



用户指南

OpenPLS 用户指南

面向研究人员与实践者的完整指南

版本 0.3.28

2026-08-04

openpls.app · 面向研究人员的现代 PLS-SEM

Contents

1	什么是 PLS-SEM?	4
1.1	基本构成要素	4
1.2	迭代过程	4
1.3	OpenPLS 为你提供什么	4
2	快速入门	5
2.1	Projects 页面	5
2.2	创建项目	5
2.3	克隆一个科学示例	5
2.4	初次进入项目	7
3	数据集	8
3.1	空的 Dataset 面板	8
3.2	数据集准备就绪后	9
3.3	编辑缺失码	9
3.4	管理多个数据集	10
3.5	删除数据集	10
4	构建模型	10
4.1	可视化编辑器	10
4.2	画布	12
4.3	表单视图	13
4.4	实时计算	13
4.5	版本	14
5	计算并阅读结果	15
5.1	Results 页头	15
5.2	Inner summary (默认)	15
5.3	Paths	16
5.4	Inner model / outer model	16
5.5	Reliability	16
5.6	HTMT	17
5.7	VIF	17
5.8	Effects	18
5.9	Model fit	18
6	进阶方法	18
6.1	Bootstrap	19
6.2	多组分析 (MGA)	19
6.3	MICOM	22
6.4	Moderation	22

6.5	Nomological validity	22
6.6	PLSc	25
6.7	Gaussian copula (内生性)	26
6.8	Common method bias (CMB, Lindell-Whitney)	27
6.9	IPMA	27
6.10	PLSpredict	27
6.11	CVPAT	27
6.12	FIMIX-PLS	32
6.13	高阶构念 (HOC)	32
7	导出	35
7.1	PDF 报告	35
7.2	Excel 工作簿 (XLSX)	35
7.3	Report JSON 包	36
7.4	模型导出 (Editor 工具栏)	36
7.5	LaTeX 片段	36
7.6	可重现性	36
8	协作	37
8.1	Team 面板	37
8.2	角色	37
8.3	发送邀请	37
8.4	活动日志	37
8.5	离开项目	39
8.6	协作者能看到什么	39
9	设置与语言	39
9.1	项目设置	39
9.2	语言	40
9.3	账户设置	42
9.4	主题	42
10	帮助与延伸阅读	42
10.1	应用内帮助	42
10.2	OpenPLS 引擎 (开源)	43
10.3	支持、Bug 和功能请求	44
10.4	引用 OpenPLS	44
11	AI Companion	44
11.1	启用	44
11.2	AI hints: 常驻可见的条带	45
11.3	Chat: 浮动抽屉	45
11.4	从研究问题设计模型	47

11.5 设计问卷	49
11.6 Method-Selection Assistant	52
11.7 报告起草	53
11.7.1 Methods / Results / Discussion	53
11.7.2 Reviewer defense	55
11.8 论文库	55
11.9 配额与隐私	55
11.10 关闭	57
12 公开调查	57
12.1 构建问卷	57
12.2 发布	57
12.2.1 版本快照	59
12.3 公开受访者视图	59
12.4 观察响应到达	60
12.5 将响应导入为数据集	60
12.6 下载原始 CSV	60
12.7 取消发布	63
12.8 限制与惯例	63
13 附录	63
13.1 术语表	63
13.2 键盘快捷键	64
13.3 参考文献	65

1 什么是 PLS-SEM?

偏最小二乘结构方程建模 (Partial Least Squares Structural Equation Modeling, PLS-SEM) 是一种用于估计 PLS-SEM 通过一次估计同时回答三个相互关联的问题:

1. 测量。可观测的指标 (问卷条目、评分、传感器读数) 在多大程度上代表了每个潜在构念?
2. 结构。每个构念对其他构念的影响强度如何? 感知质量是否驱动满意度? 信任是否在有用性对采纳的影响中起作用?
3. 预测。模型的泛化能力如何? 对于新的观测样本, 结果会是什么?

与基于协方差的 SEM (LISREL、Mplus、lavaan) 不同, PLS-SEM 不要求数据服从多元正态分布, 适用于小样本 (例如: 收入 + 教育 + 职业 → 社会经济地位)。这使其成为商业研究、市场营销、信息系统以及日益增多的健康和社会科学领域的理想选择。

1.1 基本构成要素

一个 PLS-SEM 模型由四类元素构成:

- 潜变量 (LVs): 你所关心的构念。以圆圈表示。
- 指标 / 显变量 (MVs): 数据集中用于测量每个 LV 的观测列。在传统路径图中以方框表示; 在 OpenPLS 中, 它们位于侧边栏, 可分配给 LV。
- 测量模型 / 外部模型: LV 与其指标之间的箭头。方向区分了反映型 (Mode A: LV 引起指标, 例如满意度 → sat1..sat4) 与形成型 (Mode B: 指标引起 LV, 例如收入 + 教育 → SES)。
- 结构模型 / 内部模型: LV 之间的箭头。这些是待检验的假设。

1.2 迭代过程

PLS-SEM 交替进行两个最小二乘回归, 直至收敛:

1. 使用当前的内部模型权重, 将每个 LV 估计为其指标的加权组合。
2. 通过将每个内生 LV 对其前置变量进行回归, 估计内部模型的路径系数。

收敛速度很快 (通常 < 10 次迭代)。标准误和 p 值来自 Bootstrap: 对数据进行有放回的重抽样数千次, 并重新估计模型。

1.3 OpenPLS 为你提供什么

OpenPLS 是一个现代化、开放、云原生的 PLS-SEM 工具。在一个项目中你可以:

- 上传数据集 (CSV、XLSX、SPSS .sav、Stata .dta), 或者构建问卷、将其发布为公开调查, 并一键将数据导入 OpenPLS (第 13 章)。
- 可视化地绘制模型 (从 LV 拖到 LV)、在表单中输入模型, 或者通过 AI Companion 根据研究问题生成模型 (第 12 章)。
- 计算 PLS 估计, 每次更改都可实时重新计算。
- 评估测量质量: 信度、AVE、HTMT、VIF。
- 评估结构质量: R^2 、调整后 R^2 、 f^2 、 Q^2 、SRMR、dULS、Fornell-Larcker。
- 对路径进行 Bootstrap 以获取 p 值和置信区间。

- 运行进阶方法：MGA、MICOM、调节效应、中介效应、IPMA、PLSpredict、FIMIX、HOC、PLSc、高Copula、CVPAT、法则效度（Nomological validity）、共同方法偏差。
- 在每个标签页获取情境化的AI提示，与一个对你的项目数据具有只读访问权限的助手对话，起草可发表的结果（Results）/ 讨论（Discussion）章节以及预期的审稿人答辩 — 所有内容都建立在一个精心策划的开放PLS-SEM论文库之上（第12章）。
- 将整个运行导出为PDF、XLSX、LaTeX，或用于可重现性的自包含JSON包。
- 邀请协作者、跟踪活动，并保留每个模型版本的审计跟踪。

本指南的其余部分将按你实际使用的顺序逐一介绍这些功能。

本指南是针对OpenPLS 0.3.27版本编写的，运行地址为 cloud.openpls.app。后续版本中界面可能略有不同。

2 快速入门

你需要一个账户来保存项目。前往 <https://cloud.openpls.app>，使用邮箱注册、验证并登录。你将进入Projects页面，即你的工作区。

2.1 Projects 页面

Projects页面（图1）是你的主界面。它列出了你有权访问的每个项目，并在页头指向两个起始操作，同时在空白处提供三种起始方式（图1中已编号）：

1. New project（右上角）：一个空白工作区。你上传自己的数据集并从零开始构建模型。
2. Examples：示例案例库。一键即可将一个经典的科学PLS-SEM案例（ECSI、MOBI、Russett、若干合成数据集）+ 模型已预先关联。
3. Create project（空状态卡片中）：与(1)相同对话框的快捷方式。使用这个按钮或页头按钮，哪个离你最近。

2.2 创建项目

点击任意一个New project按钮会打开图2所示的对话框。

填写内容（图2中已编号）：

1. Project name：必填， ≥ 2 个字符。从“2026年顾客满意度”到“H1实验，国别队列”均可。
2. Description：可选，最多500个字符。用于记录研究问题或正在检验的假设。当你六个月后回来时非常有帮助。
3. 点击Create。你将进入新的项目详情页面。由于尚未附加数据集，应用会打开Dataset标签页，以便你上传数据。

