel mit Iterationsanzahl, Seed und Run-Button.

schätzt das Modell erneut, die Laufzeit ist ungefähr linear in dieser Zahl.

2. **Seed:** optional, setzt den Zufallszahlengenerator. Setzen Sie einen Seed, wenn Sie perfekt reproduzierbare Standardfehler benötigen.

Klicken Sie auf **Start bootstrap**, um zu starten. Läuft auf einem Cloud-Run-Service; das Panel streamt Fortschritt + ETA während der Iterationen.

Interpretation. Für jeden Pfad liefert der Bootstrap eine Verteilung über B Replikate. Aus dieser Verteilung berechnet OpenPLS den mittleren Schätzwert, SE, $t = \text{mean}/\text{SE}$, p-Wert und Perzentil-KIs.

6.2 Multi-Group-Analysis (MGA)

Zweck. Testen, ob sich Pfadkoeffizienten zwischen Gruppen unterscheiden (Frauen vs. Männer, EU vs. USA, Treatment vs. Kontrolle) (Abbildung 16).

Eingaben.

1. **Grouping variable:** eine Spalte des Datensatzes, welche die Stichprobe aufteilt. Numerische Spalten können nach Range aufgeteilt werden; String-Spalten nach Wert.
2. **Groups:** eine Karte pro Gruppe. Bei String-Spalten die Werte ankreuzen, die in jede Gruppe gehören (z. B. Gruppe A = {Automotive}, Gruppe B = {Fashion}). Bei numerischen Spalten min/max setzen.
3. **Permutations:** 200 schnell, 1000 Standard, 5000+ für die Veröffentlichung. MGA führt pro Pfad einen permutationsbasierten Test durch.

Interpretation. Für jeden Pfad zeigt die Tabelle den Schätzwert in jeder Gruppe und einen p-Wert unter der Nullhypothese "Pfad ist über die Gruppen gleich". Ist $p < 0.05$, unterscheidet sich der Effekt sinnvoll zwischen den Gruppen.

Führen Sie immer zuerst **MICOM** aus. Wenn Messinvarianz in Schritt 2 (kompositionell) scheitert, bedeuten die LVs in jeder Gruppe etwas Unterschiedliches und der Vergleich der Pfade wird bedeutungslos.

6.3 MICOM

Zweck. Test der **Measurement Invariance of Composite Models** zwischen zwei Gruppen. Voraussetzung für MGA/Moderation über Gruppen hinweg (Abbildung 17).

Eingaben. Identisch mit MGA, Gruppierungsvariable + genau zwei Gruppen + Permutationsanzahl.

Interpretation. Drei-Schritt-Verdikt pro Konstrukt:

1. **Configural invariance:** gleiche Indikatoren, gleicher Algorithmus. Durch OpenPLS automatisch erfüllt.

AI hints Beta ⓘ

Run IPMA for Brand Loyalty
Your R^2 for Brand Loyalty = 0.29, making it a strong target construct for the Importance-Performance Map Analysis. Use the advanced tab to run IPMA and identify which experiential dimensions drive loyalty most efficiently.

Evaluate out-of-sample prediction with PLSpredict
With endogenous R^2 values above 0.19 across Brand Loyalty (0.29), Brand Personality (0.24), and Brand Satisfaction (0.24), your model meets the criteria for predictive relevance. Set up PLSpredict in the advanced tab to test case-level indicator predictions.

Test for unobserved heterogeneity using FIMIX-PLS
Your sample size of $n = 380$ provides adequate statistical power to detect unobserved customer segments. Run FIMIX-PLS in the advanced tab to check whether latent segments affect your structural paths.

Method-Selection Assistant BETA ▾

INFERENCE AI hints Beta ⓘ

- Bootstrap
- Multi-group analysis**
- MICOM
- Moderation
- Nomological validity

ROBUSTNESS

- PLSc
- Gaussian copula
- Common method bias

MODEL EXTENSIONS

- IPMA
- PLSpredict
- CVPAT
- FIMIX-PLS
- Higher-order model

Multi-Group Analysis (MGA) ⓘ
Compares path coefficients between two or more groups using Henseler's permutation test.

Model: Brand Experience (Brakus & Schmitt) Grouping variable: category (string)

Groups + ADD GROUP

Name: Group 1

Values (multi-select): Automotive

Name: Group 2

Values (multi-select): Fashion

Permutations: 200 **COMPUTE MGA**

Previous MGA runs

Run: category · 200 permutations · 2 min ago

success Column: category · Iterations: 200 · Duration: 141.3s

Path coefficients per group

Abbildung 16 Das MGA-Panel: eine Gruppierungsspalte auswählen, eine Karte pro Gruppe definieren, Permutationsanzahl setzen.

The screenshot displays the OpenPLS Method-Selection Assistant interface. At the top, there are three hint cards: 'Run IPMA for Brand Loyalty', 'Evaluate PLSpredict for Endogenous Constructs', and 'Test for Unobserved Heterogeneity via FIMIX-PLS'. Below these is the 'Method-Selection Assistant' header with a 'BETA' tag. The left sidebar lists various inference and robustness methods, with 'MICOM' selected under the 'INFERENCE' section. The main panel is titled 'MICOM – Measurement Invariance' and describes a three-step test. It includes instructions for Step 1 (Configural), Step 2 (Compositional), and Step 3 (Scalar). The configuration section shows 'Brand Experience (Brakus & Schmitt)' as the model and 'category (string)' as the grouping variable. Under 'Groups (exactly two)', Group A is 'Group 1' with 'Automotive' as a value, and Group B is 'Group 2' with 'Fashion' as a value. The number of permutations is set to 200. A 'COMPUTE MICOM' button is visible. At the bottom, a section for 'Previous MICOM runs' shows a run for 'category · Group 1 vs. Group 2' completed 1 minute ago.

AI hints **BETA**

Run IPMA for Brand Loyalty
Your R^2 for Brand Loyalty = 0.29. Run Importance-Performance Map Analysis in the PLS-SEM software to identify which experience dimensions drive loyalty relative to their performance.

Evaluate PLSpredict for Endogenous Constructs
Your endogenous constructs have meaningful R^2 values, such as Brand Loyalty at 0.29 and Brand Satisfaction at 0.24. Run PLSpredict in the PLS-SEM software to assess out-of-sample predictive relevance using case-level RMSE or MAE values.

Test for Unobserved Heterogeneity via FIMIX-PLS
Your sample size is $n = 380$, which easily clears the threshold for segmentation analysis. Run FIMIX-PLS in the PLS-SEM software to test if unobserved heterogeneity affects your path coefficients.

Method-Selection Assistant **BETA**

INFERENCE

- Bootstrap
- Multi-group analysis
- MICOM**
- Moderation
- Nomological validity

ROBUSTNESS

- PLSc
- Gaussian copula
- Common method bias

MODEL EXTENSIONS

- IPMA
- PLSpredict
- CVPAT
- FIMIX-PLS
- Higher-order model

MICOM – Measurement Invariance

Three-step test (configural, compositional, scalar) that constructs are measured equivalently across two groups before running MGA or moderation.

Step 1 (Configural) same indicators, same algorithm, same data treatment in both groups — auto-satisfied by OpenPLS.
Step 2 (Compositional) the composite scores from group A and group B must correlate close to 1. A p-value below 0.05 means measurement is not comparable.
Step 3 (Scalar) given Step 2 holds, composite means and variances must be equal across groups. If both pass: full invariance. If only Step 2 passes: partial invariance — MGA is still valid for path comparison.

Model: Brand Experience (Brakus & Schmitt) Grouping variable: category (string)

Groups (exactly two)

Group A: Group 1
Values (multi-select): Automotive

Group B: Group 2
Values (multi-select): Fashion

Permutations: 200

COMPUTE MICOM

Previous MICOM runs

Run: category · Group 1 vs. Group 2 · 1 min ago

Abbildung 17 MICOM-Panel: dieselbe Gruppierungsspalten- und Gruppendefinitions-UI wie MGA.

2. **Compositional invariance:** die Composite-Scores jeder LV korrelieren zwischen den Gruppen nahezu mit 1. Ein p-Wert < 0.05 bedeutet, dass sich die Composites sinnvoll unterscheiden, stopp, MGA-Ergebnisse sind nicht interpretierbar.
3. **Scalar invariance:** Composite-Mittelwerte und -Varianzen sind gleich. Wenn sowohl Schritt 2 als auch Schritt 3 bestehen: **vollständige Invarianz**. Wenn nur Schritt 2 besteht: **partielle Invarianz**: MGA ist noch gültig für den Vergleich von Pfaden, nur nicht für den Vergleich latenter Mittelwerte.

6.4 Moderation

Zweck. Testen, ob der Effekt von X auf Y von einem Moderator Z abhängt. Verwendet den Zwei-Stufen-Ansatz (Henseler & Chin 2010) (Abbildung 18).

Eingaben.

1. **Predictor (independent):** das X in $Y = X + Z + X \times Z$.
2. **Moderator:** das Z.
3. **Target (dependent):** das Y.
4. **Interaction LV name:** optional, Standard ist X_x_Z .

Interpretation. OpenPLS fittet das Modell mit einer zusätzlichen Interaktions-LV erneut (Produkt aus Prädiktor- und Moderator-Composite-Scores). Die Zeile **interaction effect** ist die Kernaussage: Ist ihr p-Wert < 0.05 , verändert der Moderator die $X \rightarrow Y$ -Beziehung.

Der **simple slopes**-Plot unter den Tabellen visualisiert dies: $X \rightarrow Y$ bei niedrigen ($M - 1$ SD), mittleren und hohen ($M + 1$ SD) Moderatorwerten. Fächern die Slopes auf, verstärkt der Moderator X; konvergieren sie, schwächt er ihn ab.

6.5 Nomologische Validität

Zweck. Prüfen, dass jeder hypothesierte Pfad mit dem erwarteten Vorzeichen und der erwarteten Signifikanz herauskam. Plausibilitätsprüfung vor der Ergebnisdokumentation.

Eingaben. Automatisch aus Ihrem Modell befüllt, jeder Pfad wird eine Zeile mit einem **expected sign**-Toggle: + für "hypothesiert positiv" oder – für "hypothesiert negativ". Setzen Sie das Vorzeichen für jede Hypothese, wählen Sie das Signifikanzniveau α (Standard 0.05) und klicken Sie auf Run test.

Interpretation. Für jede Hypothese zeigt die Tabelle das geschätzte Vorzeichen, den p-Wert aus dem letzten Bootstrap und ein Urteil (pass / fail) gegen Ihr erwartetes Vorzeichen bei α . Fehlschläge sind Kandidaten für eine erneute Untersuchung oder für den Bericht als Nullbefund.

6.6 PLSc

Zweck. Konsistente PLS-Schätzung. Korrigiert den Attenuation-Bias des Standard-PLS, wenn Konstrukte tatsächlich Common-Factor- statt Composite-Konstrukte sind (Abbildung 19).

Simple slopes

Abbildung 18 Das Moderation-Panel: drei LV-Dropdowns (Predictor / Moderator / Target).

Method-Selection Assistant BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

PLSc – Consistent PLS

Bias-corrected estimates for reflective constructs (Dijkstra & Henseler 2015). Returns ρ_A , dis-attenuated correlations, and corrected paths/loadings.

Model

ECSI: Customer Satisfaction

COMPUTE PLSc

Past runs

Run

PLSc · just now

success

PLSc stability warning

Maximum dis-attenuated correlation is 0.976 (≥ 0.85 ; Henseler/Ringle/Sarstedt 2015 HTMT threshold).

Maximum inner VIF is 56.68 (≥ 10 ; severe collinearity, Hair et al. 2022).

Under these conditions PLSc can return corrected paths outside $[-1, 1]$. Consider merging the highly-correlated constructs, modelling them as a higher-order construct, or dropping one.

SRMR: 0.061

d_ULS: 1296

Reliability per construct

LV	ρ_A	ρ_C	AVE	R^2	R^2 adj.	BIC
Image	0.858	0.830	0.503	—	—	—
Expectations	0.855	0.847	0.527	0.427	0.425	-129.38
Perceived Quality	0.878	0.872	0.578	0.952	0.952	-750.91
Perceived Value	0.853	0.835	0.563	0.861	0.859	-476.90
Satisfaction	0.891	0.896	0.596	0.877	0.871	-480.46

Abbildung 19 PLSc-Panel, One-Click bei Modellen mit ausschließlich Mode-A-LVs.

29

Method-Selection Assistant BETA

INFERENCE

- Bootstrap
- Multi-group analysis
- MICOM
- Moderation
- Nomological validity

ROBUSTNESS

- PLSc
- Gaussian copula**
- Common method bias

MODEL EXTENSIONS

- IPMA
- PLSpredict
- CVPAT
- FIMIX-PLS
- Higher-order model

Gaussian copula – endogeneity test ⓘ

Park & Gupta (2012) / Hult et al. (2018): tests whether a predictor correlates with the error term (endogeneity) and returns endogeneity-corrected path estimates.

Model: **ECSI: Customer Satisfaction** Endogenous LV: **Satisfaction**

Suspected predictors (optional): **200** Bootstrap resamples: **42** Seed (optional): **42**

Leave empty to test all predecessors.

COMPUTE COPULA TEST

Past runs

Run: **Endo=Satisfaction · just now**

success Endogenous LV: Satisfaction · Bootstrap resamples = 200

Copula coefficients per predictor

Predictor	γ (copula term)	SE	t	p	CvM p (non-norm.)	Decision
Image	-0.020	0.226	-0.088	0.930	0.042	no endogeneity detected
Expectations	0.133	0.226	0.591	0.554	0.077	copula not admissible (normal)
Perceived Quality	0.122	0.364	0.334	0.738	0.044	no endogeneity detected
Perceived Value	-0.400	0.255	-1.569	0.117	0.010	no endogeneity detected

Endogeneity-corrected path estimates

Path Estimate

COPY **CSV**

Abbildung 20 Das Copula-Panel: Auswahl der endogenen LV, Bootstrap-Resamples.

Eingaben. Keine, klicken Sie auf Run.

Interpretation. Führt das exakt gleiche Modell mit angewandter PLSc-Korrektur erneut aus. Pfadkoeffizienten bewegen sich typischerweise leicht weg von null (die Korrektur beseitigt die Attenuation aus endlicher Reliabilität). Berichten Sie sowohl Standard-PLS als auch PLSc, wenn Reviewer eine faktorbasierte Interpretation erwarten.

Nur gültig, wenn. Alle LVs sind Mode A (reflektiv). Mode-B-LVs brechen die PLSc-Annahme.

6.7 Gaussian Copula (Endogenität)

Zweck. Testen, ob ein Prädiktor endogen ist (mit dem Fehlerterm korreliert). Wenn ja, endogenitätskorrigierte Schätzungen über den Park- & Gupta- (2012) / Hult- et al.- (2018)-Ansatz zurückgeben (Abbildung 20).

Eingaben.

1. **Endogenous LV:** die LV auswählen, deren Prädiktoren Sie im Verdacht haben, mit dem Fehlerterm korreliert zu sein.
2. **Suspected predictors:** optional; leer lassen, um alle Vorgänger der endogenen LV zu testen.
3. **Bootstrap resamples:** Standard 500, für publikationsreife KIs auf 1000-5000 erhöhen.
4. **Seed:** optional.

Interpretation. Für jeden verdächtigten Prädiktor berichtet OpenPLS einen Copula-Term-Koeffizienten γ , seinen p-Wert und ein CvM (Cramér-von Mises)-p für Nicht-Normalität des Prädiktors (durch die Methode gefordert). Die Tabelle der augmentierten Pfade zeigt die korrigierten Koeffizienten, wenn der Copula-Term signifikant ist.

Nur gültig, wenn. Prädiktoren sind nicht normalverteilt. Ist der CvM-p eines Prädiktors > 0.10 , kann der Copula-Korrektur nicht getraut werden.

6.8 Common Method Bias (CMB, Lindell-Whitney)

Zweck. Method-Bias über die Marker-Variable-Prozedur von Lindell & Whitney (2001) erkennen (Abbildung 21).

Eingaben.

1. **Marker variable:** eine theoretisch unverwandte numerische Spalte, die KEIN Indikator einer LV im Modell ist. Beispiele: Demografie wie `tenure_years` oder `age`.
2. **Signifikanz α :** Standard 0.05.

Interpretation. OpenPLS partialisiert die kleinste beobachtete Marker-LV-Korrelation aus jeder Konstruktkorrelation heraus und berichtet, welche Pfade an Signifikanz verlieren. Pfade, die von signifikant zu nicht signifikant kippen, sind verdächtig für eine Common-Method-Inflation.

Nur gültig, wenn. Ihr Datensatz eine numerische Spalte außerhalb des Modells hat. Wenn nicht, wählen Sie eine andere Stichprobe oder führen Sie Harmans Einzelfaktor-Test manuell durch.

6.9 IPMA

Zweck. Importance-Performance Map Analysis. Beantwortet: "Welche Konstrukte sind für das Ziel am wichtigsten und welche davon sind unterversorgt?" (Abbildung 22)

Eingaben.

1. **Target LV:** gewöhnlich das Ergebnis, das Sie interessiert (Loyalty, Adoption, ...).

Interpretation. Für jeden Vorgänger des Ziels berechnet OpenPLS:

- **Importance** = totaler Effekt (direkt + indirekt) auf das Ziel
- **Performance** = Mittelwert des Composite-Scores der LV, reskaliert auf 0-100

Als Streudiagramm geplottet (Performance auf x, Importance auf y). Oben-links = hohe Importance, niedrige Performance = der größte Verbesserungshebel. Oben-rechts = beides hoch

Method-Selection Assistant BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

Common method bias (Lindell-Whitney)

Marker-variable procedure: partial out the smallest marker-LV correlation and check which path correlations lose significance (Lindell & Whitney 2001).

Model

Employee Engagement (JD-R)

Marker variable

tenure_years (numeric)

Significance level (α)

0.05

▶ RUN TEST

A theoretically unrelated numeric column from the dataset. Cannot be an indicator of any LV in the model.

Past runs

Run

tenure_years · just now

success

Marker variable: tenure_years · $r_M = 0.003$ · $\alpha = 0.05$

Marker-variable correlations

LV	$r(\text{marker, LV})$
Job Resources	0.003
Job Demands	-0.133
Work Engagement	-0.065
Job Satisfaction	0.010
Job Performance	-0.047

Pairwise LV correlations before / after adjustment

COPY LATEX

LV pair	r (unadjusted)	p (unadjusted)	r (adjusted)	p (adjusted)	Verdict
Job Resources ↔ Job Demands	0.060	0.265	0.057	0.287	unchanged
Job Resources ↔ Work Engagement	0.538	< 0.001	0.536	< 0.001	unchanged
Job Resources ↔ Job Satisfaction	0.466	< 0.001	0.464	< 0.001	unchanged
Job Resources ↔ Job Performance	0.375	< 0.001	0.374	< 0.001	unchanged
Job Demands ↔ Work Engagement	-0.174	0.001	-0.177	< 0.001	unchanged

Abbildung 21 CMB-Panel: Marker-Spalte wählen und ausführen.

Method-Selection Assistant

BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

Importance-Performance Map (IPMA)

Identify constructs and indicators with high importance but low performance, prime levers for improvement.

Model

ECSI: Customer Satisfaction

Target latent variable

Expectations

Endogenous LV whose drivers you want to rank.

Scale minimum (optional)

Scale maximum (optional)

Rescales indicator scores to a 0, 100 performance index. Defaults to data min/max.

RUN IPMA

Past IPMA runs

Run

Expectations · just now

success

Target latent variable: Expectations

Latent-variable importance / performance

LV	Importance (total effect)	Performance (0-100)
Image	0.579	73.36

Indicator-level breakdown

LV	Indicator	Outer weight	Normalized weight
Image	imag1	0.098	0.146
	imag2	0.157	0.234
	imag3	0.157	0.233
	imag4	0.077	0.114
	imag5	0.184	0.274
Expectations	expe1	0.106	0.176
	expe2	0.141	0.233
	expe3	0.119	0.197
	expe4	0.099	0.165
	expe5	0.138	0.229
Perceived Quality	qual1	0.107	0.187
	qual2	0.135	0.236
	qual3	0.117	0.206
	qual4	0.096	0.168
	qual5	0.116	0.203
Perceived Value	val1	0.175	0.302

Abbildung 22 IPMA-Panel: Ziel-LV auswählen, ausführen.

= Stärken, die zu bewahren sind.

6.10 PLSpredict

Zweck. Beurteilung der **out-of-sample**-Prognoserelevanz. R^2 ist in-sample; PLSpredict sagt Ihnen, ob das Modell generalisiert (Abbildung 23).

Eingaben.

1. **k-fold:** Standard 10. Höher = kleinere Test-Folds.
2. **Repeats:** Standard 1. Auf 10+ erhöhen, um über verschiedene Fold-Zuweisungen zu mitteln und die Varianz zu reduzieren.
3. **Seed:** für Reproduzierbarkeit.

Interpretation. Für jeden Indikator listet die Tabelle:

- **Q^2_{predict} :** kreuzvalidiertes prognostisches R^2 . Über 0 sinnvoll; höher ist besser.
- **RMSE_PLS** vs. **RMSE_LM:** PLS-Vorhersagen vs. lineare Benchmark. Ist RMSE_PLS für die Mehrheit der Indikatoren kleiner, hat PLS Prognosewert jenseits der linearen Alternative.

6.11 CVPAT

Zweck. Cross-Validated Predictive Ability Test. Vergleicht das PLS-Modell mit einer Benchmark (Indikator-Mittelwert oder lineares Modell) über einen t-Test auf out-of-sample-Verlust (Abbildung 24).

Eingaben. Benchmark + k-fold + Wiederholungen + Seed. Die Standardwerte sind sicher.

Interpretation. Für jede endogene LV berichtet OpenPLS die Verlustdifferenz (PLS – Benchmark) mit einem p-Wert. Negative Verlustdifferenz + $p < 0.05$ = PLS prognostiziert besser als die Benchmark, eine starke Aussage.

6.12 FIMIX-PLS

Zweck. Finite-Mixture-PLS. Entdeckt latente Respondentenklassen, wenn Sie Heterogenität vermuten, die Gruppierung aber nicht kennen (Abbildung 25).

Eingaben.

1. **Number of classes (K):** 2..5 ausprobieren, Fit-Kriterien vergleichen.
2. **EM restarts:** das Beste aus n Zufallsstarts. Mehr = robuster.
3. **Max iterations, tolerance:** Konvergenzsteuerungen; Standards belassen.

Interpretation. Für jedes K schätzt OpenPLS klassenspezifische Pfade + einen Mischanteil ρ pro Klasse. Vergleichen Sie **AIC**, **BIC**, **CAIC** und **entropy (EN)** über die K-Werte hinweg, der Ellbogen im BIC kombiniert mit $EN > 0.6$ identifiziert das "richtige" K.

Method-Selection Assistant BETA

INFERENCE

- Bootstrap
- Multi-group analysis
- MICOM
- Moderation
- Nomological validity

ROBUSTNESS

- PLSc
- Gaussian copula
- Common method bias

MODEL EXTENSIONS

- IPMA
- PLSpredict**
- CVPAT
- FIMIX-PLS
- Higher-order model

PLSpredict

k-fold cross-validated out-of-sample predictive power vs. a linear-model benchmark (Shmueli et al. 2019).

Model: ECSI: Customer Satisfaction

Folds (k): 10 Repeats: 1 Seed (optional): 42

Number of cross-validation folds. Common: 5 or 10. Repeat the k-fold procedure with different splits and average.

RUN PLSPREDICT

Past runs

Run: k=10, no predictive power · just now

success Overall verdict: no predictive power k = 10 · Repeats: 1

0 better than LM · 22 worse · 0 tie · 22 total

Indicator summary

Out-of-sample (k-fold cross-validation).

LV	Indicator	RMSE (PLS)	MAE (PLS)	MAPE (PLS)	RMSE (LM)	MAE (LM)	MAPE (LM)	Q ² _PLS
Expectations	expe1	1.832	1.400	27.6%	1.808	1.362	26.4%	
Expectations	expe2	1.953	1.433	31.8%	1.876	1.383	29.4%	
Expectations	expe3	2.121	1.637	36.1%	2.080	1.612	34.3%	
Expectations	expe4	1.536	1.165	17.4%	1.489	1.161	17.1%	
Expectations	expe5	2.081	1.654	32.9%	2.055	1.596	30.6%	
Perceived Quality	qual1	1.903	1.462	23.1%	1.341	1.062	16.6%	
Perceived Quality	qual2	2.068	1.602	25.4%	1.337	1.059	16.3%	
Perceived Quality	qual3	2.201	1.727	36.0%	1.622	1.197	22.0%	
Perceived Quality	qual4	1.619	1.254	21.2%	1.222	0.987	16.0%	
Perceived Quality	qual5	1.971	1.527	32.3%	1.433	1.086	19.3%	
Perceived Value	val1	1.949	1.496	24.9%	1.540	1.149	18.5%	
Perceived Value	val2	1.682	1.245	20.4%	1.560	1.144	17.4%	

Abbildung 23 PLSpredict-Panel: k-fold, Wiederholungen, Seed.

Method-Selection Assistant

BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

CVPAT (Cross-Validated Predictive Ability Test)

Compare PLS out-of-sample predictions against an indicator-average or linear-model benchmark using a paired t-test (Liengaard et al. 2021).

Model

EC SI: Customer Satisfaction

Benchmark

IA (indicator average)

Choose which benchmark the PLS model is compared against.

Folds (k)

10

Repeats

1

Seed

42

Significance level (α)

0.05

RUN CVPAT

Past runs

Run

IA, k=10 · just now

success Benchmark: IA · k = 10 · repeats = 1 · α = 0.05

Overall test

COPY LATEX

n	Loss (PLS)	Loss (benchmark)	Δ (PLS - benchmark)	t	p (1-sided)	Verdict
250	83.605	100.920	-17.315	-6.62	< 0.001	PLS better

Per-construct test

COPY LATEX

LV	Loss (PLS)	Loss (benchmark)	Δ (PLS - benchmark)	p (1-sided)	Verdict
Expectations	18.358	21.673	-3.315	< 0.001	PLS better
Loyalty	19.269	23.292	-4.023	< 0.001	PLS better
Perceived Quality	19.254	23.063	-3.809	< 0.001	PLS better
Perceived Value	14.631	17.481	-2.850	< 0.001	PLS better

Abbildung 24 CVPAT-Panel: Benchmark-Picker (IA = indicator average, LM = linear model), k-fold, Wiederholungen.

Method-Selection Assistant BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

FIMIX-PLS segmentation

Finite-mixture EM detection of unobserved latent classes (Hahn et al. 2002).

Model

ECSI: Customer Satisfaction

Number of classes (K)

3

EM restarts

5

Max iterations

500

Convergence tolerance

0.000001

Seed (optional)

42

RUN FIMIX

Past runs

Run

K = 3 · just now

success K = 3

Log-likelihood: -1068.66

Parameters: 62

Class sizes (mixing proportions ρ)

COPY CSV

Class	ρ (proportion)	n (hard-assigned)
1	0.320	91
2	0.286	73
3	0.394	8

Per-class structural-model paths

Class

Class 1

Per-class structural-model paths

COPY CSV

Path	Estimate
(intercept) → Expectations	-0.349
Image → Expectations	1.012
(intercept) → Perceived Quality	0.285
Expectations → Perceived Quality	0.945
(intercept) → Perceived Value	0.193
Expectations → Perceived Value	0.385
Perceived Quality → Perceived Value	0.456

Abbildung 25 FIMIX-Panel: Anzahl der Klassen, EM-Restarts, Toleranz.

37

Sobald K entschieden ist, exportieren Sie die Klassenzuweisungen und führen Standard-PLS pro Klasse für detaillierte Inferenz erneut aus. Oder verwenden Sie sie als Gruppierung in MGA.

6.13 Higher-Order Construct (HOC)

Zweck. Mehrere First-Order-LVs über den disjunkten Zwei-Stufen-Ansatz (Sarstedt et al. 2019) in ein einziges Second-Order-Konstrukt kombinieren (Abbildung 26).

Eingaben.

1. **Name:** der LV-Name des HOC, z. B. BrandExperience.
2. **First-order LVs:** die zu kombinierenden Komponenten (mindestens 2).
3. **HOC mode:** A (reflektiv) oder B (formativ).

Interpretation. OpenPLS führt Stufe 1 (normales PLS, um First-Order-LV-Scores zu erhalten) und dann Stufe 2 (ein neues PLS mit dem HOC anstelle der vier Dimensionen) aus. Berichtet:

- **Loadings:** First-Order-LVs auf dem HOC (Mode A) oder Gewichte (Mode B).
- **R² (stage 2):** erklärte Varianz der nachgelagerten Ziele des HOC.
- **Paths (stage 2):** strukturelle Koeffizienten im neu geschätzten Modell.

Vergleichen Sie Mode A vs. Mode B durch erneutes Fitten mit dem anderen Modus und prüfen Sie, was sauberere Reliabilität und Pfade produziert.

Method-Selection Assistant

BETA

INFERENCE

Bootstrap

Multi-group analysis

MICOM

Moderation

Nomological validity

ROBUSTNESS

PLSc

Gaussian copula

Common method bias

MODEL EXTENSIONS

IPMA

PLSpredict

CVPAT

FIMIX-PLS

Higher-order model

Higher-order construct (HOC)

Disjoint two-stage approach (Sarstedt et al. 2019): combines multiple first-order LVs into a higher-order construct and re-estimates the structural model.

Model

Brand Experience (Brakus & Schmitt)

Higher-order construct name

BrandExperience

First-order LVs

Sensory Experience

Affective Experience

Intellectual Experience

Behavioral Experience

HOC mode

Mode A (reflective)

Select at least two LVs to combine into the HOC.

COMPUTE HOC

Past runs

Run

BrandExperience · Mode A · just now

success

BrandExperience ← Sensory Experience, Affective Experience, Intellectual Experience, Behavioral Experience · Mode A (reflective)

First-order loadings on the HOC

COPY

First-order LV	Loading
Sensory Experience	0.514
Affective Experience	0.730
Intellectual Experience	0.443
Behavioral Experience	0.409

R² (stage 2)

COPY

CSV

LV	R ²
Brand Loyalty	0.290
Brand Personality	0.238
Brand Satisfaction	0.222
BrandExperience	0.000

Path coefficients (stage 2)

COPY

CSV

Path	Estimate
BrandExperience → Brand Personality	0.488

Abbildung 26 HOC-Panel: HOC benennen, 2+ First-Order-LVs auswählen, Mode A oder B wählen.

7 Exporte

Jedes Ergebnis in OpenPLS kann das Werkzeug in einem von vier Formaten verlassen. Die Export-Buttons liegen im Header des Results-Tabs, über den KPI-Karten, und sind ebenfalls über das Model-Export-Dropdown in der Editor-Toolbar verfügbar.

7.1 PDF-Report

Was Sie bekommen. Ein druckfertiges PDF mit:

- Deckblatt mit Modell- + Datensatzmetadaten
- Pfaddiagramm (gerendert aus den aktuellen Editor-Positionen)
- Jede Ergebnistabelle aus dem Results-Tab, druckformatiert
- Literaturhinweise zu den OpenPLS-Methoden (jede Metrik zitiert die Publikation, die sie definiert hat)

Wann verwenden. Als Anhang zu einer Papiereinreichung, zum Teilen mit einer nicht-technischen Person oder in einem Foliensatz.

Dateiname. `{model_name}_report.pdf`, bereinigt zur Entfernung von Leerzeichen und Schrägstrichen.

7.2 Excel-Arbeitsmappe (XLSX)

Was Sie bekommen. Eine Arbeitsmappe mit einem Sheet pro Ergebnistabelle. Sheet-Namen wie sie in der Datei erscheinen:

- Overview, Lauf-Metadaten (Compute-Zeit, GoF, SRMR, LV-/Pfadanzahl)
- Inner Summary, R^2 , R^2 adjustiert, AVE, Kommunalität pro LV
- Model Fit, SRMR + d_ULS (nur wenn berechnet)
- Path Coefficients, Schätzwerte für jeden Innenmodell-Pfad
- Inner Model, mit Bootstrap-SE, t, p, KI wenn verfügbar
- Outer Model, Ladungen (Mode A) oder Gewichte (Mode B) pro Indikator
- Effect Sizes f^2 , direkte Effektstärken für jeden Prädiktorpfad
- Predictive Relevance Q^2 , kreuzvalidierte Redundanz pro LV
- VIF, Outer und VIF, Inner, Kollinearität
- HTMT, Heterotrait-Monotrait-Matrix
- Reliability, Cronbach's α , Composite Reliability, Eigenwerte
- Model, LVs und Model, Paths, die Modellspezifikation selbst

Bedingte Sheets, nur hinzugefügt, wenn der entsprechende Advanced-Lauf existiert:

- Nomological Validity, pro Hypothese Vorzeichen + Urteil
- CMB (Lindell-Whitney), marker-korrigierte Korrelationen
- CVPAT, kreuzvalidierter Predictive-Ability-Test

Wann verwenden. Reviewer möchte Rohzahlen, oder Sie möchten Tabellen im Hausstil Ihres

Papers neu formatieren.

Dateiname. {model_name}_results.xlsx.

7.3 Report-JSON-Bundle

Was Sie bekommen. Ein in sich geschlossenes JSON mit:

- Vollständige Modellspezifikation (LVs, Indikatoren, Pfade, Positionen, Messmodus)
- Datensatz-Schema (Spaltennamen, Typen, Missing-Counts)
- Jede Ergebnistabelle verbatim
- Compute-Metadaten (Zeitstempel, Dauer, OpenPLS-Version)

Wann verwenden. Als **Zusatzmaterial** für eine Publikation. Jeder, der das JSON importiert, kann exakt nachvollziehen, was Sie geschätzt haben, ohne OpenPLS auszuführen.

Dateiname. {model_name}.report.json.

7.4 Modell-Export (Editor-Toolbar)

Die Editor-Toolbar hat ein **Download**-Symbol, das ein Menü mit zwei reinen Modell-Exportformaten öffnet:

- **JSON (.openpls.json):** vollständiges OpenPLS-Modell inklusive Positionen. Round-Trip-fähig: Sie können dies in ein anderes OpenPLS-Projekt importieren, um das Modell mit einem anderen Datensatz wiederzuverwenden.
- **SmartPLS (.splsm):** eine XML-Datei, kompatibel mit SmartPLS 3 und 4. Ermöglicht die Übergabe des Modells an einen Kollegen, der mit diesem Werkzeug arbeitet.

Diese enthalten nur das Modell, keine Ergebnisse. Verwenden Sie dafür das Report-JSON-Bundle.

7.5 LaTeX-Snippets

Jede Ergebnistabelle im Results-Tab hat oben rechts ein kleines **Copy LaTeX**-Symbol. Klicken Sie es und fügen Sie direkt in die .tex-Quelle Ihres Papers ein. Der generierte LaTeX-Code verwendet booktabs (\toprule, \midrule, \bottomrule) und entspricht dem konventionellen Tabellenstil in JMS, MIS Quarterly und ähnlichen Zeitschriften.

Die LaTeX-Snippets betten Modellname, Datensatzname und OpenPLS-Version in die Tabellenüberschrift ein, um die Nachvollziehbarkeit sicherzustellen. Passen Sie die Überschrift in Ihrem Paper an, wenn Sie ein aufgeräumteres Aussehen wünschen.

7.6 Reproduzierbarkeit

Für ein Paper ist das ideale Reproduzierbarkeits-Bundle:

1. Das **Report JSON** (Ergebnisse und Modell + Datensatz-Schema).
2. Der **Rohdatensatz** (eine Zusatzdatei, falls die Lizenz es erlaubt).
3. Ein kurzer Hinweis, der für die Schätzung auf `openpls.app` verweist.

Jeder Leser kann dann exakt dieselbe Berechnung erneut ausführen, ohne von Ihrer lokalen Installation abhängig zu sein.

Die exportierten Werte sind dieselben Zahlen, die Sie im Results-Tab sehen. Vergleiche mit SmartPLS auf publizierten PLS-SEM-Datensätzen laufen; die aktuellsten Ergebnisse sind unter `openpls.app/validation` verlinkt.

8 Zusammenarbeit

OpenPLS behandelt Projekte als teilbare Arbeitsbereiche. Laden Sie Koautoren, wissenschaftliche Assistenten oder Reviewer ein, dasselbe Modell anzusehen oder zu bearbeiten. Jede Aktion hinterlässt einen Audit-Trail, damit Sie sehen können, wer wann was getan hat.