2.3 克隆一个科学示例

如果你是PLS-SEM新手，可以从一个已知良好的模型开始。点击页头的Examples打开示例库（图3）。

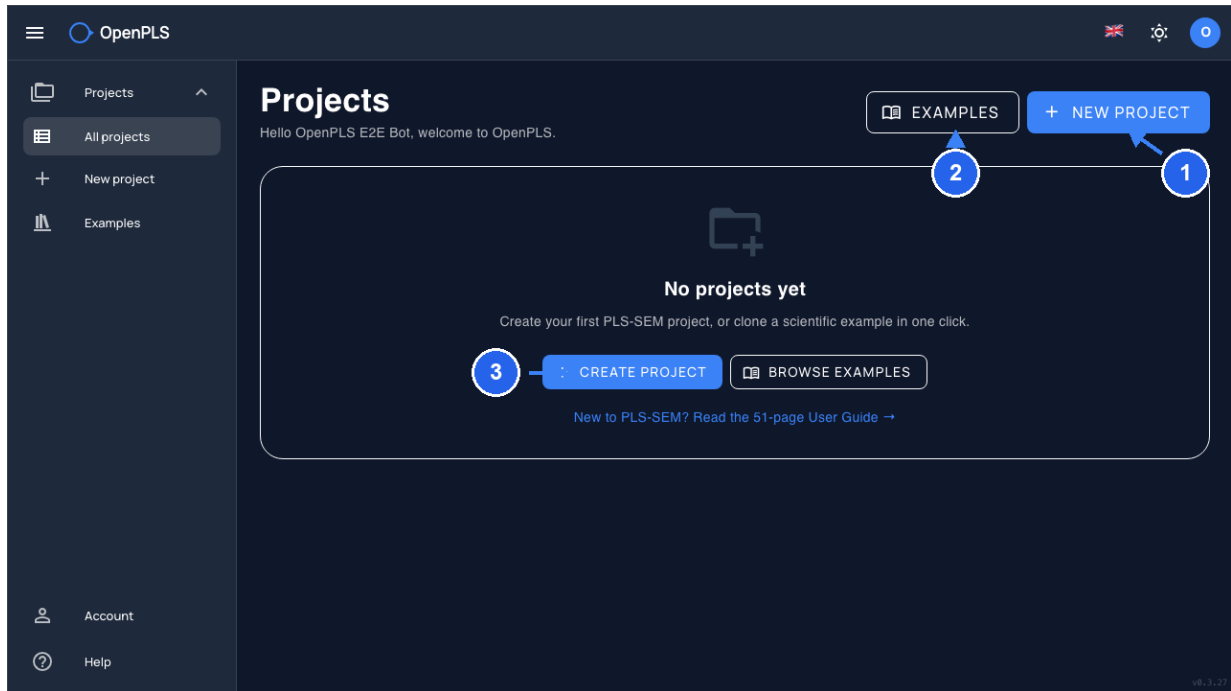


Figure 1 Projects 首页：页头提供 New project 和 Examples；中央空状态卡片提供相同的两个操作。

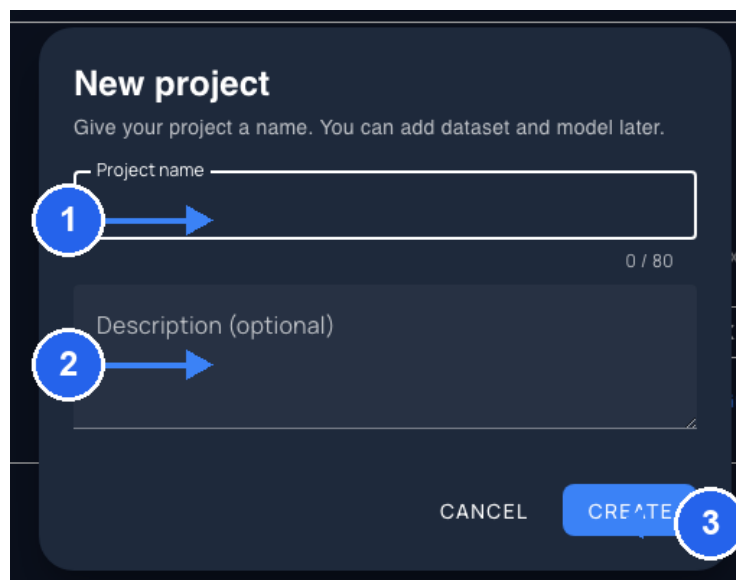


Figure 2 New project 对话框，包含名称字段、可选描述字段，以及右下角的 Create 操作。

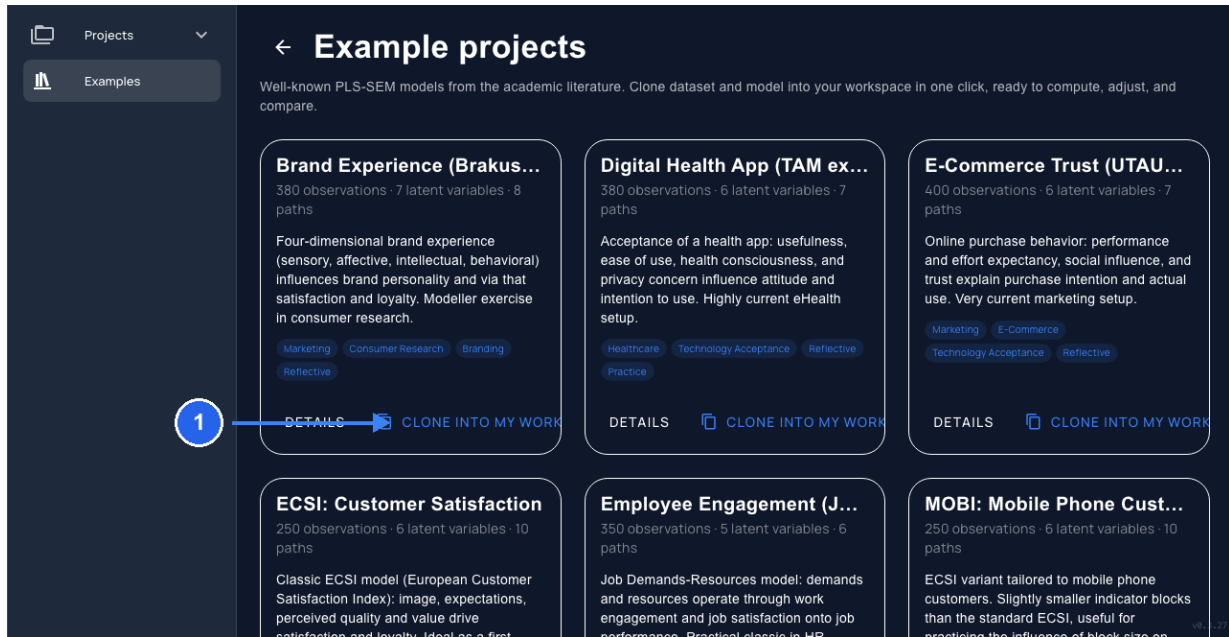


Figure 3 Examples 库。每张卡片显示数据集统计信息和引用；点击任意卡片即可打开，然后点击 Clone into my workspace 制作可编辑的副本。

1. 选择一张匹配你所在领域的卡片（顾客满意度、品牌体验、技术接受、工作投入、健康应用、政治学……+ 模型）。
2. 点击 Clone into my workspace。OpenPLS 会在你的账户下创建一个私有副本。你可以自由删除指标、LV、更改路径，而原始版本对其他用户保持完好。

这些示例是来自 PLS-SEM 文献的开放许可教学案例（Tenenhaus et al. 2005、Brakus & Schmitt 2009、Bakker & Demerouti 2017、Russett 1964、Venkatesh et al. 2012、Davis 1989）。完整引用信息在每张卡片上。合成数据集有明确标注，是根据已发表的路径结构模拟生成，仅供教学

2.4 初次进入项目

进入项目后，你会在顶部看到标签栏：

- Dataset: 上传、检查、编辑缺失码约定
- Model: 构建模型（Editor / Form / Versions）
- Results: 查看计算输出
- Advanced: 13 种进阶方法
- Manage: 团队、活动日志、项目设置

我们将依次介绍每个标签页。下一步：上传数据集。

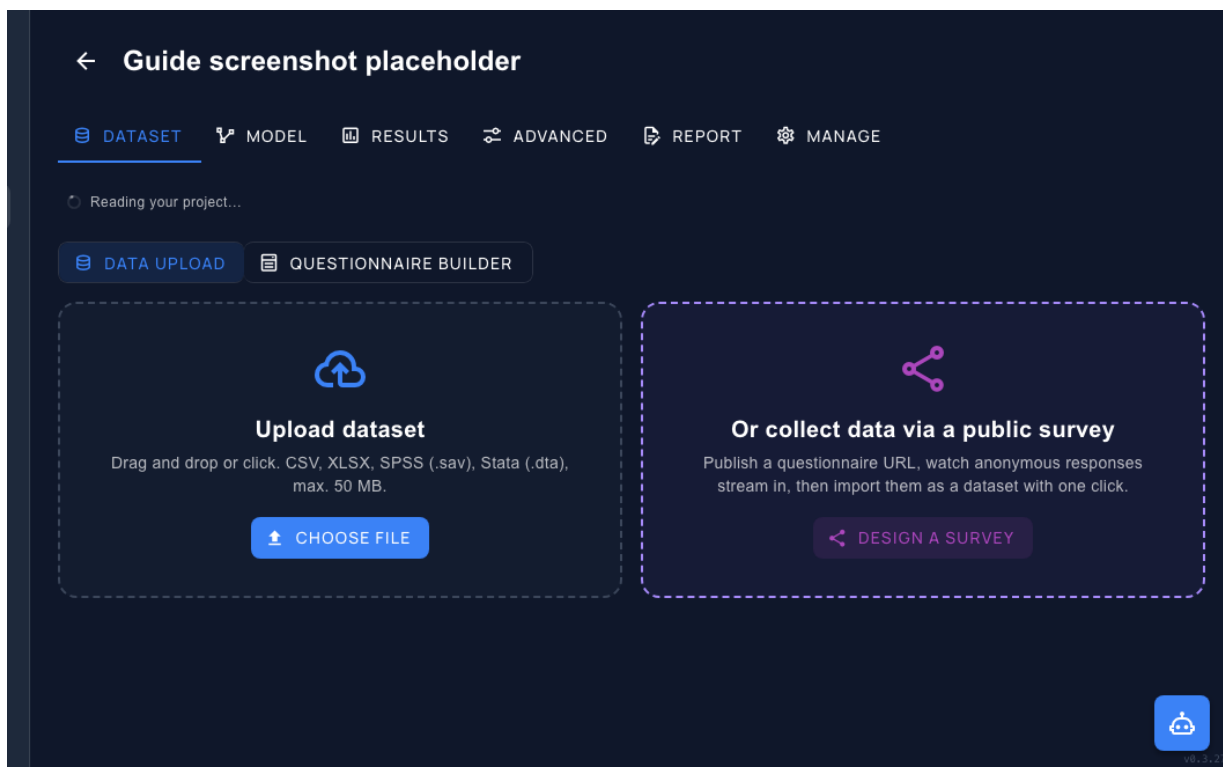


Figure 4 空的 Dataset 标签页。顶部的标签栏在 Dataset、Model、Results、Advanced、Report 和 Manage 之间切换。下方的子切换器用于选择上传文件或构建公开调查。

3 数据集

每个项目在构建模型之前都需要且只需要一个数据集。OpenPLS 支持读取 CSV、XLSX、SPSS **.sav** 和 Stata **.dta** 格式，每个文件最大 50 MB。

3.1 空的 Dataset 面板

新项目在打开时默认激活 Dataset 标签页。如果尚未附加任何数据，你会看到图 4 所示的上传区域。

两步导入数据：

1. 确认你在 Dataset 标签页。它是每个项目的第一个标签页。
2. 将文件拖到虚线拖放区，或点击它打开操作系统的文件选择器。

对于 CSV 上传，OpenPLS 接下来会打开 CSV separator 对话框。它会预览前几行，并自动建议分隔符（,、; 或 |）；确认后继续。

从欧洲版 Excel 导出的 CSV 通常使用 ‘;’ 作为分隔符。OpenPLS 会自动检测，但在点击上传之前请确认预览是否正确，分隔符错误是产生无意义数据集的最快途径。

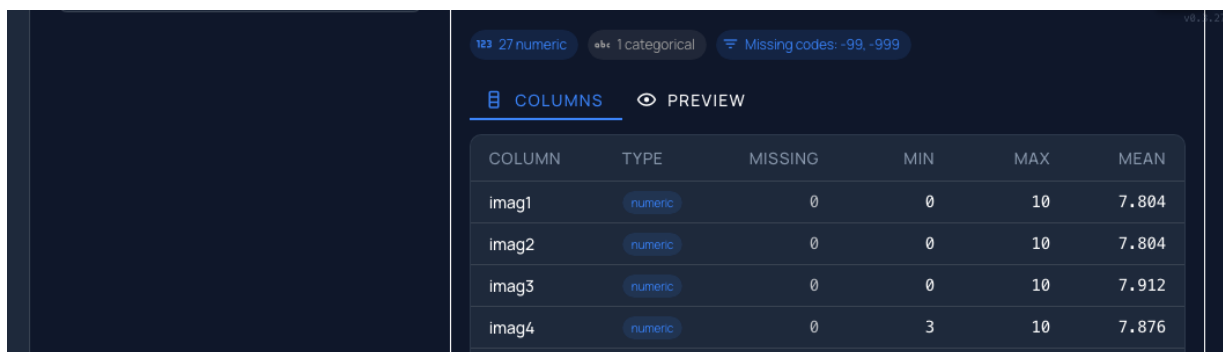


Figure 5 shows the Dataset panel after upload. The top bar indicates 27 numeric columns, 1 categorical column, and missing codes: -99, -999. The 'COLUMNS' tab is active, displaying a table with the following data:

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

Figure 5 上传后的 Dataset 面板。数据集卡片显示 Ready 状态标记和当前的缺失码哨兵值；右侧主面板显示各列统计信息。

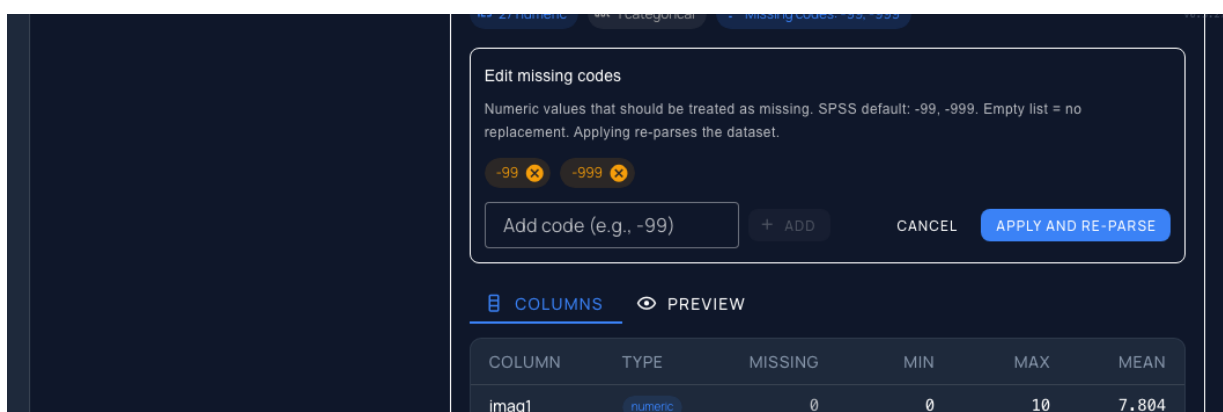


Figure 6 shows the Missing code editor. It allows editing missing codes. The text says: "Numeric values that should be treated as missing. SPSS default: -99, -999. Empty list = no replacement. Applying re-parses the dataset." Below this, there are input fields for missing codes, currently showing -99 and -999, each with a delete icon (X). There is an "Add code (e.g., -99)" button, an "ADD" button, a "CANCEL" button, and an "APPLY AND RE-PARSE" button.