Alles, was mit Zusammenarbeit zu tun hat, liegt unter **Manage** in der Tab-Leiste. Drei Unteransichten: **Team**, **Activity**, **Settings**.

8.1 Das Team-Panel

Öffnen Sie **Manage** → **Team**, um zum in Abbildung 27 gezeigten Panel zu gelangen.

Drei Abschnitte übereinander:

1. **Team**: jedes aktuelle Mitglied mit Name, E-Mail, Avatar und Rollen-Pille. Der Owner (Projekt-Ersteller) ist markiert. Rechts von jeder Zeile befinden sich die Aktionen für Rollenwechsel und Entfernen. Die Entfernungsaktion ist nur für den Owner sichtbar.
2. **Invite person**: E-Mail-Eingabe + Rollen-Dropdown + Invite-Button. Sendet eine signierte Einladungsmail über Resend. Der Empfänger kann über einen Link akzeptieren, der `/invitations/{id}` öffnet.
3. **Pending invitations**: eine Karte pro ausstehender Einladung mit der eingeladenen E-Mail, dem Zeitstempel und einer **Revoke**-Aktion.

8.2 Rollen

Drei Rollen:

- **Owner**: Ersteller, kann alles tun, einschließlich das Projekt löschen.
- **Editor**: kann Datensatz und Modell bearbeiten, Berechnungen und fortgeschrittene Methoden ausführen, Einstellungen bearbeiten. Kann das Projekt nicht löschen oder neue Mitglieder einladen (das kann nur der Owner).
- **Viewer**: nur lesend. Sieht jedes Ergebnis und jeden Export, kann nichts ändern.

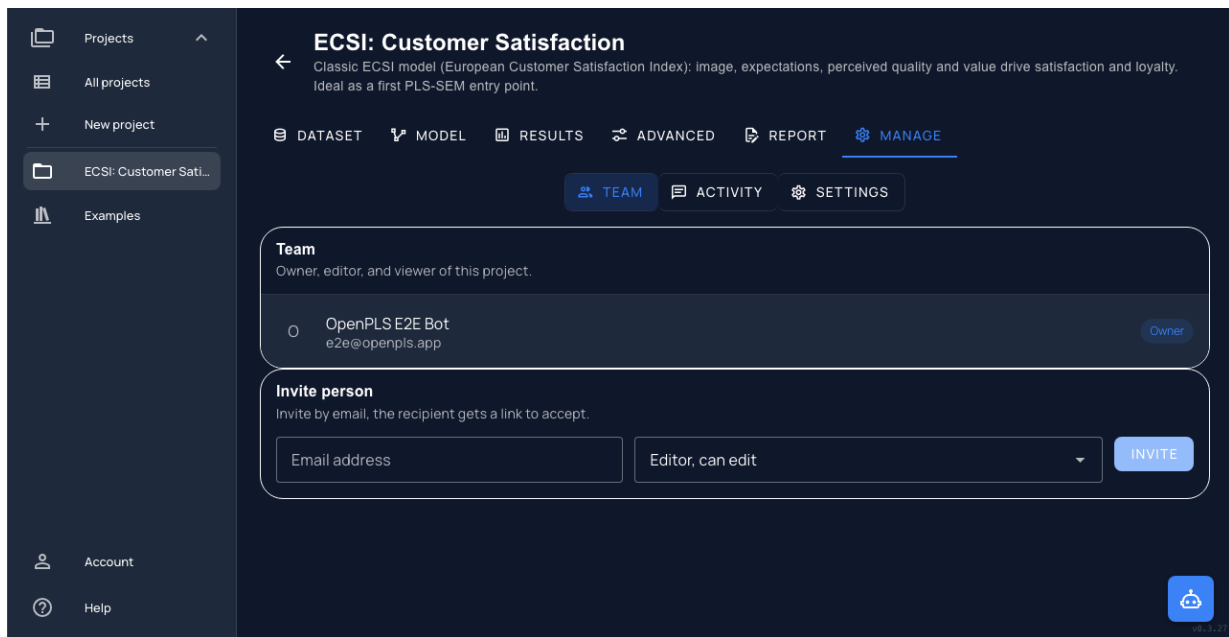


Abbildung 27 Die Team-Unteransicht: Mitgliederliste oben, Einladungsformular darunter, offene Einladungen unten.

Rollenwechsel treten sofort in Kraft und erscheinen im Activity-Log.

8.3 Eine Einladung senden

1. Geben Sie die E-Mail des Eingeladenen in das Feld **Email address** ein.
2. Wählen Sie die Rolle: **Editor, can edit** oder **Viewer, read only**.
3. Klicken Sie auf **Invite**.

OpenPLS erstellt einen Einladungseintrag in der Datenbank und sendet dem Eingeladenen eine E-Mail mit dem Accept-Link. Wenn die E-Mail-Adresse bereits einem OpenPLS-Nutzer gehört, landet die Einladung direkt in dessen eigener /invitations-Route.

Einladungslink kopieren. Jede Pending-Invitation-Karte hat ein “Copy link”-Symbol, das die Accept-URL in die Zwischenablage legt, nützlich, wenn die E-Mail-Filter des Empfängers aggressiv arbeiten.

Zurückziehen. Klicken Sie auf **Revoke** auf einer Pending-Karte. Die Einladung wird ungültig; der Accept-Link funktioniert nicht mehr.

Eingeladene Nutzer, die noch kein OpenPLS-Konto haben, können über den Accept-Flow eines erstellen (ihre E-Mail ist vorausgefüllt). Die Einladung hält den Platz frei, bis sie akzeptiert wird, es ist nicht nötig, sich zuerst zu registrieren und erneut eingeladen zu werden.

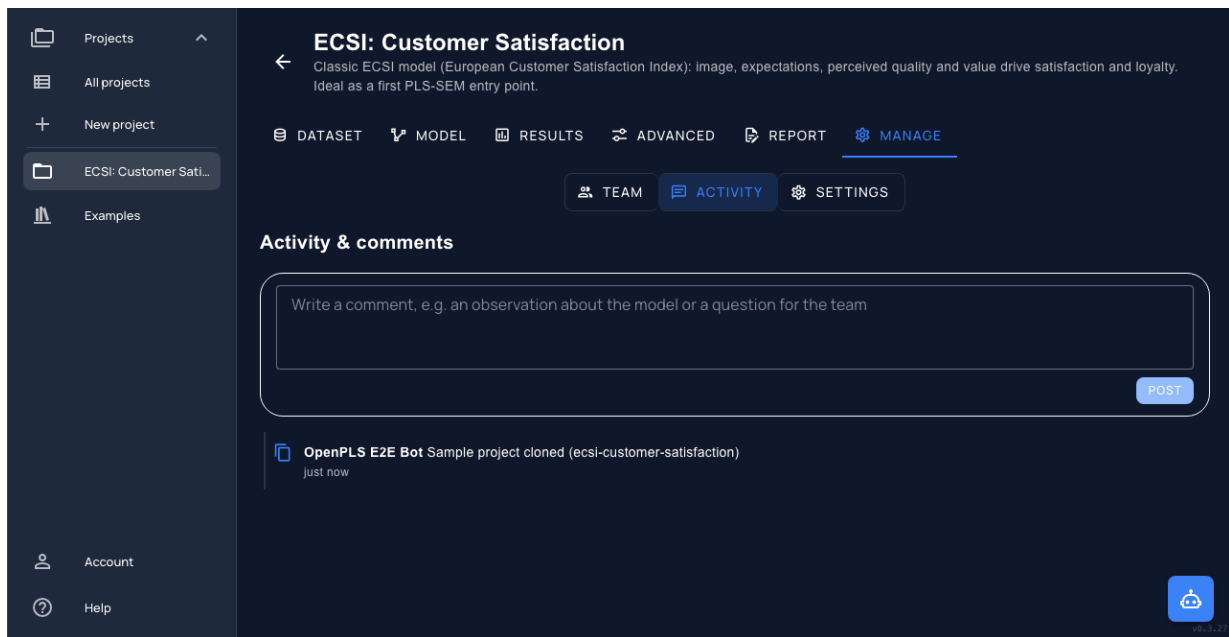


Abbildung 28 Die Activity-Unteransicht: chronologischer Feed jedes Projekt-Events, mit Icons und Akteur.

8.4 Das Activity-Log

Wechseln Sie zur **Activity**-Unteransicht (Abbildung 28) für den Audit-Trail.

Das Backend loggt einen Eintrag pro Event. Derzeit ausgelöste Events:

- **Team-Events:** Mitglied eingeladen, Einladung angenommen, Einladung zurückgezogen, Einladungsmail fehlgeschlagen, Mitgliederrolle geändert, Mitglied hat verlassen, Sample-Projekt geklont.
- **Model computed:** eine vollständige Hauptberechnung ist abgeschlossen.
- **Advanced-Method-Läufe:** MICOM, Moderation, PLSc, Copula, IPMA, PLSpredict, FIMIX, HOC, Nomological, CMB, CVPAT. Ein Eintrag pro erfolgreicher Fertigstellung.

Noch nicht geloggt: Bootstrap-Läufe, MGA-Läufe, einzelne Modellspeicherungen, Dataset-Uploads und Dataset-Re-Parsen. Bootstrap und MGA laufen auf Cloud-Run-Services, die ihre Lauf-Dokumente direkt schreiben, ohne das Activity-Log zu berühren, prüfen Sie stattdessen die "Past runs"-Historie des jeweiligen Panels.

Kommentare. Der Activity-Feed enthält auch einen Kommentar-Composer. Kommentare werden als Top-Level-Einträge gespeichert und chronologisch in den Feed eingefügt, sie sind nicht unter einem spezifischen Event verschachtelt. Verwenden Sie Kommentare, um zu annotieren, was ein Kollege bei einem Lauf beachten sollte.

8.5 Ein Projekt verlassen

Editors und Viewers sehen einen **Leave**-Button in ihrer eigenen Zeile der Teamliste. Ein Klick darauf entfernt das Projekt aus Ihrem Arbeitsbereich und gibt den Platz frei. Der Owner kann

nicht verlassen, er muss zuerst das Projekt löschen oder die Eigentümerschaft übertragen.

Die Übertragung der Eigentümerschaft ist in der UI noch nicht verfügbar. Im aktuellen Release ist der Owner die Person, die auf "New project" geklickt hat, und bleibt für die Lebensdauer des Projekts der Owner. Die Übertragung der Eigentümerschaft steht auf der Roadmap, bis dahin kontaktieren Sie den Support, wenn Sie ein Projekt übergeben müssen.

8.6 Was Mitwirkende sehen

Editors erhalten dieselbe UI wie der Owner, minus die Danger-Zone-Aktionen. Viewers sehen:

- Alle Tabs, Dataset, Model, Results, Advanced, Manage
- Alle nur-lesenden Inhalte
- Export-Buttons, Viewers können PDFs, XLSX, JSON herunterladen
- Copy-LaTeX-Symbole auf Tabellen

Viewers können nicht:

- Das Modell oder den Datensatz bearbeiten
- Berechnen oder eine fortgeschrittene Methode auslösen
- Personen einladen oder Rollen ändern
- Einstellungen ändern

Wenn ein Viewer versucht, eine Bearbeitungsaktion anzuklicken, zeigt OpenPLS ein dezentes "permission denied"-Toast an, statt still zu scheitern.

9 Einstellungen und Sprache

Die **Settings**-Unteransicht unter Manage kümmert sich um projektbezogene Metadaten und die Danger-Zone. Der **Sprachumschalter** in der Topbar behandelt die UI-Locale für die aktuelle Sitzung und persistiert sie in Ihrem Nutzerprofil, damit sie Ihnen über verschiedene Geräte hinweg folgt.

9.1 Projekteinstellungen

Öffnen Sie **Manage** → **Settings**, um zum in Abbildung 29 gezeigten Panel zu gelangen.

Zwei Abschnitte:

1. **Project details**: dieselben Namens- und Beschreibungsfelder aus dem New-project-Dialog, jederzeit bearbeitbar. Klicken Sie auf Save, um zu persistieren. Ein Success-Alert bestätigt das Schreiben.
2. **Danger zone**: der Button **Delete project**. Ein Klick öffnet einen Bestätigungsdialog, der den Projektnamen zur doppelten Überprüfung ausschreibt. Die Löschung kaskadiert: das

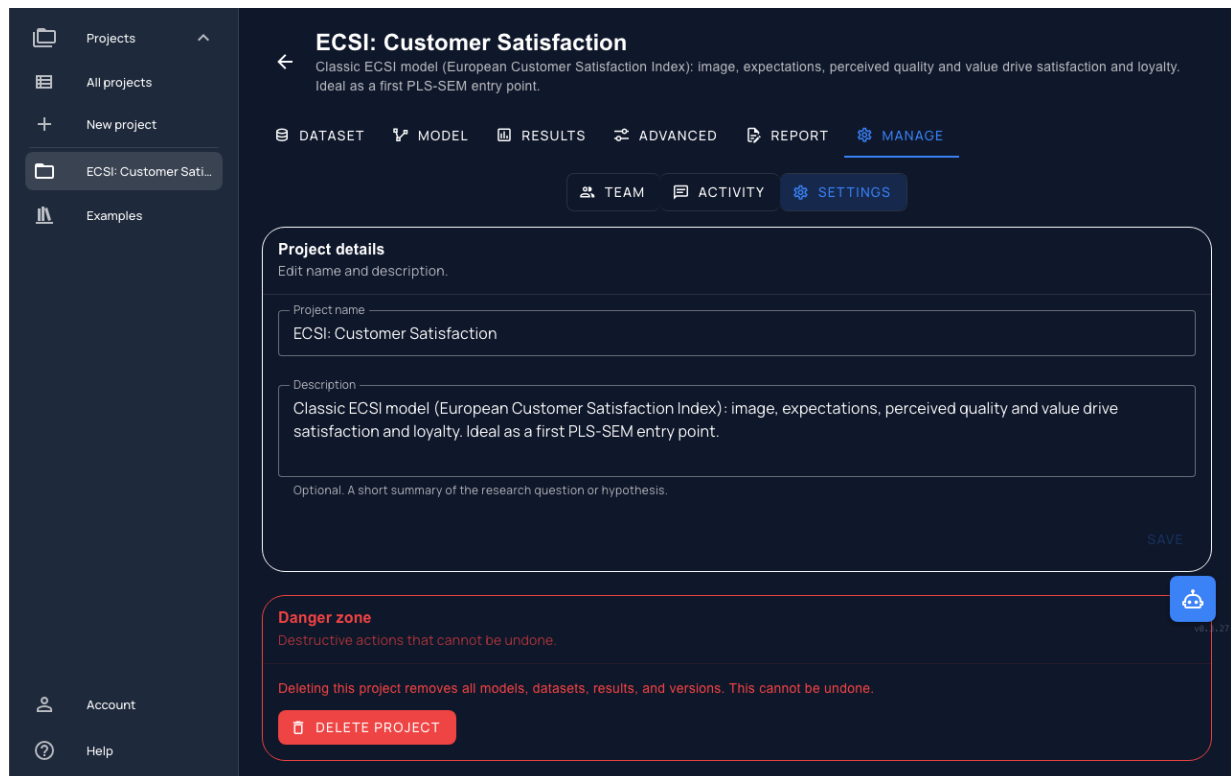


Abbildung 29 Die Settings-Unteransicht: Name- und Beschreibungsformular oben, Danger-Zone unten.

Projekt, seine Datensätze, Modelle, Ergebnisse, Versionen und alle Advanced-Methoden-Läufe werden entfernt. Kann nicht rückgängig gemacht werden.

Nur der Owner sieht die Danger-Zone-Karte. Editors und Viewers sehen nur den Abschnitt Project details.

Das Löschen eines Projekts löscht auch die Einladungshistorie und das Activity-Log. Wenn Sie ein Archiv des Audit-Trails vor der Löschung möchten, exportieren Sie zuerst die Activity-Ansicht (verfügbar über das Download-Symbol in der Activity-Toolbar).

9.2 Sprache

OpenPLS wird mit acht UI-Sprachen ausgeliefert:

- English (default)
- Deutsch
- Español
- Français
- Português (Brasil)
- Simplified Chinese
- Italiano

- Japanese

Klicken Sie in der Topbar auf die Flagge, um das Sprachmenü zu öffnen (Abbildung 30).

Wählen Sie eine Locale. OpenPLS lädt das Message-Bundle (klein, ~30-60 KB) und rendert neu. Die Auswahl wird in Ihrem Nutzerprofil persistiert, beim nächsten Login von jedem Gerät landen Sie in derselben Locale.

Was übersetzt wird.

- Alle UI-Labels, Buttons, Dialoge, Toasts
- Hilfeseiten (/help)
- Report-PDF-Untertitel und Abschnittsüberschriften
- E-Mail-Templates für Einladungen, Passwort-Reset, Verifizierung

Was nicht übersetzt wird.

- Metrikwerte und Spaltennamen, Statistiken bleiben im wissenschaftlichen Standard (lateinbasierte numerische Labels, englische Spaltenüberschriften). Dies ist beabsichtigt, damit Zahlen sauber in englischsprachige Publikationen übernommen werden können.
- Sample-Projekt-Zitationen und lange Beschreibungen, die Wissenschaftszitate sind in der Originalsprache der Publikation.

E-Mail-Templates (Einladung, Passwort-Reset, Verifizierung) sind derzeit nur auf Englisch verfügbar, selbst wenn die Profil-Locale des Empfängers auf eine andere Sprache eingestellt ist. Vollständige E-Mail-Lokalisierung steht auf der Roadmap.

9.3 Kontoeinstellungen

Ihr eigenes Profil liegt unter /account (Link in der Seitenleiste). Bearbeitbar:

- **Display name:** wird im Team-Panel und Activity-Log angezeigt.
- **Avatar:** zum Hochladen klicken; PNG oder JPEG, max. 5 MB, kreisrund zugeschnitten.
- **Preferred locale:** derselbe Effekt wie der Topbar-Umschalter, aber direkt persistiert.
- **Password:** hier ändern. Geben Sie das aktuelle Passwort einmal und das neue Passwort zweimal ein; eine Bestätigungs-E-Mail ist nicht erforderlich.
- **E-Mail-Präferenzen:** ein einzelner Toggle "Receive product updates by email". Transaktionale E-Mails (Einladungen, Passwort-Reset, Verifizierung) werden unabhängig von dieser Einstellung immer gesendet.

Melden Sie sich über die runde Avatar-Schaltfläche oben rechts in der Topbar ab (neben Theme-Toggle und Sprachumschalter).

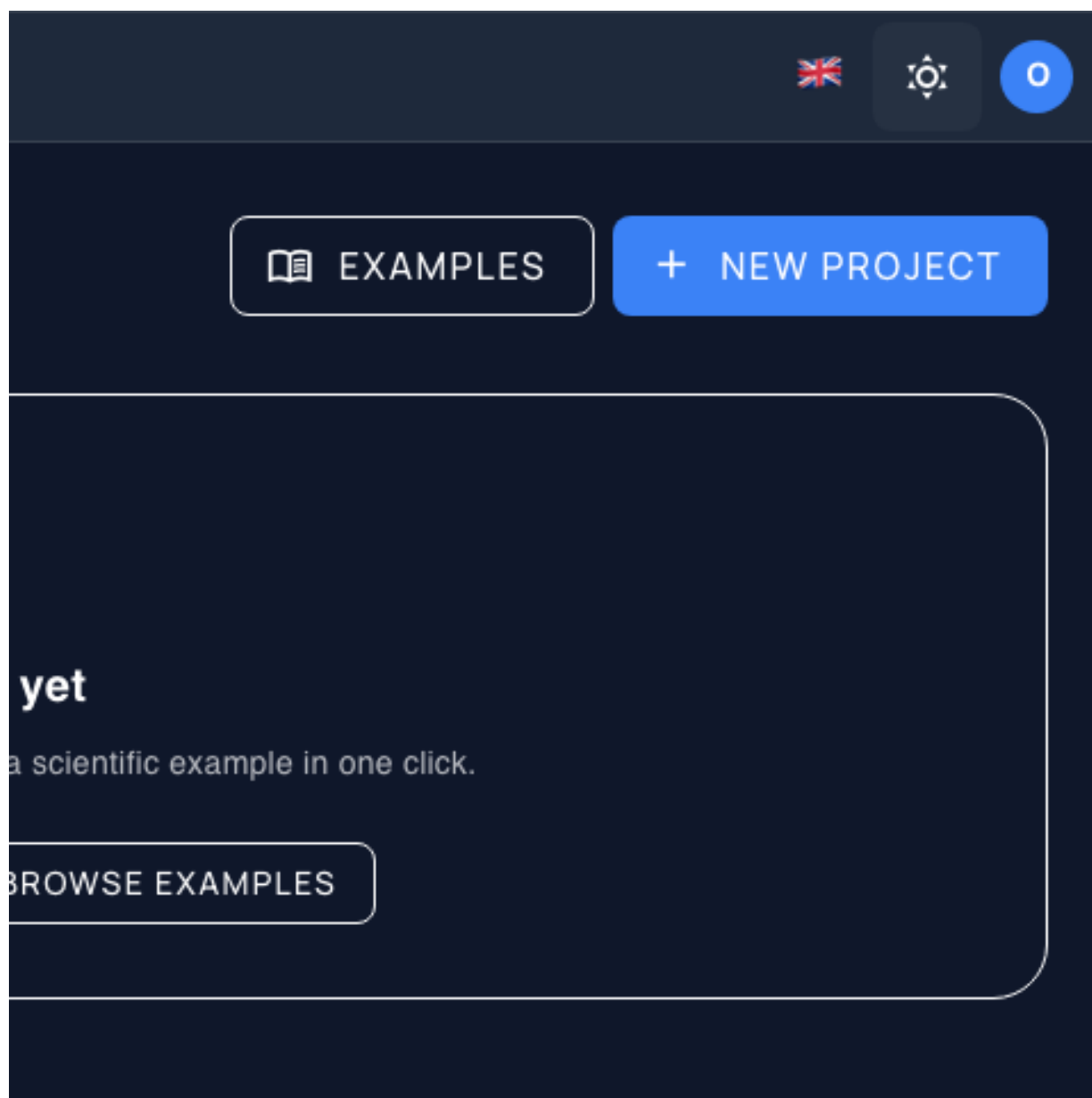


Abbildung 30 Das Sprach-Dropdown aus der Topbar. Die aktuelle Locale ist hervorgehoben.

9.4 Theme

Der Theme-Toggle in der Topbar (Sonne/Mond-Symbol) wechselt für die aktuelle Sitzung zwischen hellem und dunklem Modus. Der Dark-Mode ist Standard. Die Theme-Auswahl wird in `localStorage` gespeichert und synchronisiert daher nicht über Geräte hinweg.

Hilfe und weiterführende Literatur

10.1 In-App-Hilfe

OpenPLS wird mit einem In-App-Hilfe-Center unter `/help` ausgeliefert, zugänglich aus der Seitenleiste jedes Projekts oder über `cloud.openpls.app/help`. Abbildung 31 zeigt die Übersichtsseite.

Sechs Abschnitte:

- **How-to:** aufgabenorientierte Schritt-für-Schritt-Anleitungen, die die Kapitel dieses Leitfadens widerspiegeln, aber häppchenweise. Schnellster Weg, um “wie mache ich ...” nachzuschlagen, ohne ein ganzes PDF zu laden.
- **Metrics:** jede Metrik, die OpenPLS ausweist, mit Formel, Interpretationsschwellen und dem Paper, das sie definiert hat. Wenn Sie sich nicht sicher sind, warum R^2 adjustiert von R^2 abweicht, beginnen Sie hier.
- **Choose:** Entscheidungsbäume für häufige Fragen: “Wann ist mein Konstrukt reflektiv vs. formativ?”, “Welche fortgeschrittene Methode passt zu meiner Hypothese?”, “Wie groß sollte der Bootstrap sein?”.
- **Advanced:** eine Seite pro fortgeschrittener Methode mit demselben Inhalt wie Kapitel 6 dieses Leitfadens, plus methoden-spezifischen Beispielen und Interpretationsdurchläufen. Deep-Links aus der App: Ein Klick auf das kleine `?`-Symbol in einem Advanced-Panel springt zum passenden Help-Seiten-Anker.
- **Pitfalls:** die immer wieder auftretenden Modellierungsfehler, die wir in Support-Tickets sehen: Sentinelwerte, die valide Beobachtungen verschlingen, falscher Messmodus, Vergleich unvergleicher Gruppen. Lesen Sie diese vor Ihrem ersten ernsthaften Projekt.

Die Hilfeseiten sind in alle acht UI-Sprachen lokalisiert.

10.2 Die OpenPLS-Engine (Open Source)

Der Schätz-Code, der OpenPLS antreibt, ist ein Open-Source-Python-Paket namens **openpls-engine**, veröffentlicht unter GPL-3.0 auf <https://github.com/jojacobsen/openpls-engine>. Was Sie in der Cloud berechnen, können Sie auch lokal ausführen mit:

```
pip install openpls-engine
```

Das Paket enthält:

- den zentralen PLS-SEM-Schätzer (entspricht SmartPLS-4-Ausgaben auf 4 Nachkomma-

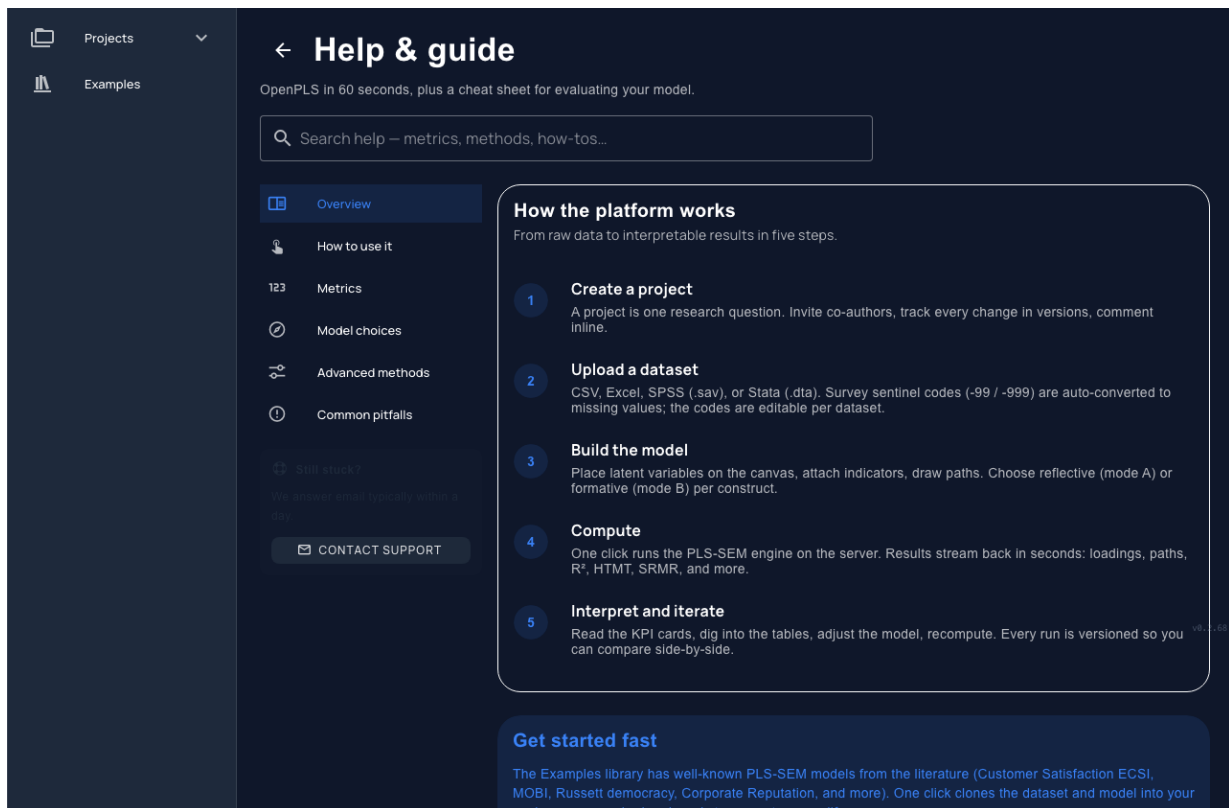


Abbildung 31 Die Help-Übersicht: Einstiegspunkte zu How-to (Bedienung), Metrics, Choose-a-Method, Advanced und Common Pitfalls.

stellen bei validierten Fällen)

- jede in Kapitel 6 beschriebene fortgeschrittene Methode
- CLI + Python-API für reproduzierbare Batch-Läufe
- ein Docker-Image für schlüsselfertiges Self-Hosting

Wenn Sie die gesamte OpenPLS-UI selbst hosten möchten, ist die Engine die Compute-Komponente; siehe die Roadmap im Repo für den aktuellen Stand der Docker-Compose-Orchestrierung.

10.3 Support, Bugs und Feature-Wünsche

Bei allem rund um die Cloud-App (`cloud.openpls.app`) schreiben Sie an hello@openpls.app. Fügen Sie Ihre Projekt-ID bei (in der URL sichtbar) und, wenn Sie einen Bug melden, eine kurze Beschreibung dessen, was Sie erwartet haben, und was passiert ist.

Engine-spezifische Bugs und Feature-Wünsche können auch als GitHub-Issues eingereicht werden unter <https://github.com/jojacobsen/openpls-engine/issues>, oder per E-Mail, wenn Sie das bevorzugen.

10.4 OpenPLS zitieren

Wenn OpenPLS Ergebnisse in Ihrem Paper produziert hat, zitieren Sie bitte:

Jacob, J. (2026). *OpenPLS Engine: An Open-Source Implementation of Partial Least Squares*

Structural Equation Modeling Validated Against SmartPLS 4 Across 14 Reference Cases. SSRN Working Paper. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6869001

Diese Zitation deckt sowohl die Schätzengine als auch die Cloud-Anwendung ab, da sie eine einzige validierte Implementierung teilen.

11 Der AI Companion

OpenPLS wird mit einem **AI Companion** ausgeliefert - einem In-App- Assistenten, der Ihren Projektzustand liest, peer-reviewte Methodikpublikationen zitiert und bei den Teilen von PLS-SEM hilft, die nicht durch Klicken von Buttons erledigt werden können: die Wahl, was als Nächstes ausgeführt werden soll, das Verfassen der Ergebnisdokumentation, das Antizipieren, was ein Reviewer bemängeln könnte.

Der Companion ist **opt-in**. Nichts wird aufgerufen, bis Sie den Einwilligungstext akzeptiert haben, und Sie können ihn jederzeit in den Account-Einstellungen deaktivieren. Dieses Kapitel geht jede Oberfläche in der Reihenfolge durch, in der Sie sie typischerweise antreffen werden.

Jede KI-Antwort ist in zwei Quellen verankert: (1) die tatsächlichen Zahlen in Ihrem Projekt (Datensatzform, Modellspezifikation, aktuelles Compute-Ergebnis), die der Assistent über Nur-Lese-Werkzeuge abruft; und (2) eine kuratierte Bibliothek von Open-Access-PLS-SEM-Publikationen, deren Abstracts per semantischer Ähnlichkeit zu Ihrer Frage abgerufen werden. Der Assistent hat keinen Schreibzugriff auf Ihr Projekt - jede angewandte Empfehlung ist ein Klick, den Sie selbst ausführen.

11.1 Aktivieren

Erstnutzer sehen nach dem Login einen Einstufigen Opt-in-Dialog. Er fasst zusammen, was der Companion liest (Ihre Projektdaten), was er weiterleitet, um verarbeitet zu werden, und das Nutzungsbudget (ein weicher Cap von 100 Turns pro Monat pro Nutzer verhindert außer Kontrolle geratene Nutzung).

Wenn Sie den Dialog früher weggeklickt haben, öffnen Sie **Account** → **AI Companion** und aktivieren dort. Dieselbe Seite trägt auch einen Sub-Toggle für **AI hints** (siehe unten): mit dem Companion an, aber Hints aus, bleiben nur proaktive Features still - Chat, Model-Designer und

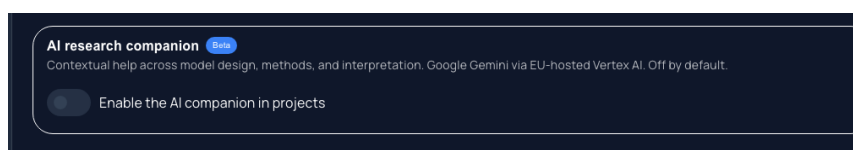


Abbildung 32 Der AI-Companion-Consent-Dialog. Lesen Sie die vier Bulletpoints und klicken Sie dann auf **Enable AI Companion**. Wenn Sie überspringen, bleibt er aus; Sie können ihn später in den Account-Einstellungen aktivieren.

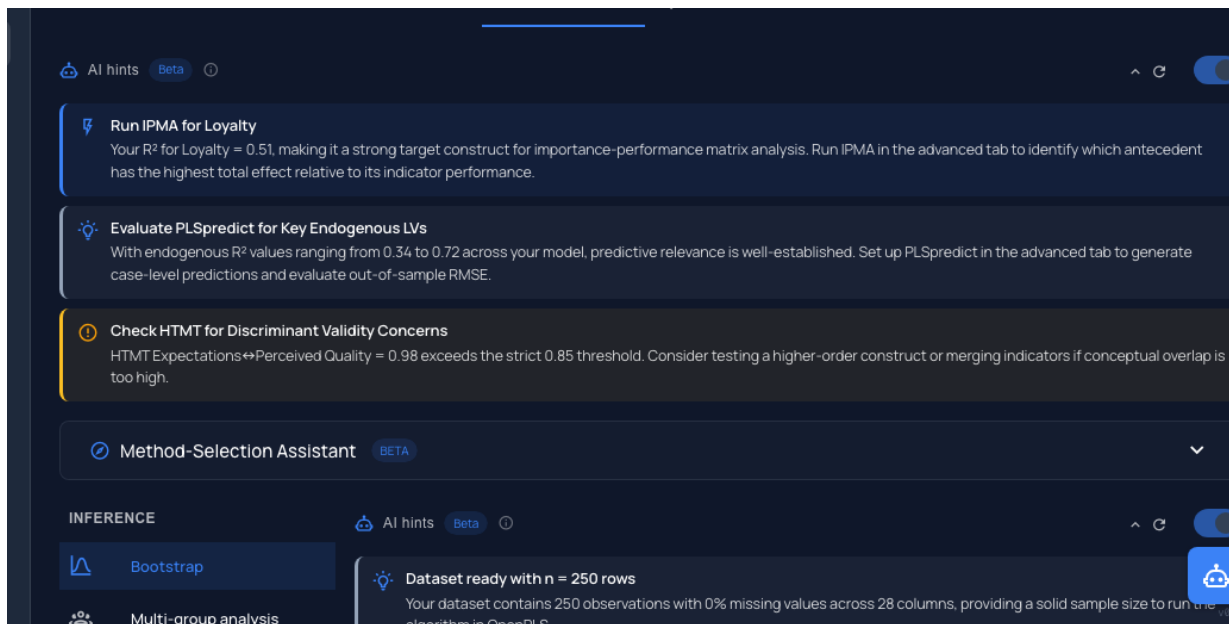


Abbildung 33 Der AI-hints-Streifen im Advanced-Tab. Header trägt einen **Beta**-Chip, einen Info-Tooltip, Refresh, Collapse und den On/Off-Toggle. Jeder Hint hat ein Schweregrad-Icon: Glühbirne für Info, Alert-Circle für Warnung, Blitz für Aktion.

Report-Drafting funktionieren weiter, wenn Sie sie aufrufen.

11.2 AI hints: der stets sichtbare Streifen

Jeder Tab, der einen plausiblen nächsten Schritt hat, zeigt oben einen kleinen **AI hints**-Streifen: ein bis drei Vorschläge, kalibriert auf das, was Sie in diesem Projekt gerade haben. Auf dem Dataset-Tab könnte er eine Spalte mit 22 % fehlenden Werten kennzeichnen; in Results könnte er bemerken, dass Ihr R^2 für Loyalty 0.42 beträgt, und vorschlagen, IPMA auszuführen; in Advanced weist er auf die Methode hin, die Ihr aktueller Zustand am überzeugendsten macht.

Hints werden einmal pro Tab pro Sitzung geladen und dann gecacht. Klicken Sie auf Refresh, um ein erneutes Lesen zu erzwingen, wenn Sie gerade einen neuen Datensatz importiert oder das Modell bearbeitet haben.