Figure 6 缺失码编辑器。草稿标记可通过 × 移除；新的哨兵值通过输入框添加；Apply and re-parse 提交更改并重新读取文件。

3.2 数据集准备就绪后

解析在服务器端进行，对于典型文件仅需一两秒。完成后，数据集卡片会出现在左栏，右侧渲染出其详细信息（图 5）。有两个标记很重要：

1. 状态标记：解析完成后变为绿色并显示 “Ready”。如果保持红色（“Error”），点击数据集查看错误信息。
2. 缺失码标记：显示 OpenPLS 当前视为缺失值的数值。默认的 -99, -999 符合 SPSS 惯例。点击该标记进行编辑。

右侧面板列出每一列及其类型（numeric、string、mixed）、缺失单元格的数量，以及（对于数值列）最小值和最大值。对于 string 的，通常意味着有一个非数值单元格混入其中。

3.3 编辑缺失码

一些数据集用空白编码缺失，另一些使用 -99（SPSS 默认值），还有的使用 9999 或 -1。这一步搞错是“为什么我的 R^2 这么低？”这类问题最常见的原因——模型会把哨兵值当作真实数据，从而拉低相关性。

点击缺失码标记，打开图 6 所示的编辑器。

图 6 中的三个元素：

1. 现有哨兵值：每个都是一个标记。点击标记上的 × 将其从集合中移除。
2. 添加代码：输入任意整数（正数或负数，只要符合你数据集的约定），点击 Add。它将以新标记的形式添加到集合中。
3. Apply and re-parse：将新的集合发送到服务器，服务器使用更新后的哨兵值重新读取文件并刷新列统计信息。

应用缺失码会从头重新解析文件。如果你已经计算过模型，该计算会失效，在重新解析后需要重新运行。

3.4 管理多个数据集

一个项目可以并排持有多个数据集（例如每个国家一个，或试点样本与完整样本）。通过拖动或选择其他文件到 Dataset 标签页显示的内容，模型仍然指向创建时的那个数据集。

3.5 删除数据集

要删除数据集，将鼠标悬停在左栏中的卡片上，然后点击垃圾桶图标。删除是即时的，无法撤销。如果项目中的

数据集会立即被删除且无法恢复。项目中仍指向已删除数据集的任何模型在计算时都会出错，请在 Model 标签页中重新关联另一个数据集，或者一并删除该模型。

4 构建模型

OpenPLS 提供三种构建模型的方式：

- Editor：可视化拖放画布。这是主要的工作流程，也是 OpenPLS 的杀手级功能。
- Form：一个包含 LV 表和路径表的普通表单。当你模型了然于胸或从论文中复制时更快捷。
- Versions：每次过往计算的审计日志。每次计算都会快照模型 + 结果，因此你可以恢复到任何先前的状态。

这三个子视图编辑的是同一个底层模型。可以自由切换。

4.1 可视化编辑器

点击 Model → Editor 打开图 7 所示的画布。工具栏位于画布上方；右侧的指标侧边栏列出数据集的每一个数值型变量。从左到右阅读图 7，工具栏包含：

1. 模型选择器：一个项目可以对同一数据集持有多个模型。从下拉框中选择一个。
2. New model：创建一个新的空模型。对话框要求命名；稍后可在 Settings 中重命名。
3. 数据集选择器：每个模型都绑定到一个数据集。如果你上传了多个 CSV，在这里选择。
4. Undo / Redo：macOS 上是 Cmd+Z / Cmd+Shift+Z，Windows 和 Linux 上是 Ctrl+Z / Ctrl+Y。Undo 会回退 LV 的创建、指标分配和路径绘制。
5. Add LV：在画布的随机位置放置一个新的潜变量圆圈并自动选中它。
6. Autolayout：运行 dagre 算法按拓扑顺序从左到右排列节点。这是从一团重叠的圆圈到达可发表水平的。
7. Fit to window：将画布重新居中于所有可见节点。
8. Export model：将模型下载为 .openpls.json（完整的 OpenPLS 格式，包含位置信息）或 .splsm（SmartPLS 兼容的 XML）。

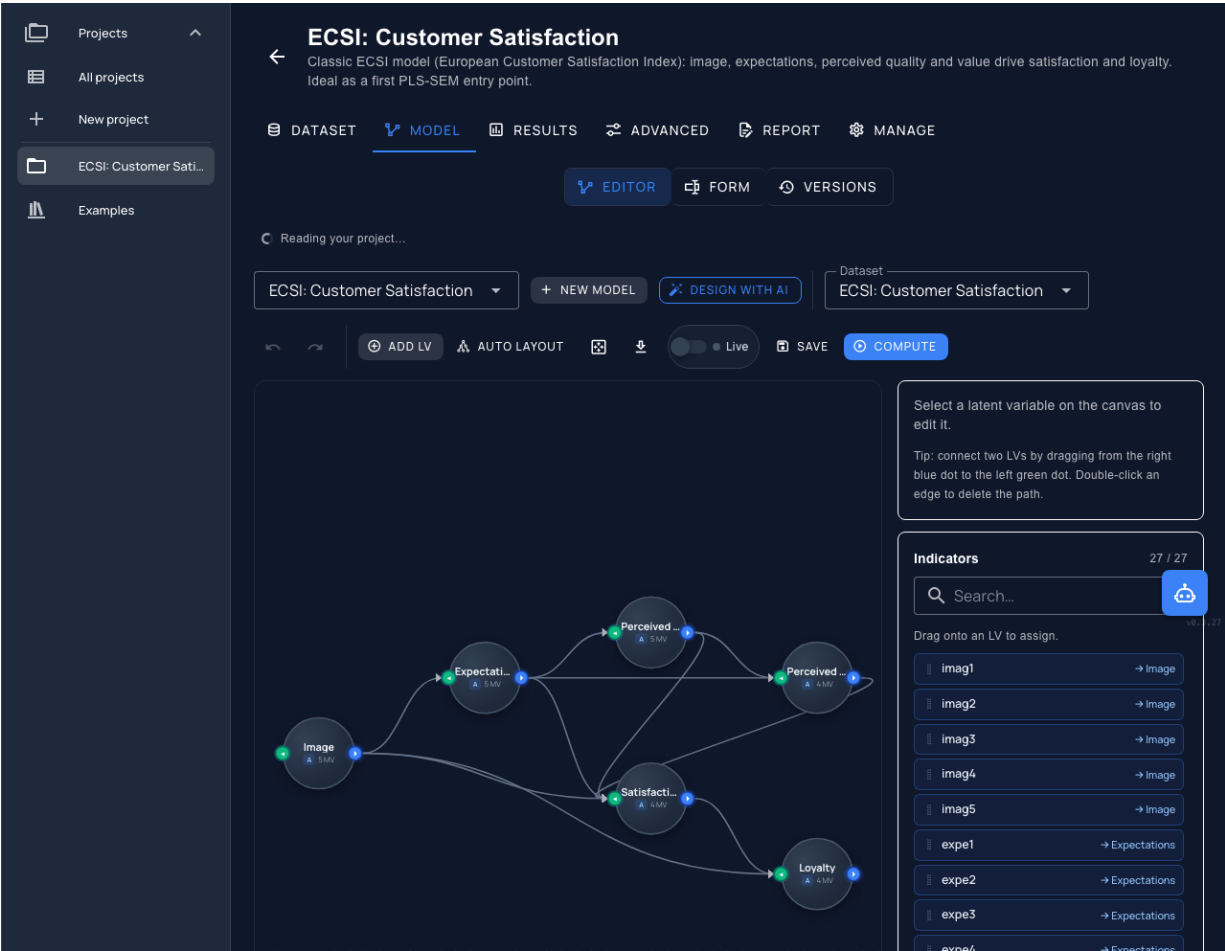


Figure 7 可视化编辑器：工具栏（顶部）、Vue Flow 画布（中央）、指标侧边栏（右侧）。

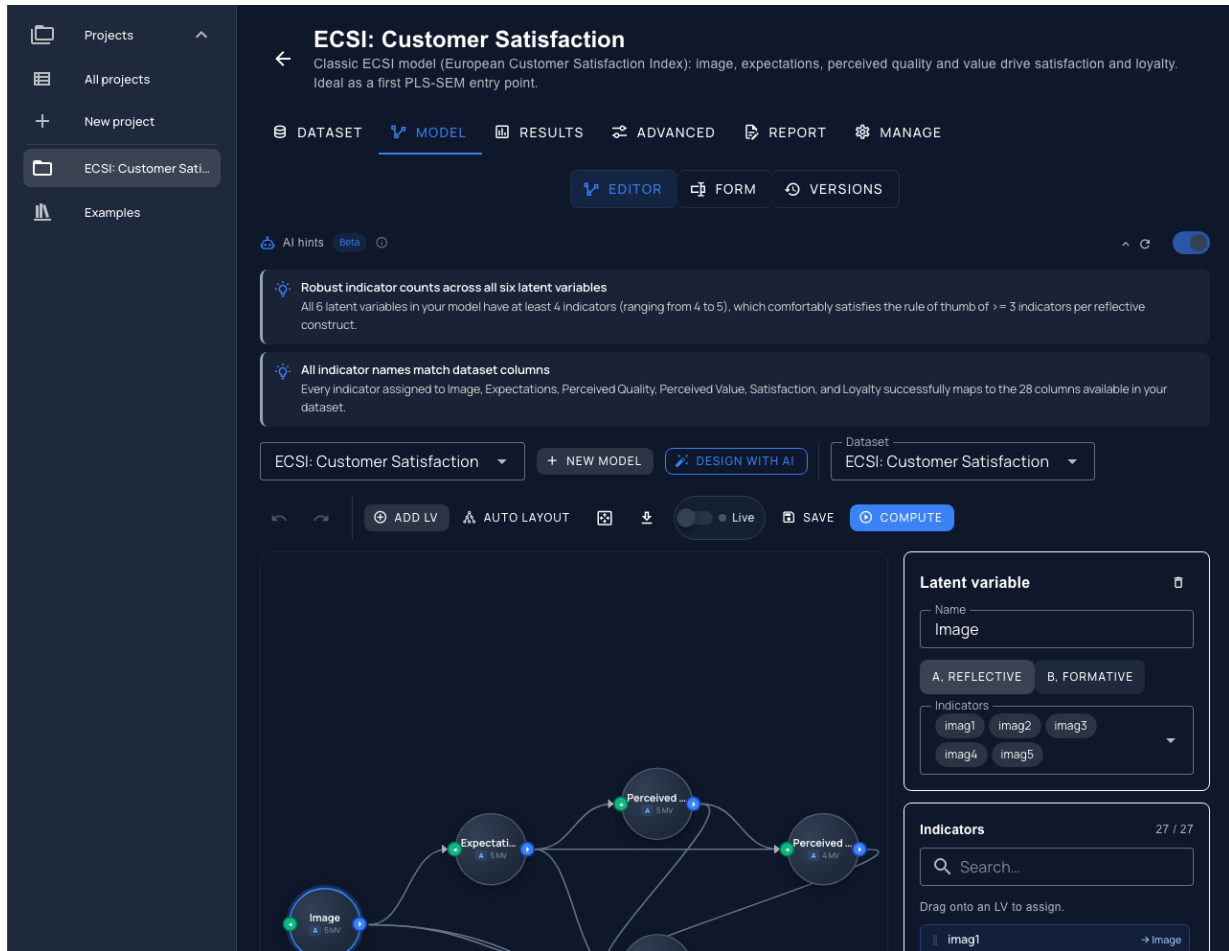


Figure 8 LV 被选中。右侧卡片让你重命名、更改测量模式和选择指标。

9. **Live 切换**：打开后，每次更改后 500ms，OpenPLS 就会重新计算模型。路径系数会渲染在边上， R^2 会出现在每个内生 LV 内。快速迭代时的即时反馈。
10. **Save**：将当前模型写入数据库。没有自动保存；点击 Save（或 Compute，它会隐式保存）以持久化你
11. **Compute**：在服务器端运行完整的估计（Bootstrap + 指标），写入一个 Result 文档，并导航到 Results 标签页。