Hints einzeln ausschalten. Nicht jeder Nutzer möchte den Streifen im Blick haben. Der Toggle rechts im Streifen schaltet ein nutzerbezogenes `aiHintsEnabled`-Flag um: Hints hören auf zu laden, aber der Rest des Companion (Chat, Designer, Report) funktioniert weiter, wenn Sie ihn öffnen. Off → der Streifen klappt in einen gedämpften Platzhalter ein, der aus dem Weg bleibt.

11.3 Chat: das schwebende Drawer

Ein kleiner Roboter-Button lebt dauerhaft in der unteren rechten Ecke. Klicken Sie ihn, um das Chat-Drawer zu öffnen.

Der Empty-State bietet drei vorbereitete Prompts:

1. **Explain my results in plain language.** Liest die aktuelle Compute-Ausgabe und schreibt

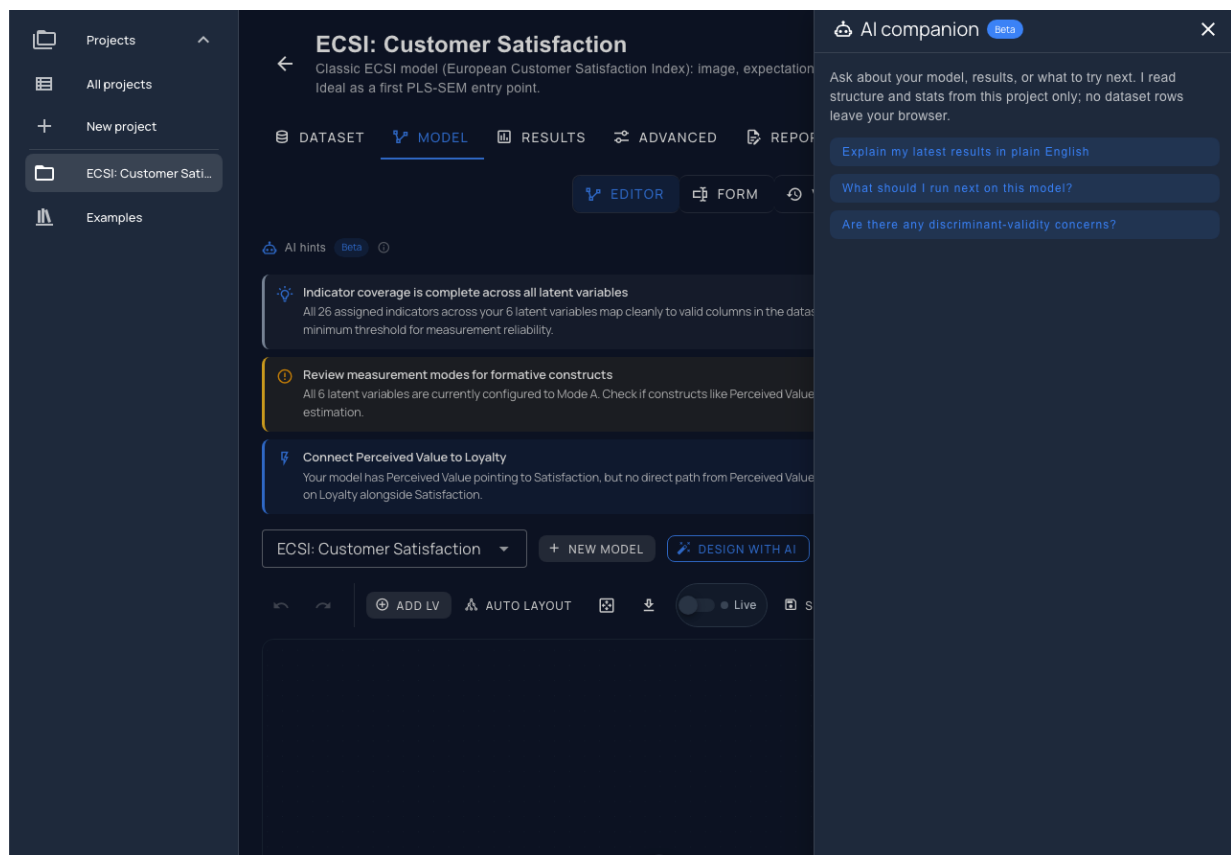


Abbildung 34 Das AI-Companion-Drawer, geöffnet über dem Results-Tab. Der Header zeigt den Titel und einen Beta-Chip, die Vorschlags-Kacheln für Erstnutzer bieten drei vorbereitete Fragen, und der Composer unten trägt einen Live-Kontingenz-Zähler.

ECSI: Customer Satisfaction
Classic ECSI model (European Customer Satisfaction Index): Image, expectation
Ideal as a first PLS-SEM entry point.

DATASET MODEL RESULTS ADVANCED REPORT

Reading your project...

ECSI: Customer Satisfaction, 08/04/2026, 10:57 AM

GOODNESS-OF-FIT
0.6097

SRMR **0.0768**
d_uls = 2.0691

PATHS
10

INNER SUMMARY PATHS INNER MODEL OUTER MODEL

R²: how much of an endogenous LV's variance the predictors explain. Above 0.10 substantial, above

LV	Type	R²	R² adj.	BIC
Expectations	Endogenous	0.3352	0.3325	-92.02
Image		0.0000	-	-
Loyalty	Endogenous	0.5099	0.5060	-162.74
Perceived Quality	Endogenous	0.7197	0.7186	-307.92
Perceived Value	Endogenous	0.5901	0.5868	-207.39
Satisfaction	Endogenous	0.7073	0.7025	-280.56

AI companion **Beta**

What is my R² for the Loyalty construct and is it substantial?

No compute results are currently available because the model calculation is still running. Once the computation completes, run your query again to inspect the R² value for Loyalty.

✖ used: getResults

Ask something about this project...

93 turns left this month

Abbildung 35 Ein Chat-Turn mit der sichtbaren Used-Tools-Zeile unter der Antwort. Der Companion hat `getResults` aufgerufen, um die tatsächlichen R^2 , β und p-Werte vor der Antwort abzurufen.

eine einabsatzige Zusammenfassung, kalibriert auf Ihre spezifischen Zahlen.

2. **Which advanced method should I run next?** Wendet dieselbe Regel-Engine an wie der Method-Selection-Assistent (siehe unten) und gibt inline eine gerankte Shortlist zurück.
3. **Are my constructs valid?** Geht die Reliability- + HTMT- + AVE-Tabellen durch und benennt jede Schwellenverletzung.

Oder tippen Sie Ihre eigene. Alles über das Projekt - Modelldesign, Datensatzqualität, Methodenwahl, Pfadinterpretation, Reviewer-Pushback - ist im Rahmen. Der Companion hat vier Nur-Lese-Werkzeuge:

- **getProject** - Name, Beschreibung, Mitgliederanzahl
- **getModel** - LVs mit Modes, Indikatoren, Strukturpfade
- **getDataset** - Zeilen- und Spaltenanzahl, Statistiken pro Spalte (Missingness, Min/Max/Mean), **niemals die tatsächlichen Zeilenwerte**
- **getResults** - R^2 pro LV, Pfadkoeffizienten mit p-Werten, HTMT-Matrizen, Reliabilität, Fornell-Larcker-Verdikte, SRMR, GoF

Die Used-Tools-Zeile unter jeder Antwort sagt Ihnen, welche davon aufgerufen wurden. Nichts anderes verlässt das Projekt.

Zitationen. Wenn eine Aussage in der Antwort auf die Methodikliteratur zurückgreift, schreibt der Companion einen eingeklammerten Rollenslug (z. B. [henseler-2015-htmt]). Das Frontend rendert diese als klickbare Pill-Chips, die Autor und Jahr zeigen (z. B. **Henseler+ 2015**). Beim Hovern erscheint der vollständige Titel; ein Klick öffnet die DOI- oder Verlagsseite in einem neuen Tab.

Die Publikationsbibliothek hinter diesen Chips ist eine kuratierte Sammlung von ~100 Open-Access-PLS-SEM-Referenzen (siehe *Die Publikationsbibliothek* unten).

Historie + Persistenz. Chat-Turns werden pro Nutzer in der Account-Settings-Storage gespeichert. Laden Sie die Seite neu, öffnen Sie das Drawer, Ihre Konversation ist immer noch da. Das **Clear conversation**-Symbol im Header löscht sie.

Kontingent. 100 Chat-Turns pro Nutzer und Monat während der Beta. Der Zähler am unteren Rand des Composers zeigt, was übrig ist. Model-Design- und Report-Drafting-Turns teilen sich denselben Topf.

11.4 Ein Modell aus einer Forschungsfrage entwerfen

Wenn Ihr Modell noch nicht existiert - ein frisches Projekt, keine LVs gezeichnet - öffnen Sie den **Editor** und klicken Sie auf **Design with AI**. Der Dialog fragt nach Ihrer Forschungsfrage und optional nach den Konstrukten, die Sie bereits im Sinn haben, und einem Disziplin-Hinweis.

Sie erhalten einen strukturierten Vorschlag zurück:

- **3–6 latente Variablen** mit ihrem Modus (A = reflektiv, B = formativ), einer kurzen Begründung und 3–5 Indikatorhinweisen pro LV (Phrasen, die beschreiben, was Items dieses Konstrukt messen würden - nicht die vollständige Fragebogenwording).
- **5–12 Strukturpfade** jeweils mit einer verständlichen Hypothese und einer kurzen Literaturnotiz.
- **2–3 Blueprint-Modelle** aus publizierter PLS-SEM-Literatur, die die Struktur teilen.
- **Notes** zu Caveats oder alternativen Spezifikationen zur Überlegung.

Der Vorschlag wird **nicht** in Ihr Projekt geschrieben, bis Sie **Apply** klicken. Sehen Sie sich eine Vorschau an, verfeinern Sie Ihren Prompt bei Bedarf (**Refine** generiert erneut) oder verwerfen Sie ihn und starten Sie von vorn.

Sobald angewendet, erscheint das Modell im Editor als normale draggable Leinwand - Sie können LVs umbenennen, Pfade hinzufügen, Indikatoren neu verteilen. Der Vorschlag ist ein Startpunkt, kein finales Artefakt.

11.5 Einen Fragebogen entwerfen

Sobald Sie ein Modell haben (LVs mit zugewiesenen Modi), kann OpenPLS eine vollständige Item-Batterie generieren: 4–6 Fragebogen-Items pro LV mit Wortlaut, Likert-Skala, Reverse-Coding-Flags und Zitationen zu validierten Skalen, wenn die LV auf ein bekanntes Konstrukt in der Literatur abbildet (Trust, Satisfaction, Loyalty, TAMs Perceived Usefulness und so weiter).

ECSI: Customer Satisfaction
Classic ECSI model (European Customer Satisfaction Index): Image, expectation
Ideal as a first PLS-SEM entry point.

DATASET MODEL RESULTS ADVANCED REPORT

EDITOR FORM

ECSI: Customer Satisfaction + NEW MODEL DESIGN WITH

Model name: ECSI: Customer Satisfaction Dataset: ECSI: Customer Satisfaction

Latent variables
Each LV has a name, a measurement mode (reflective/formative), and indicators from

Name	Mode	Indicators
Image	A, reflective	imag1 imag2 imag3
Expectations	A, reflective	expe1 expe2 expe3
Perceived Quality	A, reflective	qual1 qual2 qual3
Perceived Value	A, reflective	val1 val2 val3
Satisfaction	A, reflective	sat1 sat2 sat3
Loyalty	A, reflective	loy1 loy2 loy3

+ ADD LV

Structural model (paths)
Each path goes from one LV (source) to another (target).

From	To
Image	Expectations
Image	Satisfaction
Image	Loyalty
Expectations	Perceived Quality
Expectations	Perceived Value
Expectations	Satisfaction
Perceived Quality	Perceived Value
Perceived Quality	Satisfaction
Perceived Value	Satisfaction
Satisfaction	Loyalty

+ ADD PATH

AI companion Beta

What is my R² for the Loyalty construct and is it substantial?

No compute results are currently available because the model calculation is still running. Once the computation completes, run your query again to inspect the R² value for Loyalty.

How should I address common method bias for a reviewer? Suggest citations.

Common method bias (CMB) often arises when data for both independent and dependent variables are collected from the same source using a single measurement instrument (Podsakoff+ 2003). Because self-reported cross-sectional data can inflate observed relationships, reviewers frequently request procedural or statistical remedies (Fuller+ 2016).

To address this in your project, consider the following steps directly in OpenPLS or your broader workflow:

- Procedural remedies:** Detail the design steps taken during data collection to separate measurement contexts, such as assuring respondent anonymity, separating predictor and criterion items temporally or psychologically, and improving scale clarity.
- Statistical assessment:** Run a full collinearity test based on variance inflation factors (VIFs) to check if the data are contaminated by common method bias (Kock+ 2021). A VIF above 3.3 for any lateral collinearity check can indicate potential distortion.

Ask something about this project...

92 turns left this month

Abbildung 36 Chip-Zitationen inline in einer Chat-Antwort gerendert. Jede Pille verlinkt auf die DOI-Seite der Originalpublikation.

AI model designer Beta

Describe your research question. The AI proposes a starter model with latent variables, hypothesised paths and blueprint references from the PLS-SEM literature.

Research question:

55 / 2000

Constructs you already have in mind (optional)

One per line or comma-separated. Leave empty to let the AI decide.

Discipline (optional) ▼ Model name (optional)

CANCEL DESIGN MODEL

Abbildung 37 Der AI-Model-Designer-Dialog. Geben Sie eine Forschungsfrage ein (erforderlich), listen Sie bereits vorhandene Konstrukte auf (optional), wählen Sie eine Disziplin. Klicken Sie auf **Submit**, um einen strukturierten Vorschlag zu erhalten.

AI model designer Beta

Online Grocery Loyalty Model

A PLS-SEM structural and measurement model investigating the drivers of customer loyalty in online grocery shopping, incorporating service quality, perceived value, trust, and satisfaction.

Latent variables (5)

ServiceQuality Reflective exogenous

Reflective measurement of overall service delivery standards perceived by the e-grocery consumer.

Suggested items:

- on-time delivery
- product freshness accuracy
- order fulfillment correctness
- responsive customer support

PerceivedValue Reflective mediator

Reflective indicator set capturing the trade-off between perceived benefits and monetary/time costs.

Suggested items:

- worth for money spent
- fair pricing compared to alternatives
- reasonable delivery fees

Trust Reflective mediator

Reflects consumer confidence in the online grocer's reliability.

Satisfaction Reflective mediator

Reflective measure of cumulative contentment with past

← CHANGE INPUTS DISCARD ✓ APPLY AS NEW MODEL

Abbildung 38 Die Vorschlags-Karte. Jede LV zeigt Modus, Indikatorhinweise und Begründung; Pfade listen Quelle → Ziel mit Hypothese. **Apply** schreibt das Modell als neue Version in Ihr Projekt.

Abbildung 39 Der AI-Fragebogen-Designer-Dialog. Wählen Sie eine Antwortskala (5- oder 7-Punkt-Likert), Zielgruppensprache und klicken Sie auf **Submit**. Der Dialog zeigt, für welche LVs er Items bauen wird; solche, die bereits Items haben, werden übersprungen, es sei denn, Sie kreuzen **Overwrite existing** an.

Öffnen Sie den **Model**-Tab und dann die **Form**-Unteransicht. Wenn irgendeine LV keine Items hat, sehen Sie einen **Design items with AI**-Button. Klicken Sie ihn, um den Fragebogen-Designer-Dialog zu öffnen.

Das Ergebnis ist eine Tabelle von Items pro LV: Item-ID, Wortlaut, Reverse-Coded-Flag, Quellenangabe. Überfliegen Sie sie, passen Sie den Wortlaut an, entfernen oder fügen Sie Items hinzu, und wenden Sie sie dann an - die Items landen als Indikatoren auf jeder LV.

Von hier haben Sie zwei Optionen, um tatsächlich Daten zu sammeln: Bauen Sie einen Fragebogen, den Respondenten unter einer öffentlichen URL ausfüllen (siehe das nächste Kapitel, *Öffentliche Umfragen*) oder nehmen Sie die Item-Batterie in ein externes Umfragewerkzeug und importieren die Antworten als CSV.

11.6 Method-Selection-Assistant

Sobald Sie Compute-Ergebnisse haben, zeigt der Advanced-Tab oben ein **Method-Selection-Assistant**-Panel. Klicken Sie es an, um es auszuklappen.

Der Assistant betrachtet Ihren Projektzustand - Stichprobengröße, Gruppierungskandidaten im Datensatz, R^2 pro endogener LV, HTMT-Verletzungen, Reliabilitätsschwellen - und gibt eine gerankte Shortlist der fortgeschrittenen Methoden zurück, die als Ihr nächster Schritt sinnvoll sind:

- **Bootstrap** (High) wann immer Sie ein Compute-Ergebnis, aber noch keine Konfidenzintervalle haben.

ECSI: Customer Satisfaction

Classic ECSI model (European Customer Satisfaction Index): image, expectations, perceived quality and value drive satisfaction and loyalty. Ideal as a first PLS-SEM entry point.

[DATASET](#)
[MODEL](#)
[RESULTS](#)
[ADVANCED](#)
[REPORT](#)
[MANAGE](#)

[DATA UPLOAD](#)
[QUESTIONNAIRE BUILDER](#)

ECSI: Customer Satisfaction EXPORT REGENERATE

[Share as public survey](#)
Not published

Publish the questionnaire to a public URL so anyone can fill it in. No account needed for respondents. Anonymous responses stream back into the builder — download as CSV any time.

[PUBLISH](#)

[English](#)
[5-point Likert](#)
[Aug 4, 2026, 10:51 AM](#)

[Public intro](#)
Respondents see this

Shown at the top of the public survey. Explain the purpose, expected duration, and how the data will be used.

Thank you for taking part in this short survey. It should take about 5 minutes and helps us understand...

[Internal notes](#)
Team only

For your research team only. Never shown to respondents. Use for design decisions, exclusion criteria, sample size targets.

Ensure ethical review and IRB approval are obtained prior to fielding the survey. Conduct a pilot test with a minimum of 30 respondents to check indicator reliability and multicollinearity before final data collection. Data can be exported and analyzed seamlessly using OpenPLS.

[Demographic questions](#)

Optional standard fields that get exported alongside your items.

[Age](#)
[Gender](#)
[Country](#)
[Education](#)
[Profession](#)
[Language](#)

[Image](#)
Reflective
5-point

[Source: Martenson \(2007\), Journal of Retailing and Consumer Services](#)

Adapted from the European Customer Satisfaction Index (ECSI) framework to measure overall company image.

imag1	The organization has a very positive reputation in the market.	1	X
imag2	I consider this organization to be stable and trustworthy.	1	X
imag3	This organization stands for quality and customer focus.	1	X
imag4	The organization contributes positively to the community.	1	X
imag5	Overall, this organization has a poor public image.	1	X

[+ ADD ITEM](#)
Anchors: Strongly disagree → Strongly agree

[Expectations](#)
Reflective
5-point

[Source: Oliver \(1997\), McGraw-Hill](#)

Measures customer expectations prior to consumption or service delivery, following standard satisfaction models.

Abbildung 40 Die generierte Fragebogen-Karte. Jeder LV-Block listet Items mit Wortlaut + Reverse-Chip + Zitation. Bearbeiten Sie in-place und klicken Sie dann **Apply**, um sie in Ihr Modell zu schreiben.

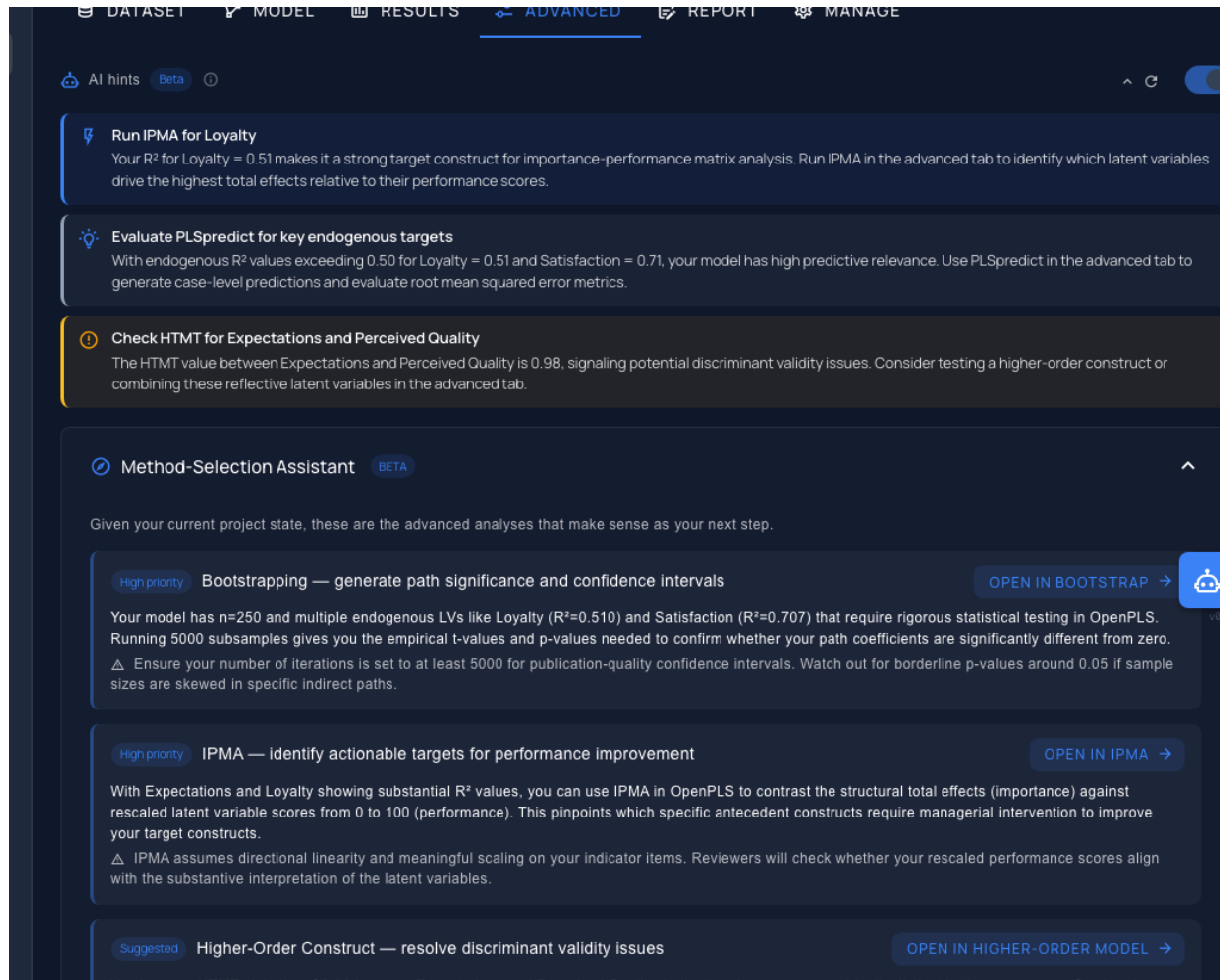


Abbildung 41 Der Method-Selection-Assistant ausgeklappt. Priority-Chips (High / Suggested / Optional), eine Ein-Satz-Begründung pro Methode, eine Risikoonotiz und ein “Open in ...”-Button, der zum entsprechenden Advanced-Untertab navigiert.

- **MGA + MICOM** (High) wenn der Datensatz eine kategoriale Spalte mit 2–5 Levels und genug Stichprobengröße pro Gruppe hat.
- **IPMA** (High) wenn eine Ziel-LV $R^2 \geq 0.33$ hat - genügend substanzielle Varianz, um das Antezedens-Ranking handlungsleitend zu machen.
- **PLSpredict** (Suggested) wenn mindestens eine endogene LV $R^2 \geq 0.19$ und $n \geq 100$ hat.
- **FIMIX-PLS** (Suggested) für $n \geq 200$ als Heterogenitätsprüfung.
- **HOC** (Suggested) wenn zwei LVs HTMT nahe 1 haben - sie könnten First-Order-Facetten desselben Second-Order-Konstrukts sein.
- **PLSc** (Optional) wenn alle LVs reflektiv sind - eine Common-Factor-Robustheitsprüfung.
- **Copula / CMB / Nomological / CVPAT** (Optional) als Standard-Reviewer-Präemptions-Diagnostiken.

Jeder Vorschlag trägt die in Ihren tatsächlichen Zahlen verankerte Begründung (z. B. “Target LV Loyalty has $R^2 = 0.42$ - IPMA identifies which antecedents matter most for practical action”). Der **Open in ...**-Button (beschriftet mit dem spezifischen Methodennamen) leitet zum passenden Advanced-Untertab, damit Sie konfigurieren und ausführen können.

Das Panel ist Fetch-on-Click; nichts wird ausgeführt, bis Sie es öffnen. Klicken Sie auf **Regenerate suggestions**, um Ihren Projektzustand erneut zu lesen.

11.7 Report-Drafting

Ein neuer Tab - **Report** - sitzt zwischen Advanced und Manage. Seine vier Sub-Buttons produzieren veröffentlichungsreife Texte, gezogen aus Ihren tatsächlichen Daten.

11.7.1 Methods / Results / Discussion

Wählen Sie einen Abschnitt und klicken Sie auf **Generate section**. Sie erhalten:

- **Markdown** auf der linken Seite, bereit zum Kopieren in einen Entwurf oder Blog-Post. Verwendet Unicode-Statistiknotation (R^2 , β , χ^2 , f^2) - kein LaTeX innerhalb der Markdown-Ausgabe.
- **LaTeX** auf der rechten Seite, bereit für die Paper-Quelle. Verwendet R^2 , β , χ^2 , f^2 Konventionen für ein passendes Journal-Template.
- **Zitations-Slots** ganz rechts: eine Zeile pro eingeklammerter Rolle, die der Abschnitt referenziert, jede verlinkt auf die DOI- oder Verlagsseite.
- **Notes** - ein oder zwei Sätze Caveat, den der Drafter denkt, dass der menschliche Autor wissen sollte (z. B. “your SRMR of 0.11 is above the 0.08 threshold - consider commenting on fit”).

Der Drafter folgt den Konventionen von Hair/Ringle/Sarstedt:

- **Methods** berichtet Stichprobe, Messmodell, Schätzoptionen, Reliabilitäts- + Validitätskriterien, Model-Fit-Indikatoren und kündigt fortgeschrittene Verfahren an.
- **Results** berichtet Messmodell-Qualität (α , ρ_c , AVE), Diskriminanzvalidität (HTMT mit be-

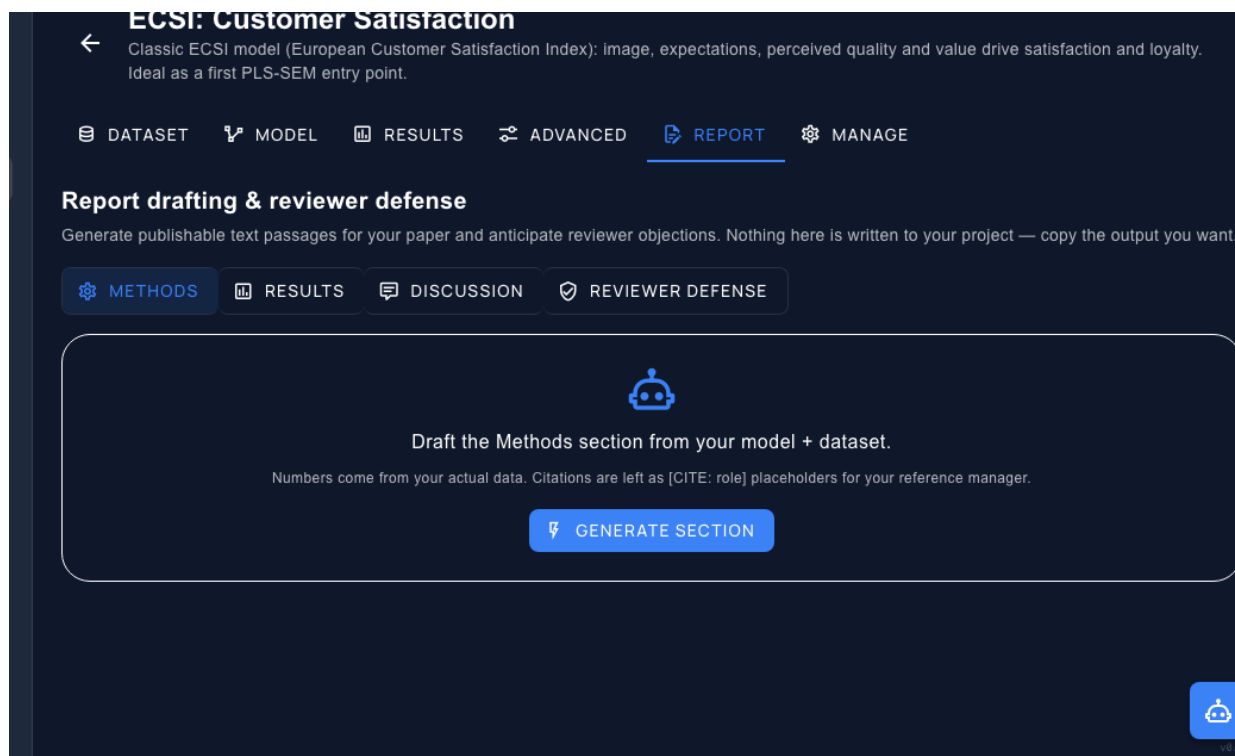


Abbildung 42 Der Report-Tab-Überblick. Vier Sub-Buttons - Methods, Results, Discussion, Reviewer defense - und eine große Karte, die den entsprechenden Drafter öffnet.

nannten Verletzungen), Pfadkoeffizienten mit Signifikanz, R^2 pro endogener LV mit dem verbalen Deskriptor nach Chin 1998, und SRMR.

- **Discussion** interpretiert signifikante Pfade substanziell, adressiert nicht signifikante, spellt praktische Implikationen für die Ziel-LV mit dem höchsten R^2 aus, listet Limitationen und schlägt zwei konkrete Erweiterungen vor.

Kopieren Sie einen der beiden Bereiche mit dem Copy-Icon-Button in seinem Header. Generieren Sie neu mit dem **Regenerate**-Button im Abschnittsheader - der Drafter liest die aktuellsten Ergebnisse erneut, nützlich nach einem frischen Bootstrap oder einer Modellanpassung.

11.7.2 Reviewer defense

Der vierte Sub-Button - **Reviewer defense** - ist der adversarielle Modus des Drafters. Klicken Sie auf **Generate reviewer defense**.

Sie erhalten vier bis sieben antizipierte Einwände, die ein Peer-Reviewer wahrscheinlich vorbringen wird, jeweils mit:

- **Concern** in ein oder zwei Sätzen unter Verwendung der exakten Zahlen aus Ihrem Snapshot.
- **Reviewer persona** - eine kurze Charakterisierung ("CB-SEM traditionalist", "Statistical purist", "Senior methods editor"), damit Sie wissen, welche Rahmung zu erwarten ist.
- **Severity** - Critical, Moderate oder Nitpick.
- **Defense** - das tatsächliche Gegenargument, mit Bezug auf Ihre Zahlen und inline-Zitationen

ECSI: Customer Satisfaction

Classic ECSI model (European Customer Satisfaction Index): image, expectations, perceived quality and value drive satisfaction and loyalty. Ideal as a first PLS-SEM entry point.

DATASET
 MODEL
 RESULTS
 ADVANCED
 REPORT
 MANAGE

Report drafting & reviewer defense

Generate publishable text passages for your paper and anticipate reviewer objections. Nothing here is written to your project — copy the output you want!

METHODS
 RESULTS
 DISCUSSION
 REVIEWER DEFENSE

Methods n = 250

Markdown

To evaluate the European Customer Satisfaction Index (ECSI) model, partial least squares structural equation modeling (PLS-SEM) was conducted using OpenPLS. The sample comprised 250 respondents, with Likert-type items serving as the measurement scale. The structural model consisted of six latent variables estimated using mode A: Image (5 indicators), Expectations (5 indicators), Perceived Quality (5 indicators), Perceived Value (4 indicators), Satisfaction (4 indicators), and Loyalty (4 indicators). All indicators were standardised prior to estimation, and the centroid weighting scheme was applied [Hair+ 2019](#). Statistical inference was performed using bootstrapping with 5,000 resamples to generate robust standard errors and confidence intervals [Streuken+ 2016](#). Measurement model reliability and convergent validity were assessed via Cronbach's alpha, composite reliability (ρ_c), and average variance extracted (AVE). Discriminant validity was examined using the Heterotrait-Monotrait (HTMT) ratio of correlations, applying the standard threshold of 0.85 [Henseler+ 2015](#). Global model fit was evaluated using the standardized root mean squared residual (SRMR), referencing the conservative threshold of < 0.08 for approximate fit [Henseler+ 2016](#). Furthermore, additional procedures such as full collinearity assessment were utilized to check for common method bias [Kock 2015](#).

LaTeX

```
\section{Methods}
To evaluate the European Customer Satisfaction Index (ECSI) model, partial least squares structural equation modeling (PLS-SEM) was conducted using OpenPLS. The sample comprised 250 respondents, with Likert-type items serving as the measurement scale. The structural model consisted of six latent variables estimated using mode A: Image (5 indicators), Expectations (5 indicators), Perceived Quality (5 indicators), Perceived Value (4 indicators), Satisfaction (4 indicators), and Loyalty (4 indicators). All indicators were standardised prior to estimation, and the centroid weighting scheme was applied \cite{hair-2019-when-to-use}. Statistical inference was performed using bootstrapping with 5,000 resamples to generate robust standard errors and confidence intervals \cite{streuken-2016-bootstrap}. Measurement model reliability and convergent validity were assessed via Cronbach's alpha, composite reliability (\rho_c), and average variance extracted (AVE). Discriminant validity was examined using the Heterotrait-Monotrait (HTMT) ratio of correlations, applying the standard threshold of 0.85 \cite{henseler-2015-hmtt}. Global model fit was evaluated using the standardized root mean squared residual (SRMR), referencing the conservative threshold of $< 0.08$ for approximate fit \cite{henseler-2016-srmr}. Furthermore, additional procedures such as full collinearity assessment were utilized to check for common method bias \cite{kock-2015-cmb}.
```

COPY

Citation slots

[Hair+ 2019](#)

When to use and how to report the results of PLS-SEM

[Streuken+ 2016](#)

Bootstrapping and PLS-SEM: A step-by-step guide to get more out of your bootstrap results

[Henseler+ 2015](#)

A new criterion for assessing discriminant validity in variance-based structural equation modeling

[Henseler+ 2016](#)

Testing for goodness-of-fit in structural equation modeling: an assessment of the standardized root mean squared residual (SRMR)

[Kock 2015](#)

Common method bias in PLS-SEM: A full collinearity assessment approach

Your HTMT values for Expectations-Perceived Quality (0.98), Perceived Quality-Perceived Value (0.88), and Perceived Value-Satisfaction (0.93) exceed the 0.85 threshold, which indicates potential discriminant validity issues that should be addressed in the results.

Abbildung 43 Der generierte Methods-Abschnitt. Markdown-Bereich links, LaTeX-Bereich unten, Zitations-Rail rechts. Jede Zahl im Text lässt sich zum Projekt zurückverfolgen.

←

ECSI: Customer Satisfaction

Classic ECSI model (European Customer Satisfaction Index): image, expectations, perceived quality and value drive satisfaction and loyalty. Ideal as a first PLS-SEM entry point.

DATASET

MODEL

RESULTS

ADVANCED

REPORT

MANAGE

Report drafting & reviewer defense

Generate publishable text passages for your paper and anticipate reviewer objections. Nothing here is written to your project — copy the output you want

METHODS

RESULTS

DISCUSSION

REVIEWER DEFENSE

Overall strategy

Lead the revision by acknowledging the severe HTMT violations and immediately performing item purification or higher-order modeling in OpenPLS. Concurrently, execute the missing advanced procedures—specifically bootstrapping, CMB full collinearity assessment, and PLSpredict—while dropping obsolete metrics like GoF in favor of SRMR and contemporary PLS-SEM standards.

REGENERATE

Critical Measurement purist

The HTMT values between Expectations and Perceived Quality (0.98), Perceived Value and Satisfaction (0.93), and Perceived Quality and Perceived Value (0.88) severely violate the standard 0.85 threshold, indicating severe discriminant validity problems.