4.2 画布

每个潜变量渲染为一个圆圈，包含：

- LV 名称
- 显示测量模式的胶囊：A（反映型）或 B（形成型）
- 已分配指标的计数，例如 “5 MV”
- 有计算结果后，圆圈内会显示 R^2 值

每条路径渲染为两个 LV 之间的箭头。在源 LV 上，右侧的小蓝点是源手柄；在目标 LV 上，左侧的小绿点是目标手柄。右键或选中一个 LV 会在右侧打开编辑器（图 8）。

1. **Name**：自由文本。重命名会传播到每一条引用该 LV 的路径。

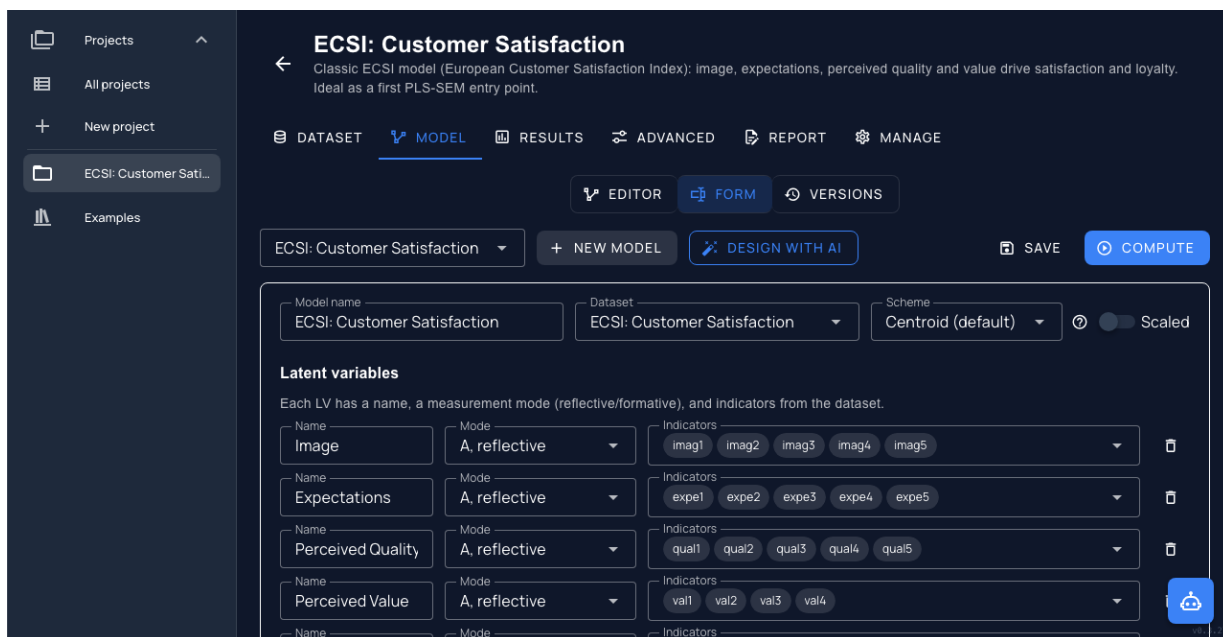


Figure 9 Form 子视图：顶部是 LV 表（名称、模式、指标），下方是路径表，右上角是 Save + Compute。

2. Mode A / Mode B：反映型 vs. 形成型。反映型意味着 LV 引起其指标（改变满意度，所有 sat 项一起变化）；形成型意味着指标定义 LV（收入、教育、职业共同定义 SES）。选错的话，信度指标 B 却按反映型解释），要么无法收敛（Mode A 用于形成型构念）。
3. Indicators：一个针对数据集所有数值列的多选。在此添加一个指标会立即在 LV 圆圈内显示为箭头。或者 Indicators 侧边栏直接把标记拖到 LV 圆圈上，OpenPLS 会在释放时放置该指标。
4. Delete LV：标题中的小垃圾桶图标。移除该 LV 以及指向它或从它出发的任何路径。

右侧的 Indicators 侧边栏列出每一个剩余的数值列，并带有一个徽章显示已在使用它的 LV（若有）。一个指标可以出现在多个 LV 中，但这通常表明建模有误，请在继续之前调查清楚。

4.3 表单视图

并非每个用户都想点击圆圈。切换到 Form 查看基于表格的编辑器（图 9）。

表单是一个普通的 HTML 表单。使用 LV 表底部的按钮添加 LV；从每行的多选框中选择指标。使用路径表下方的 LV 下拉框中选择。你在编辑器中能做的这一切在这里也能做。

何时使用表单：从论文中复制粘贴模型、键盘用户，或自动化模型构建。

4.4 实时计算

工具栏中的 Live 切换开启增量重新计算。你一有任何改动（添加 LV、连接路径、勾选指标、重命名），OpenPLS 就会防抖 500 ms，然后在服务器端运行完整的 PLS 估计。结果流回并绘制：

- 每条边上的路径系数

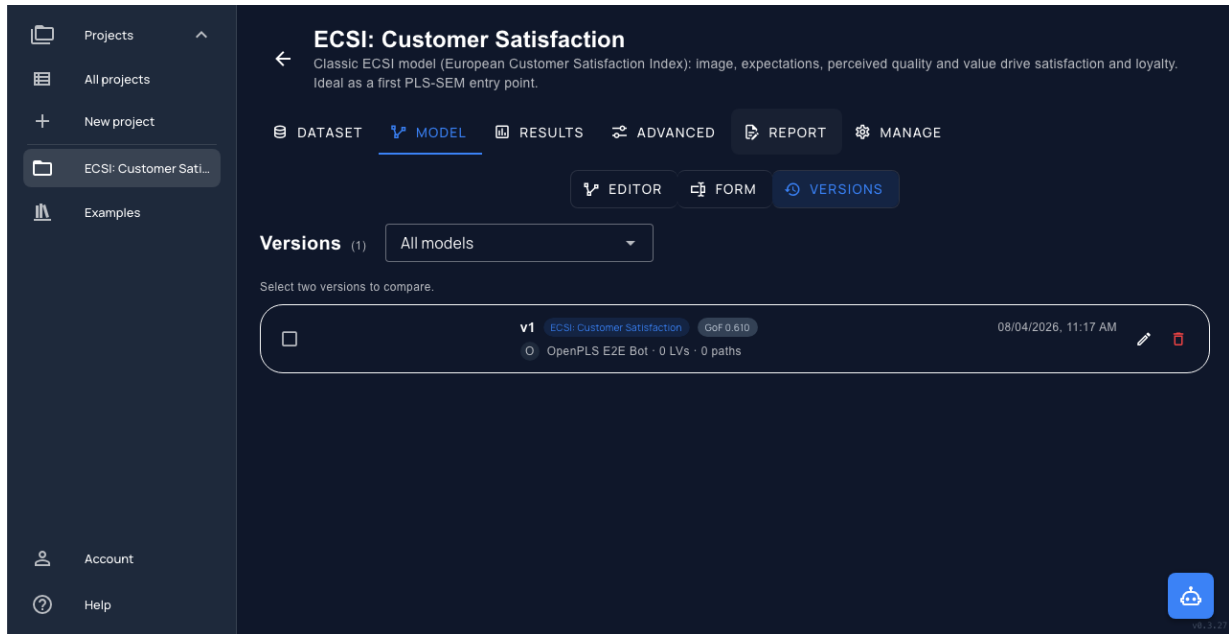


Figure 10 Versions 列表。每次先前的计算都是一张卡片，显示 LV/路径计数、GoF 以及创建者的头像。

- 每个内生 LV 内部的 R^2 与调整后 R^2
- 路径上的 p 值（来自最后一次完整 Compute 时运行的 Bootstrap，不会在每次按键时重新 Bootstrap）

工具栏中的 Live 标记会在 computing...、waiting... 和 error 之间循环，让你了解后台任务的状态。

实时计算是探索”添加一条新路径或重新分配一个指标会发生什么”最快的方法。对于希望在手动按下 Compute 前做许多编辑的大型模型，请关闭该功能。

4.5 版本

每次你点击 Compute，OpenPLS 都会将完整的模型 + 结果快照到一个版本化条目中。切换到 Versions 子视图（图 10）浏览它们。

每张版本卡片显示：

- v#：连续的版本编号
- 它所属的模型（一个标记）
- GoF：该版本结果的拟合优度（若成功）
- 创建者的头像和一行统计信息（LV 数、路径数）
- Rename（铅笔图标）：为该版本打标签以便未来参考，例如 “pre-invariance test” 或 “final submission”
- Delete：移除该版本。底层的模型 + 结果文档也会一并删除。

恢复版本会用快照替换当前实时的编辑器状态。路径位置、指标、模式，一切都会替换。

比较两个版本：勾选两张卡片上的复选框。右侧会打开一个并排对比视图，显示两个版本之间新增、删除或更改

恢复版本会覆盖你的实时编辑器状态。如果你有未保存的编辑，请先点击 Save（或做一次新的 Compute），以便有一个可回退的快照。

5 计算并阅读结果

一旦模型至少有一个内生 LV，并且每个 LV 都拥有 ≥ 1 个指标，Editor/Form 工具栏中的 Compute 按钮会变为实心蓝色。点击它。计算在服务器端运行；典型的 ECSI 规模计算在 3-6 秒内完成。成功后，C 会导航到 Results 标签页。

5.1 Results 页头

Results 标签页（图 11）打开时，顶部有四张 KPI 卡片、右侧有三个导出按钮，下方是详细表格的标签栏。

页头显示四张 KPI 卡片：

1. Goodness-of-Fit (GoF): $\sqrt{(\text{mean}(\text{communality}) \times \text{mean}(R^2))}$ 。粗略的整体拟合指标；低于 0.10 为弱，0.25 为中等，0.36 以上为强。
2. SRMR：标准化残差均方根。观测相关矩阵与模型隐含相关矩阵之间的差异。低于 0.08 是良好拟合的传统 0.10 可接受。副标题中的 dULS 值是另一种差异度量。
3. Paths：估计模型中内部模型路径的数量。
4. LVs/Indicators：潜变量的数量和指标（MV）总数。在比较多次运行时用作合理性检查。

KPI 上方是一个结果选择下拉框，列出当前模型的每一次过往 Compute，带有时间戳。可在不离开该标签页的选择器右侧是三个导出按钮：

- PDF：可打印的报告（路径图、表格、参考文献）。
- XLSX：一个工作簿包含所有结果表，每个指标一个工作表（Paths、Inner summary、Outer loadings、Reliability、HTMT、Fornell-Larcker、VIF、Effects……）。
- Report JSON：一个自包含的 JSON 包，包含模型、数据集架构和每一个结果表。是论文补充材料的理想

5.2 Inner summary（默认）

Inner Summary 表列出每个潜变量以及：

- R^2 ：由其前置变量解释的 LV 方差。0.10 为弱，0.25 为中等，0.50 为强。外生 LV（无前置变量）显示为 0。
- $R^2 \text{ adj.}$ ：根据预测变量数量进行的调整。跨模型比较时优先使用此指标。
- BIC：LV 外部模型的贝叶斯信息准则。越低越好；只有在嵌套规范之间才有比较意义。
- Block communality：LV 指标的平均平方载荷。越高 = 测量越强。
- Mean redundancy：区块中有多少方差与 LV 的前置变量“冗余”。
- AVE：平均方差抽取量。高于 0.5 意味着 LV 捕获的方差多于误差。