Defense:

While these HTMT values exceed the 0.85 threshold, discriminant validity concerns can sometimes arise in closely related consumer-behavior constructs like quality and expectations. To address this prior to resubmission, you should run a discriminant validity remediation or higher-order structural analysis in OpenPLS.

Evidence:

HTMT is 0.98 for Expectations and Perceived Quality, 0.93 for Perceived Value and Satisfaction, and 0.88 for Perceived Quality and Perceived Value.

Mitigation:

Examine item cross-loadings in OpenPLS, consider dropping problem items, or re-specify the model with higher-order constructs.

Critical Statistical reviewer

No bootstrap analysis has been run yet, meaning there are no p-values or confidence intervals to validate the statistical significance of the 10 structural paths.

Defense:

This is simply an artifact of the initial model run snapshot; bootstrapping must be executed to confirm the significance of the 7 significant paths.

Evidence:

Bootstrap is currently set to false in the project snapshot.

Mitigation:

Run bootstrapping with 5,000 resamples in OpenPLS to generate bias-corrected confidence intervals.

Critical Methods editor

Common method bias has not been assessed, yet cross-sectional survey data with 250 respondents is highly susceptible to CMB.

Defense:

CMB must be explicitly evaluated before final publication to ensure the structural paths are not artifactual [Kock 2015](#).

Evidence:

CMB is listed as false in the advanced runs performed.

Mitigation:

Perform a full collinearity assessment in OpenPLS to check if inner VIF values remain below 3.3 [Kock 2015](#).

Suggested citations

Kock 2015

Moderate CB-SEM traditionalist

The project reports the Goodness-of-Fit (GoF) index of 0.610, which is known to be methodologically flawed for PLS-SEM validation.

Defense:

Although GoF is reported in the snapshot, modern PLS-SEM guidelines discourage its use for model validation, favoring SRMR instead [Henseler+ 2013](#).

Evidence:

GoF is 0.610 while SRMR is 0.077.

Mitigation:

Drop GoF from the manuscript reporting and rely solely on SRMR and exact model fit diagnostics [Henseler+ 2013](#).

Suggested citations

Henseler+ 2013

Moderate Senior econometrician

Endogeneity and predictive relevance have not been tested, leaving open the possibility of omitted variable bias and unproven out-of-sample predictive power.

Defense:

The current model establishes foundational explanatory power with an R^2 up to 0.720, which can be complemented by out-of-sample and endogeneity checks.

Evidence:

Copula and PLSpredict are both marked as false in advanced runs.

Mitigation:

Run PLSpredict and Gaussian copula tests in OpenPLS to verify out-of-sample performance and check for endogeneity.

Abbildung 44 Eine Reviewer-Defense-Karte. Severity-Chip (Critical / Moderate / Nitpick), Reviewer-Persona, Concern, Defense, Evidence, Mitigation und suggested-citation-Chips.

64

aus der Publikationsbibliothek.

- **Evidence** - die spezifischen Snapshot-Werte, die die Verteidigung stützen.
- **Mitigation** - 0–2 Sätze dazu, was noch vor der Einreichung ausgeführt werden sollte, falls überhaupt.
- **Suggested citations** - klickbare Chips für die Publikationen, auf die sich die Verteidigung stützt.

Ein kurzer **Overall strategy**-Absatz oben auf der Karte listet die zwei oder drei Züge, die Ihr Revisionsschreiben verankern - meistens “lead with X, concede Y as a limitation, commit to Z in a follow-up”.

Der Defender ist ein Simulator, kein Orakel. Echte Reviewer werden Kontext mitbringen, den Sie nicht vorab kennen können - theoretische Einwände, die mit Ihrer spezifischen Literatur verknüpft sind, methodische Eigenheiten Ihrer Zielzeitschrift, persönliche Steckpferde. Nutzen Sie die Ausgabe als Probe, nicht als Garantie.

11.8 Die Publikationsbibliothek

Jede KI-Oberfläche, die Zitationen produziert, greift auf dieselbe kuratierte Bibliothek zurück - rund 100 Open-Access-PLS-SEM-Publikationen, die Methodik, Kritik, Anwendungen und die Modellierungsklassiker (TAM, UTAUT, TPB, SERVQUAL, ACSI, JD-R) abdecken. Jeder Eintrag ist ein Titel, eine Autorenliste, ein Jahr, eine DOI und das Abstract, wie es auf der Verlagsseite erscheint. Kein Volltext - nur die Metadaten, die Reviewer benötigen, um das Paper selbst zu finden.

Die Bibliothek ist in derselben Datenbank wie Ihre Projekte gespeichert und wird über **semantische Vektorsuche** abgefragt: Wenn Sie eine Frage stellen, bettet der Assistent Ihre Frage mit einem mehrsprachigen Embedding-Modell ein, findet die fünf ähnlichsten Paper-Abstracts per Kosinus-Ähnlichkeit und gibt diese dem LLM als Referenzkontext. Das Modell zitiert dann mit den eingeklammerten Rollenslugs (z. B. [henseler-2015-htmt]), die das Frontend durch Pill-Chips ersetzt.

Wenn ein Chip grau/flach erscheint (nicht klickbar), bedeutet das, dass der Assistent eine Rolle zitiert hat, die noch nicht in der Bibliothek ist - entweder ein älteres Paper außerhalb unseres Kurationsfensters oder gelegentlich eine Halluzination. Graue Chips sind ein Signal, die Referenz manuell doppelt zu prüfen.

11.9 Kontingent und Datenschutz

Kontingent. 100 Turns pro Nutzer und Monat während der Beta. Gilt separat für interpretative Turns (Chat, Model-Designer, Report-Drafter) und Hints (100 Hint-Fetches pro Monat). Turns übertragen sich nicht.

Datenschutz. Der Companion liest Ihre Projektstruktur und Aggregate (Zeilenanzahlen, Spal-

tenstatistiken, R^2 , Pfadkoeffizienten), niemals die tatsächlichen Fragebogenzeilen. Die gesamte Verarbeitung läuft unter vertraglichen Datenschutzbedingungen mit dem Modellanbieter, gehostet in der EU, und Ihre Daten werden nie zum Training zukünftiger Modelle verwendet.

11.10 Ausschalten

Account → **AI Companion** → **Toggle aus** deaktiviert jede KI-Oberfläche in der App: das Drawer verschwindet, der Hints-Streifen klappt ein, die Designer- + Report-Tabs existieren weiter, aber ihre Generate-Buttons werden inaktiv mit einem “not enabled”-Hinweis.

Die gespeicherte Konversation bleibt erhalten (nichts wird gelöscht); aktivieren Sie später erneut und Ihre Chat-Historie ist genau dort, wo Sie sie verlassen haben.

12 Öffentliche Umfragen

OpenPLS kann Ihren Fragebogen auf einer öffentlichen URL hosten, anonyme Antworten sammeln und sie direkt als Datensatz in Ihr Projekt zurückführen. Sie benötigen kein Drittanbieter-Umfragewerkzeug für den normalen PLS-SEM-Workflow.

Die typische Schleife:

1. **Bauen** Sie die Item-Batterie im Model-Tab (von Hand oder mit dem AI-Fragebogen-Designer - siehe das vorherige Kapitel).
2. **Veröffentlichen** Sie den Fragebogen unter einer öffentlichen `/q/{slug}`-URL.
3. **Teilen** Sie die URL mit Respondenten. Neue Antworten landen in nahezu Echtzeit im Builder.
4. **Importieren** Sie die gesammelten Antworten mit einem Klick als Datensatz.
5. **Berechnen** Sie das Modell auf Ihrem frischen Datensatz.

Dieses Kapitel geht jeden Schritt durch.

12.1 Den Fragebogen bauen

Öffnen Sie den **Model**-Tab eines Projekts und dann die **Form**-Unteransicht. Wenn Ihre LVs noch keine Indikatoren haben, startet das Panel leer; verwenden Sie den AI-Fragebogen-Designer (siehe Kapitel 12) oder fügen Sie Items von Hand hinzu. Jedes Item ist eine Spalte in der zukünftigen Antwort-CSV.

Felder pro Item:

- **Wording** - der Satz, den Respondenten sehen. Reverse-codierte Items erhalten einen R-Chip.
- **Item ID** - automatisch generiert (z. B. TRUST1, SAT2), im exportierten Datensatz als Spaltenname verwendet.
- **Reverse coded** - Toggle umlegen, wenn die Polarität des Items invertiert ist (Respondenten negieren gedanklich). Reverse-Codierung wird beim Import automatisch angewandt.

The screenshot shows the 'ECSI: Customer Satisfaction' questionnaire builder interface. At the top, there is a title 'ECSI: Customer Satisfaction' and a description: 'Classic ECSI model (European Customer Satisfaction Index): image, expectations, perceived quality and value drive satisfaction and loyalty. Ideal as a first PLS-SEM entry point.' Below this is a navigation bar with tabs: DATASET, MODEL, RESULTS, ADVANCED, REPORT, and MANAGE. The 'MODEL' tab is active. Under the 'MODEL' tab, there are two buttons: 'DATA UPLOAD' and 'QUESTIONNAIRE BUILDER'. The 'QUESTIONNAIRE BUILDER' button is highlighted. Below the buttons, the title 'ECSI: Customer Satisfaction' is repeated. A text box says: 'Generate 4–6 validated survey items per latent variable (6 LVs in this model). You can then edit anything inline before exporting.' Below this, there is a section 'Building items for:' with six buttons: 'Image · Mode A', 'Expectations · Mode A', 'Perceived Quality · Mode A', 'Perceived Value · Mode A', 'Satisfaction · Mode A', and 'Loyalty · Mode A'. Below these buttons, there are two dropdown menus: 'Response scale' (set to '5-point Likert') and 'Language' (set to 'English'). Below the dropdowns, there is a text box: '5 answer options: "Strongly disagree, Disagree, Neither, Agree, Strongly agree". Faster to answer, less resolution — good for short surveys, mobile audiences.' Below this, there is a large text box: 'Target audience or context (optional)'. Below the text box, there is a note: 'Helps the assistant adapt wording culturally.' Below this, there is a section 'Standard demographics (optional)' with a note: 'Toggle which fields to include alongside your latent-variable items.' Below the note, there are five buttons: '✓ Age', '✓ Gender', '+ Country', '+ Education', '+ Profession', and '+ Language'. At the bottom right, there are two buttons: 'START MANUALLY' and 'GENERATE ITEMS'.

Abbildung 45 Das Fragebogen-Builder-Panel. Links: LVs mit ihren Items (per Drag umsortieren, per × entfernen, Wortlaut zum Bearbeiten anklicken). Rechts: die Publish-Karte mit Antwortskala, Sprache, Zielgruppe und Demografie-Toggles.

- **LV-Bindung** - vom übergeordneten LV-Block geerbt; ziehen Sie Items zwischen Blöcken, um sie neu zuzuweisen.

Globale Felder (rechte Spalte):

- **Response scale** - 5- oder 7-Punkt-Likert. Gilt für alle Likert-artigen Items; offene und demografische Items behalten ihren eigenen Typ.
- **Audience language** - die Sprache, die Respondenten sehen (Englisch, Deutsch, Spanisch, Französisch, Italienisch, Japanisch, Portugiesisch, Chinesisch). Der Wortlaut der Items bleibt in der Sprache, in der Sie ihn geschrieben haben; nur das Umfrage-Chrome (Buttons, Fortschritt, Dankeseite) übersetzt.
- **Demographics** - wählen Sie null oder mehr aus den sechs eingebauten optionalen Feldern (Alter, Geschlecht, Land, Bildung, Sprache, Beruf). Jedes wird eine Spalte im exportierten Datensatz.
- **Public intro** - ein Absatz, den Respondenten über der ersten Frage sehen. Getrennt von Ihren internen Notizen (die niemals auf der öffentlichen Seite erscheinen).

12.2 Veröffentlichen

Der **Publish**-Abschnitt des Panels schaltet die Umfrage live. Klicken Sie auf **Publish**: OpenPLS erzeugt einen zufälligen Slug, schreibt einen öffentlichen Snapshot des Fragebogens in die Response-Collection und gibt Ihnen die URL zurück:

`https://cloud.openpls.app/q/f3a8k29b`

Jeder mit der URL kann antworten. Kein Konto erforderlich - Anonymität ist Standard. Die öffentliche Seite zeigt nur den öffentlichen Intro und die Items selbst, responsiv formatiert für Telefone und Desktops.

Preview öffnet die Umfrage in einem neuen Tab, damit Sie sie als Respondent durchgehen können (Ihre Antworten im Preview-Modus zählen nicht).

12.2.1 Versions-Snapshots

Sobald Antworten hereinkommen, führt eine Änderung des Fragebogens zu einem Datenintegritäts-Dilemma: Wenn Sie einen Item mittendrin umbenennen oder seinen Wortlaut ändern, werden die gesammelten Antworten inkonsistent. OpenPLS behandelt dies, indem es Sie beim Speichern von Änderungen an einer Umfrage mit vorhandenen Antworten bittet, einen Pfad zu wählen:

- **Publish as new version** (empfohlen). Zieht die aktuelle URL zurück, erzeugt einen neuen Slug (/q/...) und setzt den Zähler auf null. Alte Antworten bleiben unter v1 archiviert und stehen zum Download bereit.
- **Download responses, then save in place.** Selbe URL, aber Sie laden vorher eine CSV-Archivkopie des Vor-Änderungs-Batches herunter.
- **Save in place.** Selbe URL; neue und alte Antworten vermischen sich. Nur für Tippfehler-Korrekturen verwenden.

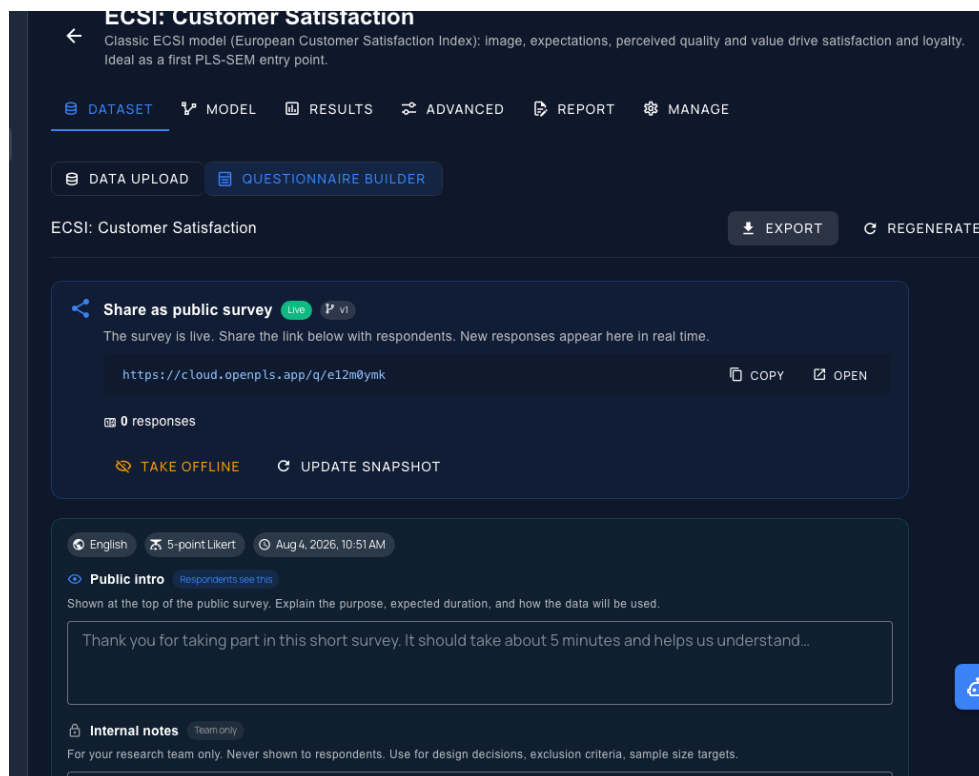


Abbildung 46 Eine veröffentlichte Umfrage-Karte. Die URL wird mit Copy-to-Clipboard angezeigt, ein Live-Response-Counter, drei Aktionen (Preview / Unpublish / Publish as new version) und der Versions-Chip.

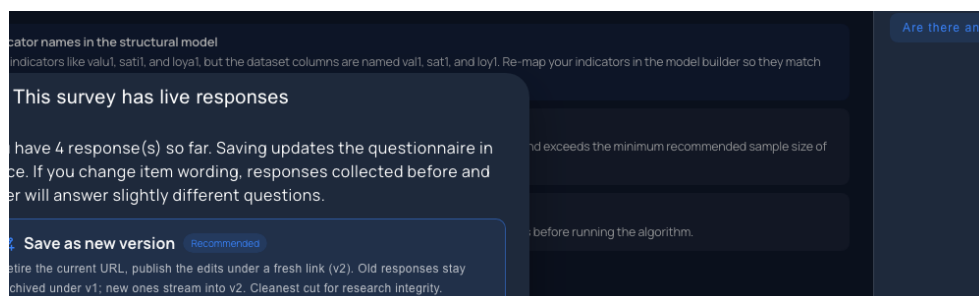


Abbildung 47 Der Save-Warn-Dialog, wenn eine veröffentlichte Umfrage Antworten hat. Drei Pfade, jeder mit der konkreten Konsequenz ausgeschrieben.

Zurückgezogene Slugs bleiben online, akzeptieren aber keine neuen Antworten mehr - Respondenten sehen eine “no longer accepting responses”-Nachricht. Die Publish-Karte listet alle vorherigen Versionen mit ihrer finalen Antwortanzahl, damit nichts verloren geht.

12.3 Die öffentliche Respondentenansicht

Respondenten sehen ein sauberes, einspaltiges Formular mit einer Fortschrittsleiste. Items sind nach LV gruppiert, damit die kognitive Last überschaubar bleibt; der LV-Name selbst ist ausgeblendet (Respondenten benötigen das theoretische Label nicht).

Erwähnenswerte Design-Entscheidungen:

- **Sprachumschalter.** Respondenten können unabhängig von der Audience-Language, die Sie gesetzt haben, durch die acht unterstützten Sprachen wechseln - nützlich für internationale Deployments.
- **One-response-per-browser-Guard.** Sobald abgesendet, verhindert ein kleines Flag im LocalStorage des Browsers versehentliches erneutes Absenden vom selben Gerät. Nicht manipulationssicher (jeder entschlossene Respondent kann es umgehen, indem er den Storage löscht oder einen anderen Browser verwendet), aber es stoppt ehrliche Doppelabsendungen.
- **Dankesseite.** Nach Submit sehen Respondenten eine kurze Bestätigung. Kein Prompt zur Kontoerstellung, kein Marketing.
- **Mobile-first.** Das Formular funktioniert bis hinunter zu ~320 px Viewports; Item-Karten stapeln sich, der Sprachumschalter wird zu einem kompakten Menü.