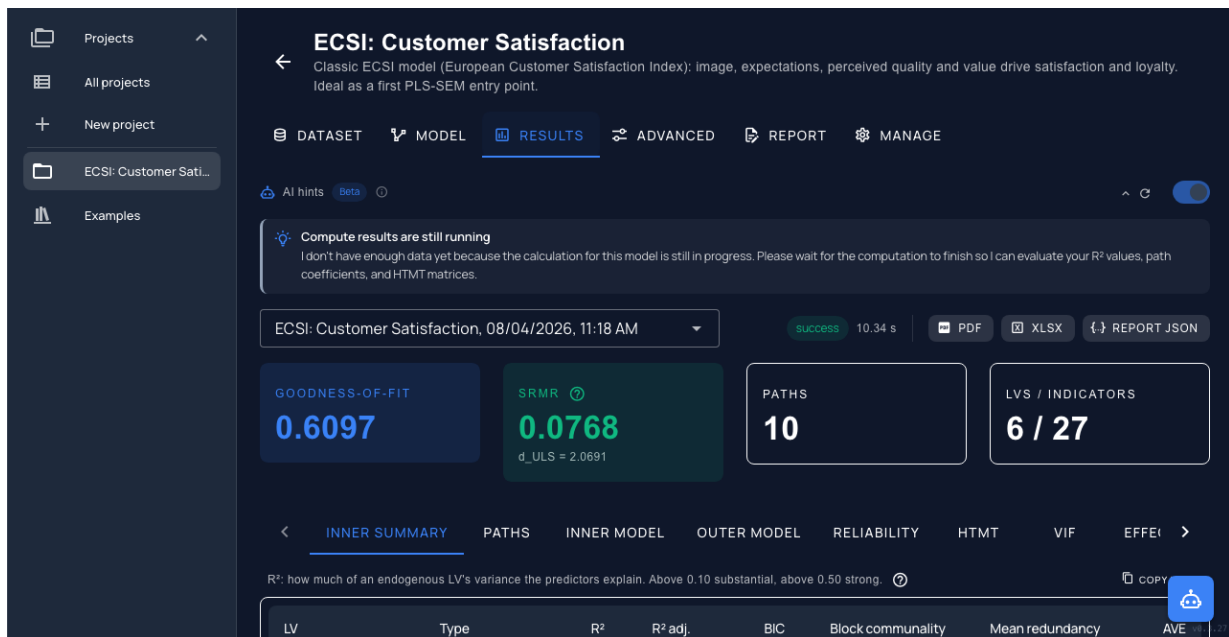


Figure 11 Results 标签页：拟合优度、SRMR、路径数、LV/MV 数的 KPI 卡片。下方是分标签的详细内容：Inner Summary、Paths、Inner Model、Outer Model、Reliability、HTMT、VIF、Effect Sizes、IPMA。

5.3 Paths

切换到 Paths 子标签页（图 12）查看结构系数表。

对于每条路径（源 → 目标），表中列出：

- Estimate：标准化的路径系数。像 OLS 回归中的 β 一样解释。
- Std. err.：来自 Bootstrap。
- t：estimate / std. err. 大于 1.96 在双侧检验中于 $\alpha = 0.05$ 显著。
- p-value：Bootstrap 的 p，与 t 相对应。
- 95% CI (lower, upper)：百分位 Bootstrap 置信区间。如果它排除零，则该路径显著。

通过表格右上角的小复制图标可复制 LaTeX 就绪版本。

5.4 Inner model / outer model

Inner model 标签页在一个更宽的表中重述路径，LV 对作为行。Outer model 反过来：每个指标一行，显示其 LV、载荷（Mode A）或权重（Mode B）以及 t 统计量。

解读载荷：高于 0.708 ($\sqrt{0.5}$) 意味着 LV 解释了 $\geq 50\%$ 的指标方差。低于 0.4 的载荷是删除的候选；在 0.4 到 0.7 之间的仅在删除会降低 AVE 时保留。

5.5 Reliability

Reliability 标签页显示每个 LV 的两个指数以及两个区块级诊断：

- Cronbach's α ：传统指数。假设 tau 等价载荷；更保守。

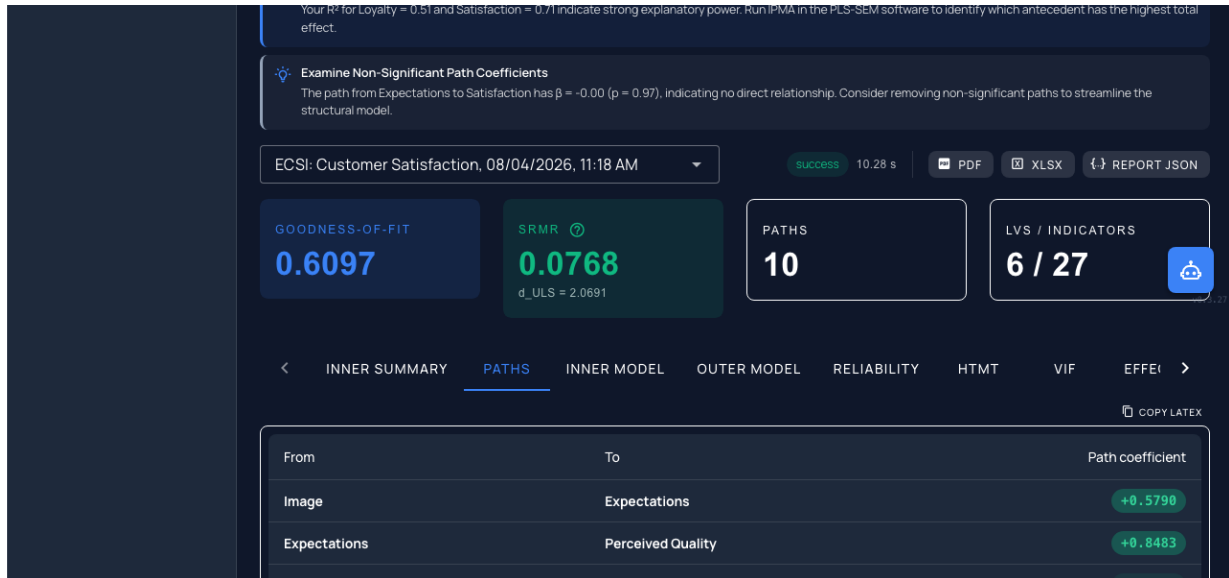


Figure 12 Paths 子标签页：每条内部模型路径及其估计值、标准误、t 统计量、p 值和 95% CI。

- 组合信度 (Dillon-Goldstein ρ)：现代 PLS-SEM 默认指标。0.7-0.9 可接受，>0.95 表明指标冗余。
- 第 1 / 第 2 特征值：区块相关矩阵的特征值。两者之间有明显差距（第一个远大于 1，第二个 ≈ 1 ）支持反映型区块的单维性。

反映型构念的目标是 $\alpha > 0.7$ 且 $\rho > 0.7$ 。对于形成型构念，信度没有意义，请检查 VIF。

Dijkstra & Henseler 的一致信度 (ρ_A) 仅在你运行 PLSc 时报告（见 6.6 节）。

5.6 HTMT

HTMT 子标签页（图 13）显示每对构念的异质一同质性比率，红色单元格突出显示超出区分效度阈值的值。

HTMT 是现代的区分效度检验。对于每对构念，它计算”异质”相关性（跨构念）与”同质”相关性（构念内部），两个构念就难以区分。

阈值：

- < 0.85 (Henseler、Ringle & Sarstedt 2015, 保守)：安全
- < 0.90 (Kline 2011, 宽松)：可接受
- ≥ 0.90 ：合并构念、删除其一或重新审视指标

OpenPLS 还运行 Bootstrap 版本 (HTMT-BST)，为每对返回置信区间。如果 CI 的上界排除 1.0，则构念在统计上可区分。

5.7 VIF

方差膨胀因子。两种类型：

- Outer VIF：用于形成型 LV 中的指标。检测形成型条目之间的冗余；保持 < 5 (< 3 更严格)。
- Inner VIF：用于每个内生区块中的预测 LV。检测预测变量之间的共线性；阈值相同。

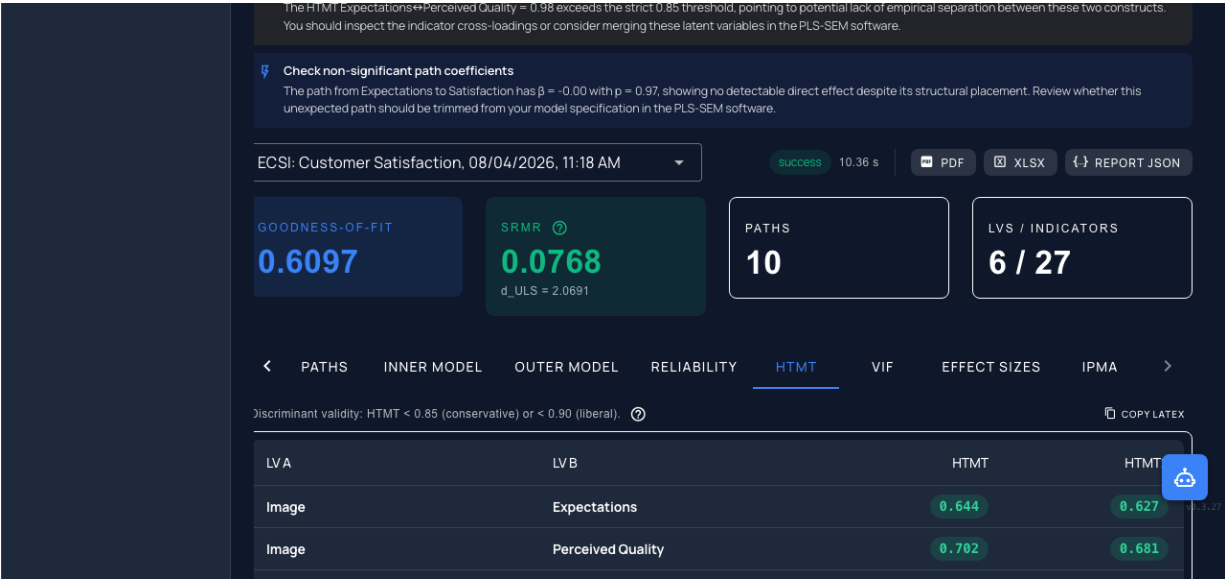


Figure 13 HTMT 子标签页：三角矩阵中的异质一同质性比率。超过阈值（0.85 保守，0.90 宽松）的单元格会被突出显示。

5.8 Effects

每个源 → 目标对的直接、间接和总效应。在报告中介效应时很有用：直接效应是你绘制的箭头，间接效应是通过 LV 的路径之和，总效应是二者之和。

5.9 Model fit

在 KPI 带下方，OpenPLS 还显示：

- SRMR 和 dULS：已在上文介绍。
- Fornell-Larcker 矩阵：HTMT 的经典替代方案。为满足区分效度，非对角线上的平方相关系数必须低于对 $\sqrt{AVE}$ 。

在 HTMT 和 Fornell-Larcker 之间，HTMT 是更新且更敏感的诊断。如果你的审稿人期待其中任一种，请两者都报

6 进阶方法

Advanced 标签页是 OpenPLS 超越基础 PLS 运行的地方。侧边导航中的三个组下共有 13 种方法：

- 推断：Bootstrap、MGA、MICOM、Moderation、Nomological validity
- 稳健性：PLSc、Gaussian copula、Common method bias
- 模型扩展：IPMA、PLSpredict、CVPAT、FIMIX-PLS、Higher-order model

每种方法的操作流程相同：在顶部配置输入，点击 Run，等待状态标记变绿，然后查看下方的结果表格。每次在 “Past runs” 下拉框中选择（图 14）。

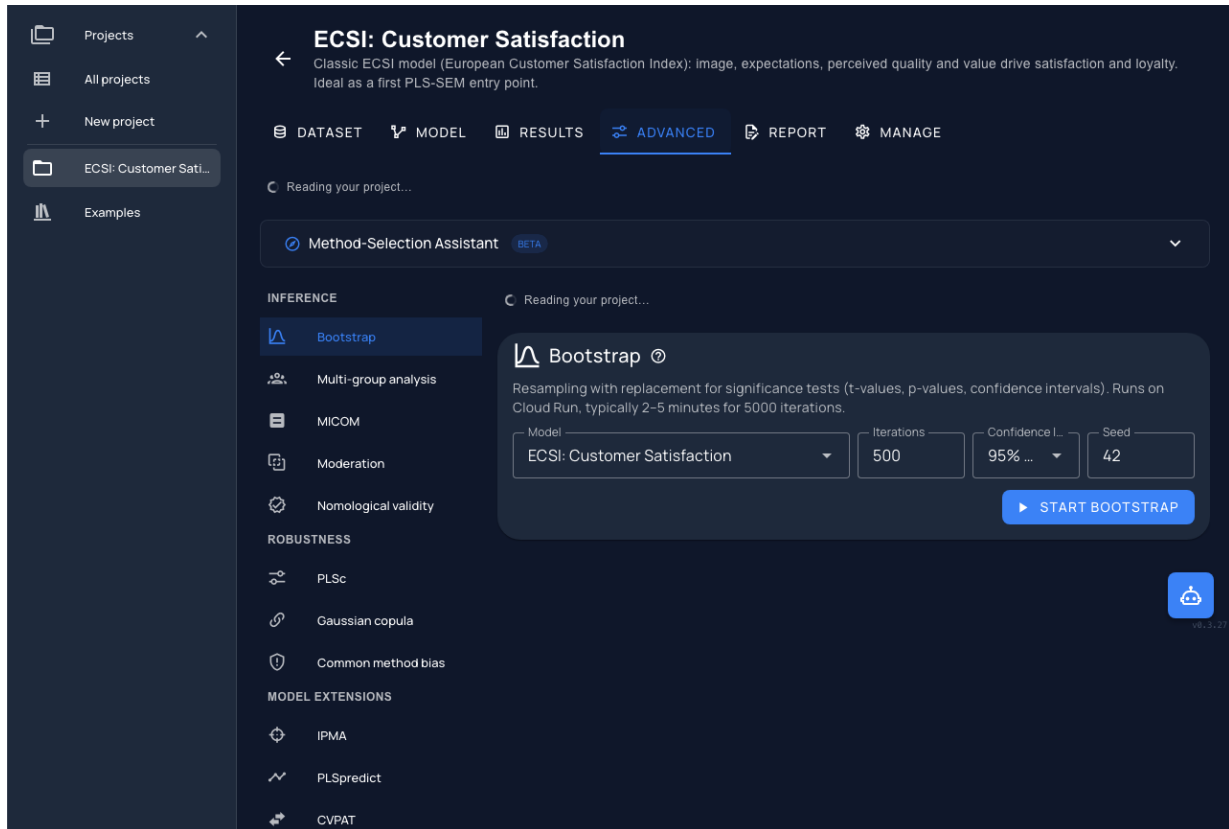


Figure 14 Advanced 标签页：侧导航中有 13 种方法，右侧是当前面板。

6.1 Bootstrap

目的。为每个路径系数和外部载荷计算标准误、t 统计量、p 值和置信区间。还为中介效应计算偏差校正加速（图 15）。

输入。

1. Iterations: 默认 500，可发表级别 ≥ 5000 。每次迭代都会有放回地重抽样并重新估计模型，所以运行时间较长。
2. Seed: 可选，设置 RNG 种子。如果需要完全可重现的标准误，请设置种子。

点击 Start bootstrap 启动。运行在 Cloud Run 服务上；面板在迭代时会流式显示进度 + ETA。

解释。对于每条路径，Bootstrap 提供了在 B 次重复上的一个分布。OpenPLS 从该分布中计算均值估计、SE、t 统计量、p 值和百分位 CI。

6.2 多组分析 (MGA)

目的。检验路径系数在不同组之间是否有差异（女性 vs. 男性、欧盟 vs. 美国、处理组 vs. 对照组）（图 16）。

输入。

1. 分组变量：将样本分组的数据集列。数值列可按范围分组；字符串列按值分组。
2. 组：每组一张卡片。对于字符串列，勾选属于每组的值（例如 Group A = {Automotive}，Group B = {Electronics}）。

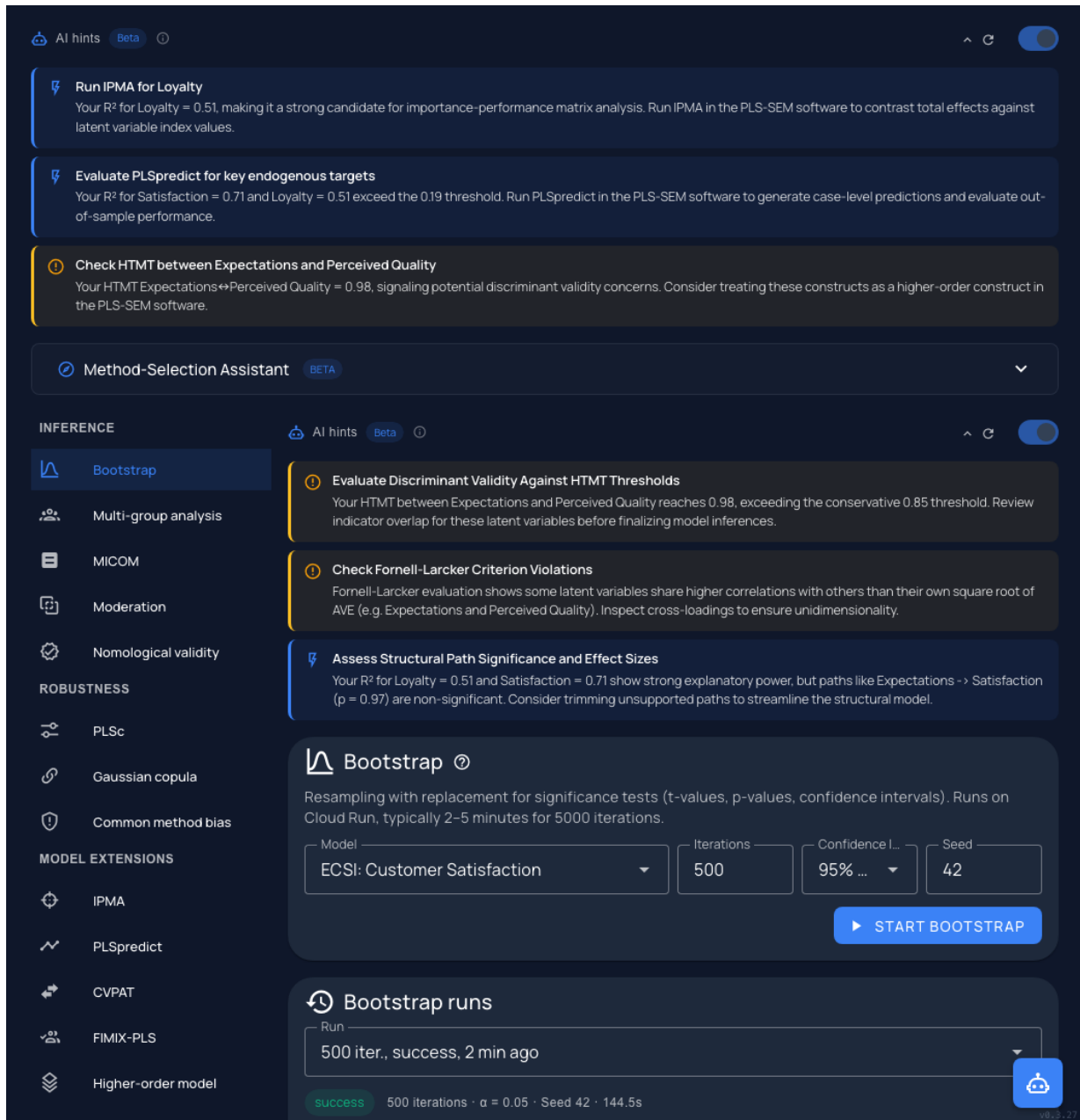


Figure 15 Bootstrap 面板，包含迭代次数、种子和 Run 按钮。

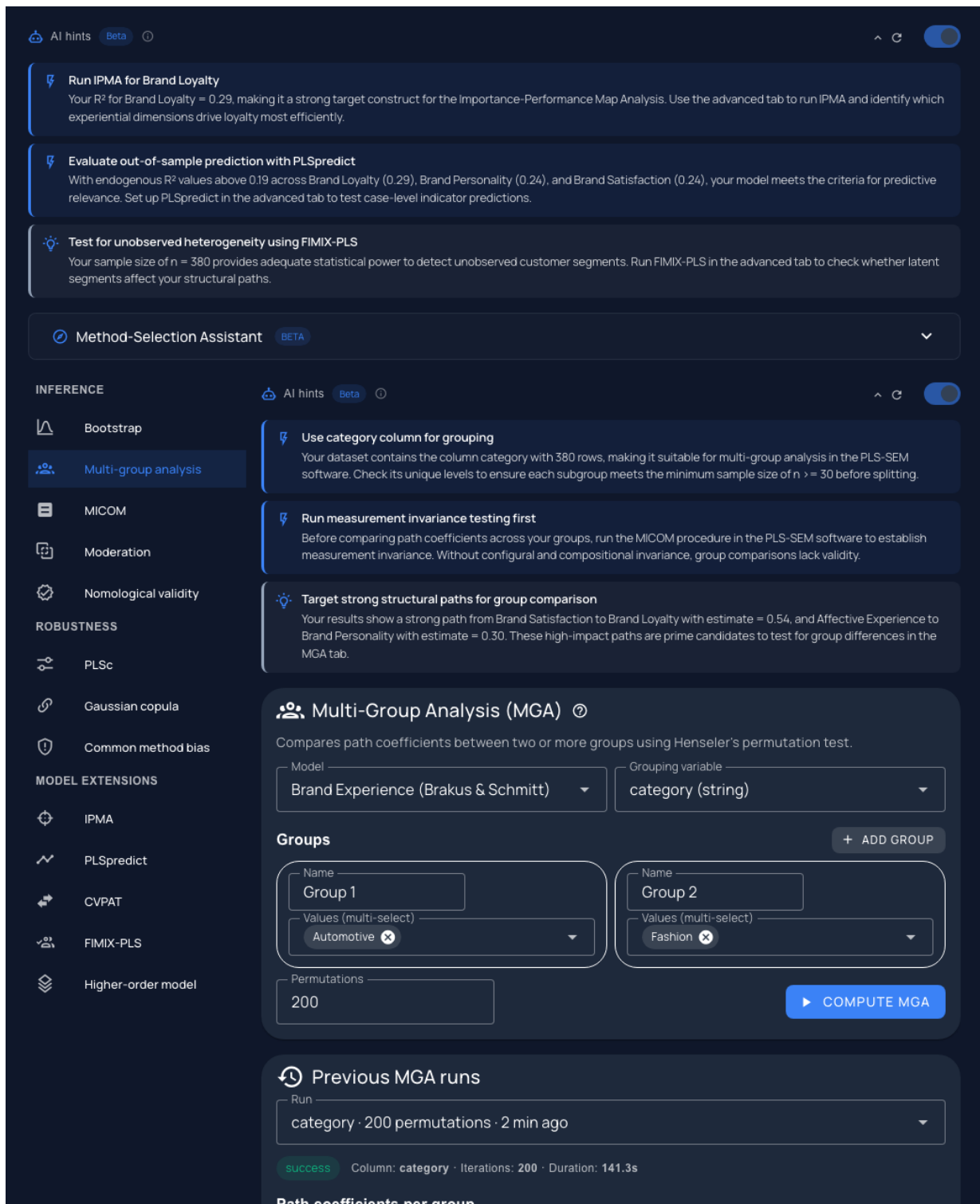


Figure 16 MGA 面板：选择分组变量列，为每组定义一张卡片，设置置换次数。

$B = \{\text{Fashion}\}$)。对于数值列，设置最小/最大值。

3. 置换次数：200 快速，1000 默认，5000+ 用于发表。MGA 对每条路径运行一次基于置换的检验。

解释。对于每条路径，表格显示每组中的估计值，以及原假设“路径在两组中相等”下的 p 值。如果 $p < 0.05$ ，则该效应在两组之间存在实质性差异。

始终先运行 MICOM。如果测量不变性在第 2 步（组合不变性）未通过，那么 LV 在每组中含义不同，比较路径就变得没有意义。

6.3 MICOM

目的。检验两组之间的组合模型测量不变性（Measurement Invariance of Composite Models）。跨组进行 MGA/调节效应分析的前提条件（图 17）。