12.4 Antworten beobachten

Die Publish-Karte im Builder zeigt die aktuelle Antwortanzahl in nahezu Echtzeit (Aktualisierung innerhalb weniger Sekunden nach jeder Einreichung). Unter dem Zähler listet ein kleiner **Recent**-Abschnitt die letzten fünf Einreichungen mit ihrem Zeitstempel und ihrer Sprache.

12.5 Antworten als Datensatz importieren

Wenn Sie genug Antworten haben, um das Modell zu berechnen - für eine typische PLS-SEM-Studie mindestens 100 für explorative Arbeit; 200+ für publikationsqualitative Inferenz - klicken Sie auf **Import as dataset** auf der Publish-Karte.

OpenPLS erstellt einen frischen Datensatz in Ihrem Projekt. Spaltenlayout:

- **item_id-Spalten** - eine pro Fragebogen-Item, unter Verwendung der von Ihnen definierten Item-IDs TRUST1, TRUST2, Der Zellwert ist die Likert-Antwort als Integer (1..5 oder 1..7). Reverse-codierte Items werden beim Import umgedreht, damit nachgelagerter Code sich nicht daran erinnern muss.
- **Demografie-Spalten** - eine pro Demografie-Feld, das Sie aktiviert haben (age, gender, country, ...). String oder Integer je nach Feldtyp.

ECSI: Customer Satisfaction

About you

Age *

Gender *

Image

Rate each statement on a 5-point scale.

The organization has a very positive reputation in the market.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

I consider this organization to be stable and trustworthy.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

This organization stands for quality and customer focus.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

The organization contributes positively to the community.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

Overall, this organization has a poor public image.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

Expectations

Rate each statement on a 5-point scale.

I expected the products and services to be highly reliable.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

I anticipated receiving personalized attention from the staff.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

I expected the overall experience to meet my high standards.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

I anticipated encountering numerous problems with the service.

☐ 1☐ 2☐ 3☐ 4☐ 5

← Strongly disagree Strongly agree →

Abbildung 48 Die öffentliche Umfrage-Seite auf dem Desktop gerendert. Sprachumschalter oben rechts, Fortschrittsleiste, Item-Karten mit Radio-Buttons, Previous- / Next-Navigation und eine finale Submit-Aktion.

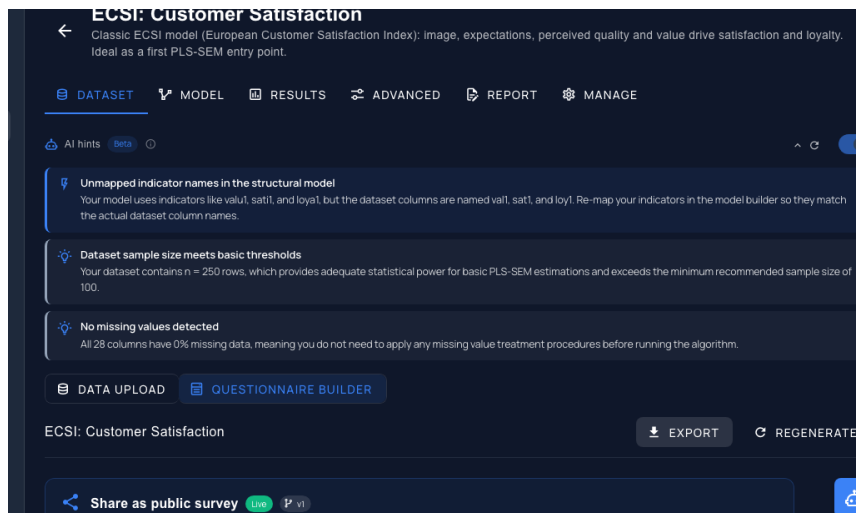


Abbildung 49 Der Response-Counter und die Recent-Liste. Der Zähler aktualisiert sich live; die Recent-Liste zeigt submitted_at plus Sprachflagge.

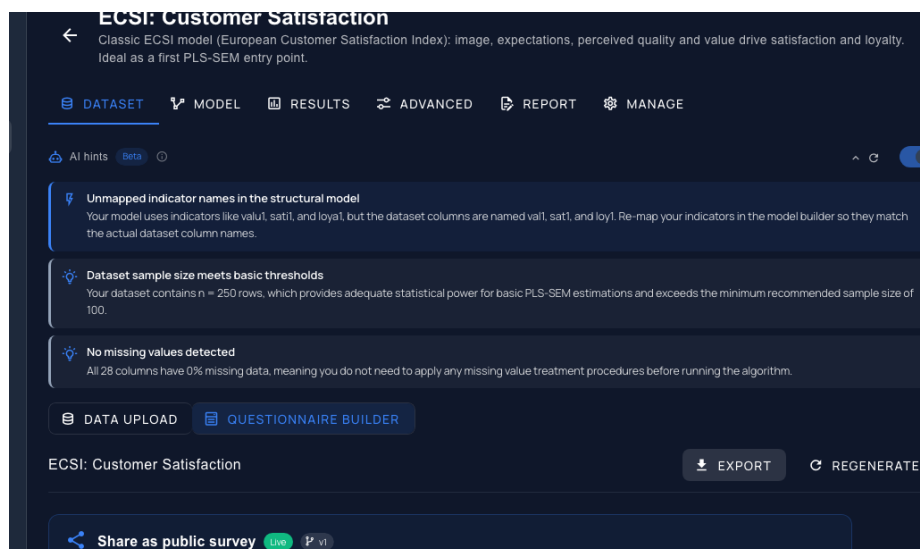


Abbildung 50 Der Import-as-dataset-Button auf der Publish-Karte. Der Bestätigungsdialog zeigt: Item-Spalten, Demografie-Spalten, Stichprobengröße und Version. Klicken Sie auf **Create dataset**, um zu schreiben.

- **submitted_at** - ISO-8601-Zeitstempel, wann die Antwort abgesendet wurde, nützlich für zeitliche Filter oder Wellenanalyse.
- **language** - die Sprache, die der Respondent verwendet hat, nützlich für sprachübergreifende MICOM/MGA-Vergleiche.

Öffnen Sie den Data-Tab; der neue Datensatz ist Ready. Berechnen Sie das Modell wie gewohnt. Wenn Sie später eine neue Version veröffentlichen, schreibt ein erneuter Import einen **separaten** Datensatz - der alte bleibt erhalten, damit Sie Schätzungen vor und nach der Fragebogenänderung vergleichen können.

12.6 Rohe CSV herunterladen

Wenn Sie die Antworten lieber in ein externes Werkzeug bringen möchten (R, Stata, SPSS, Python), klicken Sie auf **Download responses (CSV)** auf der Publish-Karte. Die CSV hat dieselben Spalten wie der importierte Datensatz. Jede Zeile ist eine Einreichung; leere Zellen bedeuten, dass der Respondent das optionale Feld übersprungen hat.

Downloads zählen nur abgeschlossene Einreichungen. Teilweise Antworten (Respondent hat den Tab vor Submit geschlossen) erreichen die Datenbank nicht.

12.7 Veröffentlichung zurücknehmen

Unpublish nimmt die URL offline. Vorhandene Antworten bleiben im Builder zugänglich und können weiterhin heruntergeladen oder importiert werden; neue Besucher sehen eine “not accepting responses”-Seite. Die URL wird nicht freigegeben - ein erneutes Klicken auf Publish auf demselben Fragebogen verwendet nach Möglichkeit denselben Slug erneut.

Vorsicht: Wenn Sie veröffentlichen → sammeln → zurücknehmen → Fragebogen ändern → erneut veröffentlichen → sammeln, haben Sie gemischte Antworten, die an denselben URL-Slug gebunden sind. Der oben beschriebene Save-Warn-Dialog verhindert dies im normalen Bearbeitungsablauf, aber Unpublish gefolgt von Bearbeiten gefolgt von Republish ist eine Umgehung, die ihn umgeht. Bevorzugen Sie **Publish as new version**, wann immer sich der Fragebogen substantiell ändert.

12.8 Limits und Konventionen

- **Antwortlimit pro Umfrage**: kein hartes Limit während der Beta. Die realistische weiche Obergrenze liegt bei einigen Tausend pro Umfrage pro Projekt, bevor der Live-Zähler merklich hinterherhinkt. Für 10 000+ Antworten teilen Sie in mehrere Umfragen auf oder verwenden einen externen Panel-Provider.
- **Antwortretention**: unbefristet, solange das Elternprojekt existiert. Das Löschen des Projekts kaskadiert die Antworten.
- **Anonymität**: gespeichert wird nur, was Respondenten eingeben. IP-Adresse, User-Agent und andere identifizierende Header werden nicht persistiert. Wenn Sie identifizierbare Ant-

worten wünschen (z. B. für eine Panelstudie), fügen Sie selbst ein Respondenten-ID-Feld als Umfrage-Item hinzu.

13 Anhang

13.1 Glossar

AVE: Average Variance Extracted. Mittelwert der quadrierten Indikatorladungen einer reflektiven LV. Über 0.5 = die LV erklärt mehr Varianz in ihren Indikatoren als der Fehler.

BIC: Bayesian Information Criterion. Modellfit-Index, der zusätzliche Parameter bestraft; niedriger ist besser; nur zwischen genesteten Spezifikationen vergleichbar.

Bootstrap: Resampling mit Zurücklegen zur Ableitung von Standardfehlern, p-Werten und Konfidenzintervallen für Statistiken, denen eine analytische Verteilung fehlt.

Composite Reliability (ρ_c , $\rho_{\theta c}$): der moderne Standard für Reliabilität in PLS-SEM. Über 0.7 akzeptabel, über 0.95 deutet auf redundante Indikatoren hin.

CMB: Common Method Bias. Systematischer Fehler, der durch das Messinstrument statt durch die gemessenen Konstrukte induziert wird. Diagnostiziert über den Lindell-Whitney-Marker oder Harmans Einzelfaktor-Test.

CVPAT: Cross-Validated Predictive Ability Test. Vergleicht den Out-of-Sample-Verlust von PLS mit einer Benchmark (Indikator-Mittelwert oder lineares Modell).

Endogenität: Ein Prädiktor ist mit dem Fehlerterm korreliert. Verletzt OLS-Annahmen; in PLS über den Gaussian-Copula-Ansatz korrigiert.

Endogene LV: Ein Konstrukt mit mindestens einem eingehenden Pfad. Im Gegensatz zu exogen.

Exogene LV: Ein Konstrukt mit nur ausgehenden Pfaden. Seine Varianz wird nicht durch das Modell erklärt.

FIMIX-PLS: Finite-Mixture-PLS. Entdeckt latente Klassen, wenn Sie unbeobachtete Heterogenität vermuten.

Formativ (Mode B), Messmodell, in dem die Indikatoren die LV verursachen (Einkommen + Bildung + Beruf $\rightarrow$ SES).

Fornell-Larcker: klassischer Test für Diskriminanzvalidität: $\sqrt{AVE(A)} > correlation(A, B)$ für jedes Paar A, B.

GoF: Goodness-of-Fit. $\sqrt{(\text{mean}(\text{communality}) \times \text{mean}(R^2))}$. Proxy für die Gesamtmodellqualität.

HTMT: Heterotrait-Monotrait-Verhältnis. Moderner Test für Diskriminanzvalidität. Über 0.85 (konservativ) oder 0.90 (liberal) deutet darauf hin, dass die Konstrukte zu ähnlich sind.

HOC: Higher-Order-Construct. Ein Konstrukt, das aus mehreren First-Order-LVs zusammengesetzt ist (z. B. Brand Experience = Sensorisch + Affektiv + Intellektuell + Verhaltensbezogen).

Indikator / manifeste Variable (MV): eine beobachtete Datensatzspalte, die eine latente Variable misst.

Inneres Modell: die strukturellen Pfade zwischen LVs. Im Gegensatz zum äußeren Modell.

IPMA: Importance-Performance Map Analysis. Identifiziert, welche Vorgänger einer Ziel-LV hohe Wichtigkeit mit niedriger Performance kombinieren (größter Verbesserungshebel).

Latente Variable (LV): ein unbeobachtetes Konstrukt, das indirekt über Indikatoren gemessen wird.

Ladung: die Korrelation zwischen einem Indikator und seiner LV in einem reflektiven Messmodell.

Messmodell: die Pfeile zwischen LVs und ihren Indikatoren. Reflektiv (Mode A) oder formativ (Mode B).

MGA: Multi-Group-Analysis. Testet, ob sich Pfadkoeffizienten zwischen Gruppen unterscheiden.

MICOM: Measurement Invariance of Composite Models. Voraussetzung für MGA/Moderation über Gruppen hinweg.

Mode A: reflektive Messung (LV → Indikatoren).

Mode B: formative Messung (Indikatoren → LV).

Moderation: Xs Effekt auf Y hängt von einer dritten Variablen Z ab. Getestet über einen Interaktionsterm $X \times Z$.

Äußeres Modell: die Pfeile zwischen LVs und Indikatoren. Im Gegensatz zum inneren Modell.

Pfadkoeffizient (β): die standardisierte Stärke eines Innenmodell-Pfeils. Interpretation wie ein Beta aus einer OLS-Regression.

PLSc: Consistent PLS. Korrigiert den Attenuation-Bias, wenn Konstrukte tatsächlich Common-Factor- statt Composite-Konstrukte sind.

PLSpredict: Out-of-Sample-Prognosebewertung per k-facher Kreuzvalidierung.

Q^2 : Stone-Geisser kreuzvalidierte Redundanz. Über 0 = Prognoserelevanz.

Q^2_{predict} : die PLSpredict-Variante von Q^2 unter Verwendung von k-fold-CV.

Reflektiv (Mode A), die LV verursacht ihre Indikatoren. Ändern Sie das Konstrukt und alle Indikatoren ändern sich mit.

R^2 : Varianz in einer endogenen LV, die durch ihre Vorgänger erklärt wird. In-sample.

SRMR: Standardized Root Mean Square Residual. Unter 0.08 guter Fit, unter 0.10 akzeptabel.

Strukturmodell: gleichbedeutend mit innerem Modell.

t-Statistik: Pfadschätzwert geteilt durch seinen Bootstrap-Standardfehler. Über ~1.96 signifikant bei $\alpha = 0.05$.

VIF: Variance Inflation Factor. Unter 5 akzeptable Kollinearität, unter 3 strenger. Wird sowohl auf formative Indikatoren (Outer VIF) als auch auf Prädiktor-LVs in jedem endogenen Block (Inner VIF) angewandt.

13.2 Tastaturkürzel

Editor

- Cmd+Z / Ctrl+Z, Undo
- Cmd+Shift+Z / Ctrl+Y, Redo
- Delete oder Backspace auf einer ausgewählten LV, entfernt die LV und ihre Pfade
- Doppelklick auf eine Kante, entfernt den Pfad

Editor-Drag-Operationen

- Ziehen vom Source-Handle einer LV (blauer Punkt, rechte Seite) zum Target-Handle einer LV (grüner Punkt, linke Seite), zeichnet einen neuen Pfad
- Ziehen eines Indikator-Chips aus der rechten Seitenleiste auf einen LV-Kreis, weist den Indikator dieser LV zu

App-weit

- Cmd+K / Ctrl+K, zukünftig: Command-Palette (Roadmap)

13.3 Referenzen

Bakker, A. B., & Demerouti, E. (2017). Job demands-resources theory: Taking stock and looking forward. *Journal of Occupational Health Psychology*, 22(3), 273-285.

Brakus, J. J., Schmitt, B. H., & Zarantonello, L. (2009). Brand experience: What is it? How is it measured? Does it affect loyalty? *Journal of Marketing*, 73(3), 52-68.

Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340.

Dijkstra, T. K., & Henseler, J. (2015). Consistent partial least squares path modeling. *MIS Quarterly*, 39(2), 297-316.

Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. 3rd ed. Sage.

Hair, J. F., Sarstedt, M., Ringle, C. M., & Gudergan, S. P. (2018). *Advanced Issues in Partial Least Squares Structural Equation Modeling*. Sage.

Henseler, J. (2021). *Composite-Based Structural Equation Modeling: Analyzing Latent and Emergent Variables*. Guilford Press.

Henseler, J., & Chin, W. W. (2010). A comparison of approaches for the analysis of interaction effects between latent variables using partial least squares path modeling. *Structural Equation Modeling*, 17(1), 82-109.

- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115-135.
- Hult, G. T. M., Hair, J. F., Proksch, D., Sarstedt, M., Pinkwart, A., & Ringle, C. M. (2018). Addressing endogeneity in international marketing applications of partial least squares structural equation modeling. *Journal of International Marketing*, 26(3), 1-21.
- Lindell, M. K., & Whitney, D. J. (2001). Accounting for common method variance in cross-sectional research designs. *Journal of Applied Psychology*, 86(1), 114-121.
- Park, S., & Gupta, S. (2012). Handling endogenous regressors by joint estimation using copulas. *Marketing Science*, 31(4), 567-586.
- Russett, B. M. (1964). Inequality and Instability: The Relation of Land Tenure to Politics. *World Politics*, 16(3), 442-454.
- Sarstedt, M., Hair, J. F., Cheah, J.-H., Becker, J.-M., & Ringle, C. M. (2019). How to specify, estimate, and validate higher-order constructs in PLS-SEM. *Australasian Marketing Journal*, 27(3), 197-211.
- Schuberth, F., Rademaker, M. E., & Henseler, J. (2023). Assessing the overall fit of composite models estimated by partial least squares path modeling. *European Journal of Marketing*, 57(6), 1678-1702.
- Shmueli, G., Ray, S., Estrada, J. M. V., & Chatla, S. B. (2016). The elephant in the room: Predictive performance of PLS models. *Journal of Business Research*, 69(10), 4552-4564.
- Tenenhaus, M., Vinzi, V. E., Chatelin, Y.-M., & Lauro, C. (2005). PLS path modeling. *Computational Statistics & Data Analysis*, 48(1), 159-205.
- Venkatesh, V., Thong, J. Y., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*, 36(1), 157-178.