输入。与 MGA 相同：分组变量 + 恰好两个组 + 置换次数。

解释。每个构念得出三步判定：

1. 构型不变性：相同的指标，相同的算法。OpenPLS 自动满足。
2. 组合不变性：每个 LV 的组合分数在两组之间的相关性接近 1。 p 值 < 0.05 意味着两组的组合有实质差异，结果不可解释。
3. 标量不变性：组合的均值和方差相等。如果第 2 步和第 3 步都通过：完全不变性。如果只有第 2 步通过：部分不变性：MGA 对比较路径仍然有效，只是不能比较潜在均值。

6.4 Moderation

目的。检验 X 对 Y 的效应是否取决于调节变量 Z 。使用两阶段方法（Henseler & Chin 2010）（图 18）。

输入。

1. Predictor（自变量）： $Y = X + Z + X \times Z$ 中的 X 。
2. Moderator: Z 。
3. Target（因变量）： Y 。
4. 交互 LV 名称：可选，默认为 X_x_Z 。

解释。OpenPLS 重新拟合模型，加上一个交互 LV（预测变量和调节变量组合分数的乘积）。交互效应行是关键 p 值 < 0.05 ，则调节变量改变了 $X \rightarrow Y$ 关系。

表格下方的单斜率图将此可视化： $X \rightarrow Y$ 在调节变量低 ($M - 1\text{SD}$)、均值和高 ($M + 1\text{SD}$) 处。如果斜率发散 X ；如果收敛，则减弱。

6.5 Nomological validity

目的。检查每一条假设路径的符号和显著性是否符合预期。撰写结果前的合理性检查。

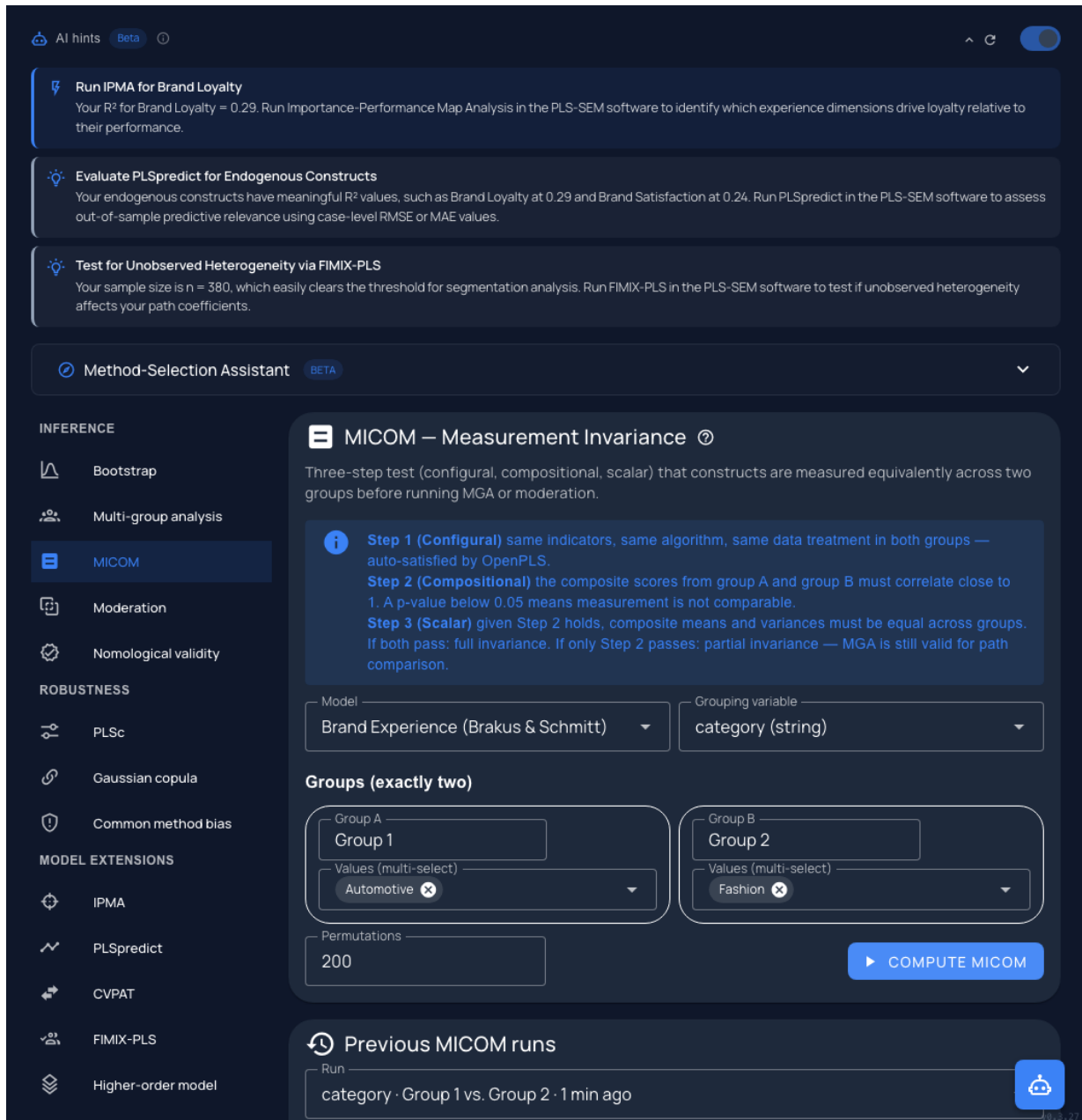


Figure 17 MICOM 面板：与 MGA 相同的分组变量和组定义 UI。

AI hints Beta

Evaluate importance-performance priorities for Loyalty
Your R^2 for Loyalty = 0.51 and Satisfaction = 0.71. Run IPMA in the advanced tab to identify which latent variables drive Loyalty with high total effects but lag in performance.

Assess out-of-sample predictive power
Your endogenous R^2 values are all > 0.33 , exceeding the 0.19 threshold. Run PLSpredict to generate case-level predictions and evaluate whether your model generalizes beyond the training sample.

Check discriminant validity for Expectations and Perceived Quality
HTMT Expectations $\leftrightarrow$ Perceived Quality = 0.98, pointing to potential collinearity. Consider testing whether they form a higher-order construct or need structural refinement.

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

Moderation analysis

Two-stage interaction effect of a moderator on a predictor $\rightarrow$ target path (Henseler & Chin 2010).

Model
ECSI: Customer Satisfaction

Predictor (independent)
Image

Moderator
Perceived Quality

Target (dependent)
Satisfaction

Interaction LV name (optional)
Default: Image_x_Perceived Quality

RUN MODERATION

Past runs

Run
Image $\times$ Perceived Quality $\rightarrow$ Satisfaction · just now

success Image $\times$ Perceived Quality $\rightarrow$ Satisfaction

Base model path coefficients

CSV

Path	Estimate
Image $\rightarrow$ Expectations	0.579
Expectations $\rightarrow$ Perceived Quality	0.848
Expectations $\rightarrow$ Perceived Value	0.105
Perceived Quality $\rightarrow$ Perceived Value	0.677
Image $\rightarrow$ Satisfaction	0.201
Expectations $\rightarrow$ Satisfaction	-0.003
Perceived Quality $\rightarrow$ Satisfaction	0.122
Perceived Value $\rightarrow$ Satisfaction	0.589
Image $\rightarrow$ Loyalty	0.275
Satisfaction $\rightarrow$ Loyalty	0.495

Simple slopes

Refit model with interaction term

CSV

Path	Estimate
Image $\rightarrow$ Expectations	0.579
Expectations $\rightarrow$ Perceived Quality	0.848
Expectations $\rightarrow$ Perceived Value	0.105
Perceived Quality $\rightarrow$ Perceived Value	0.677
Image_x_Perceived Quality $\rightarrow$ Satisfaction	-0.058
Image $\rightarrow$ Satisfaction	0.190
Expectations $\rightarrow$ Satisfaction	-0.008
Perceived Quality $\rightarrow$ Satisfaction	0.110
Perceived Value $\rightarrow$ Satisfaction	0.584
Image $\rightarrow$ Loyalty	0.276
Satisfaction $\rightarrow$ Loyalty	0.495

Figure 18 Moderation 面板：三个 LV 下拉框（Predictor / Moderator / Target）。

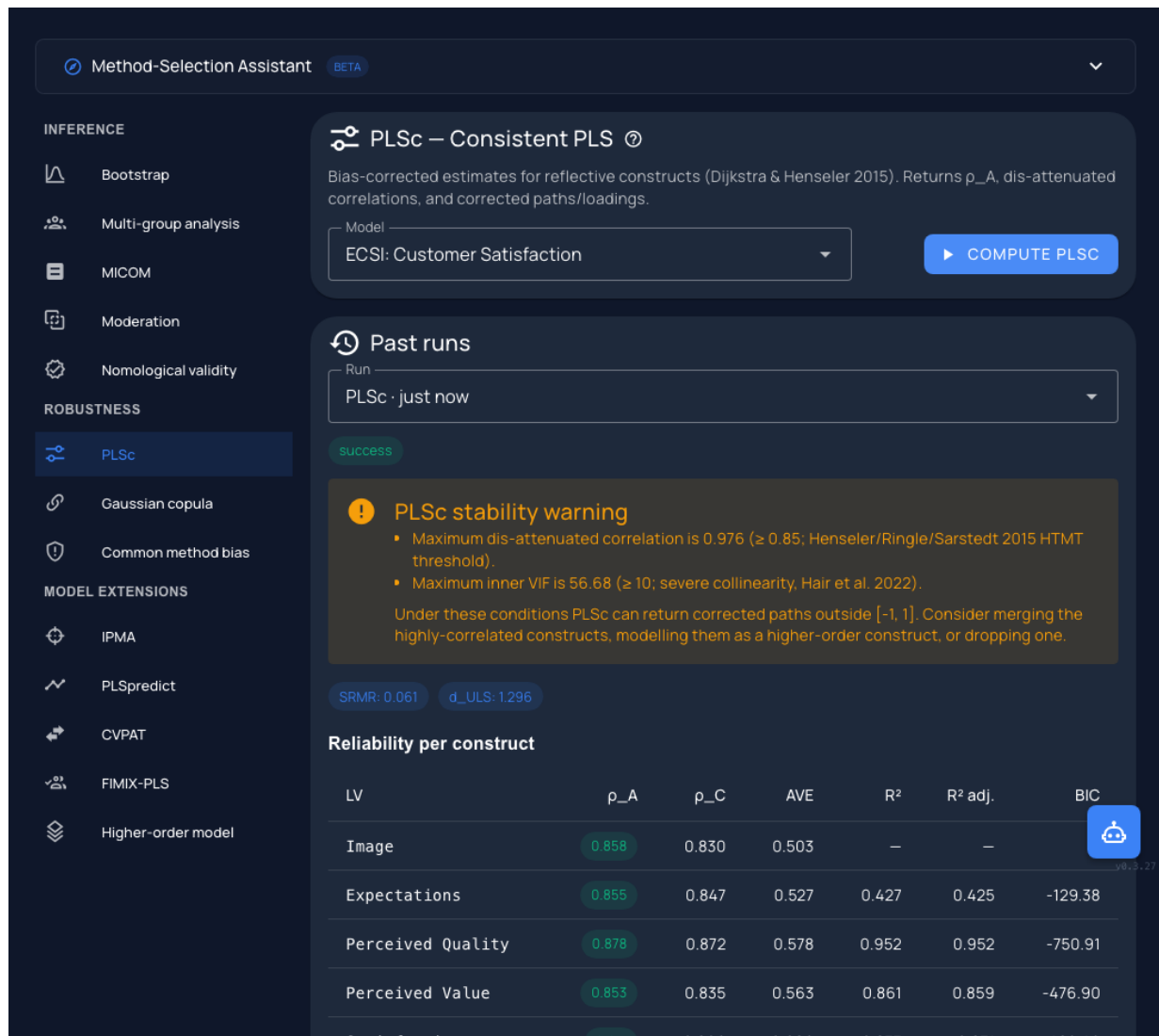


Figure 19 PLSc 面板，对仅含 Mode-A LV 的模型一键运行。

输入。根据你的模型自动填充，每条路径成为一行，带有预期符号切换：+ 表示“假设为正”，- 表示“假设为负”。为每个假设设置符号，选择显著性水平 α （默认 0.05），点击 Run test。

解释。对于每个假设，表格显示估计符号、最近一次 Bootstrap 的 p 值，以及在 α 水平下相对于你的预期符号的（通过 / 未通过）。未通过的是需要重新审视或作为零结果报告的候选。

6.6 PLSc

目的。一致 PLS 估计。当构念真正是共因子而非组合体时，纠正标准 PLS 的衰减偏差（图 19）。

输入。无，点击 Run。

解释。在应用 PLSc 校正的情况下重新运行完全相同的模型。路径系数通常略微远离零（校正消除了因有限信度导致的衰减偏差）。PLS 和 PLSc。

仅在以下情况有效。所有 LV 都是 Mode A（反映型）。Mode B LV 违反 PLSc 假设。

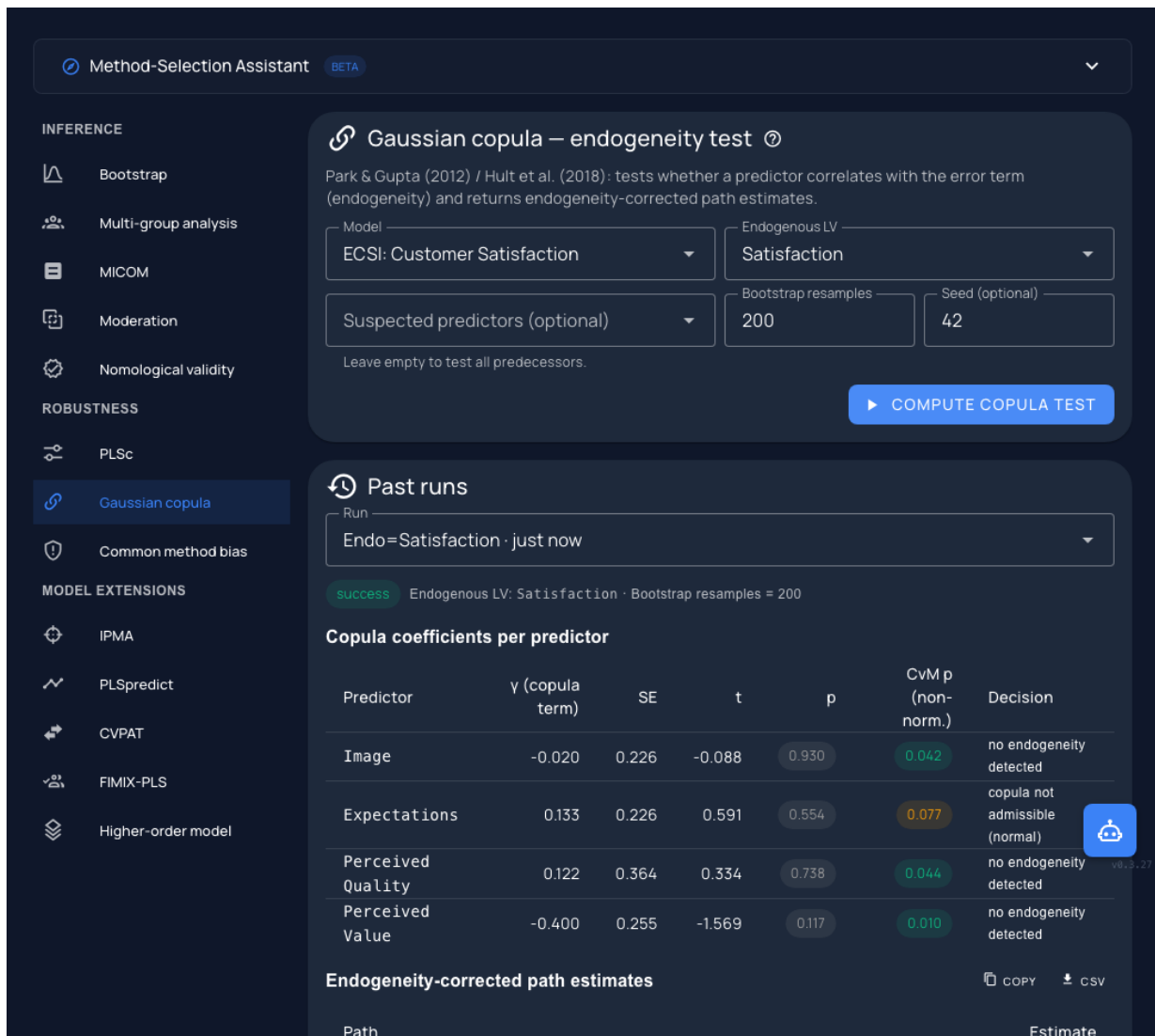


Figure 20 Copula 面板：内生 LV 选择器、Bootstrap 重抽样次数。

6.7 Gaussian copula（内生性）

目的。检验预测变量是否存在内生性（与误差项相关）。如果是，通过 Park & Gupta (2012) / Hult et al. (2018) 方法返回经内生性校正的估计（图 20）。

输入。

1. 内生 LV：选择你怀疑其预测变量与误差项相关的 LV。
2. 可疑的预测变量：可选；留空则测试内生 LV 的所有前置变量。
3. Bootstrap 重抽样：默认 500，为发表级 CI 提高到 1000-5000。
4. Seed：可选。

解释。对于每个可疑的预测变量，OpenPLS 报告一个 copula 项系数 γ 、它的 p 值以及一个 CvM（Cramér-von Mises）p 值，用于检验预测变量的非正态性（该方法所要求）。如果 copula 项显著，增强路径表会显示。仅在以下情况有效。预测变量为非正态。如果预测变量的 CvM p > 0.10，则 copula 校正不可靠。

6.8 Common method bias (CMB, Lindell-Whitney)

目的。通过 Lindell & Whitney (2001) 标记变量法检测方法偏差 (图 21)。

输入。

1. 标记变量：一个理论上不相关的数值列，且不是模型中任何 LV 的指标。示例：像 `tenure_years` 或 `age` 之类的人口学变量。
2. 显著性 α ：默认 0.05。

解释。OpenPLS 从每个构念相关性中剥离最小的观察到的标记-LV 相关性，并报告哪些路径失去了显著性。从

仅在以下情况有效。你的数据集中有一个模型外的数值列。否则，请选择不同的样本或手动运行 Harman 单因子检验。

6.9 IPMA

目的。重要性-绩效地图分析。回答“哪些构念对目标最重要，哪些表现不佳？” (图 22)

输入。

1. 目标 LV：通常是你所关心的结果 (Loyalty、Adoption……)。

解释。对于目标的每个前置变量，OpenPLS 计算：

- Importance = 对目标的总效应 (直接 + 间接)
- Performance = LV 组合分数的均值，重标度到 0-100

绘制为散点图 (Performance 为 x 轴，Importance 为 y 轴)。左上 = 重要性高、绩效低 = 最大的改进杠杆。右上 = 重要性高、绩效高 = 需要保持的优势。

6.10 PLSpredict

目的。评估样本外预测相关性。 R^2 是样本内的；PLSpredict 告诉你模型是否具有泛化能力 (图 23)。

输入。

1. k-fold：默认 10。越大 = 测试折越小。
2. Repeats：默认 1。增加到 10+ 可对不同折分配求平均并降低方差。
3. Seed：用于可重现性。

解释。对于每个指标，表格列出：

- Q^2_{predict} ：交叉验证的预测 R^2 。大于 0 有意义；越大越好。
- RMSE_PLS vs. RMSE_LM：PLS 预测 vs. 线性基准。如果大多数指标的 PLS RMSE 更低，则 PLS 相对于线性替代方案具有预测价值。

6.11 CVPAT

目的。交叉验证预测能力检验。通过对样本外损失的 t 检验，将 PLS 模型与基准 (指标平均或线性模型) 进行比较 (图 24)。

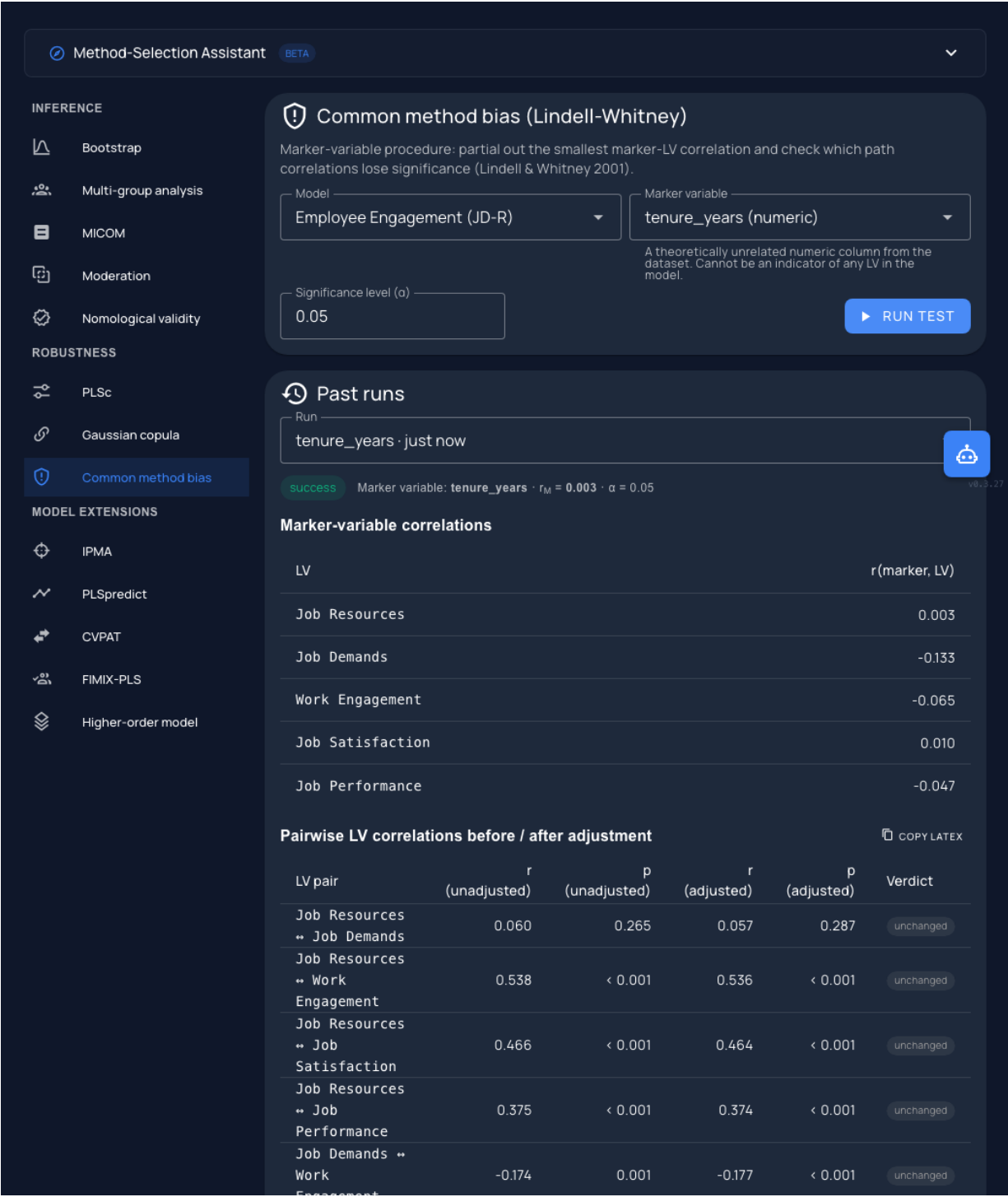


Figure 21 CMB 面板：选择一个标记列并运行。

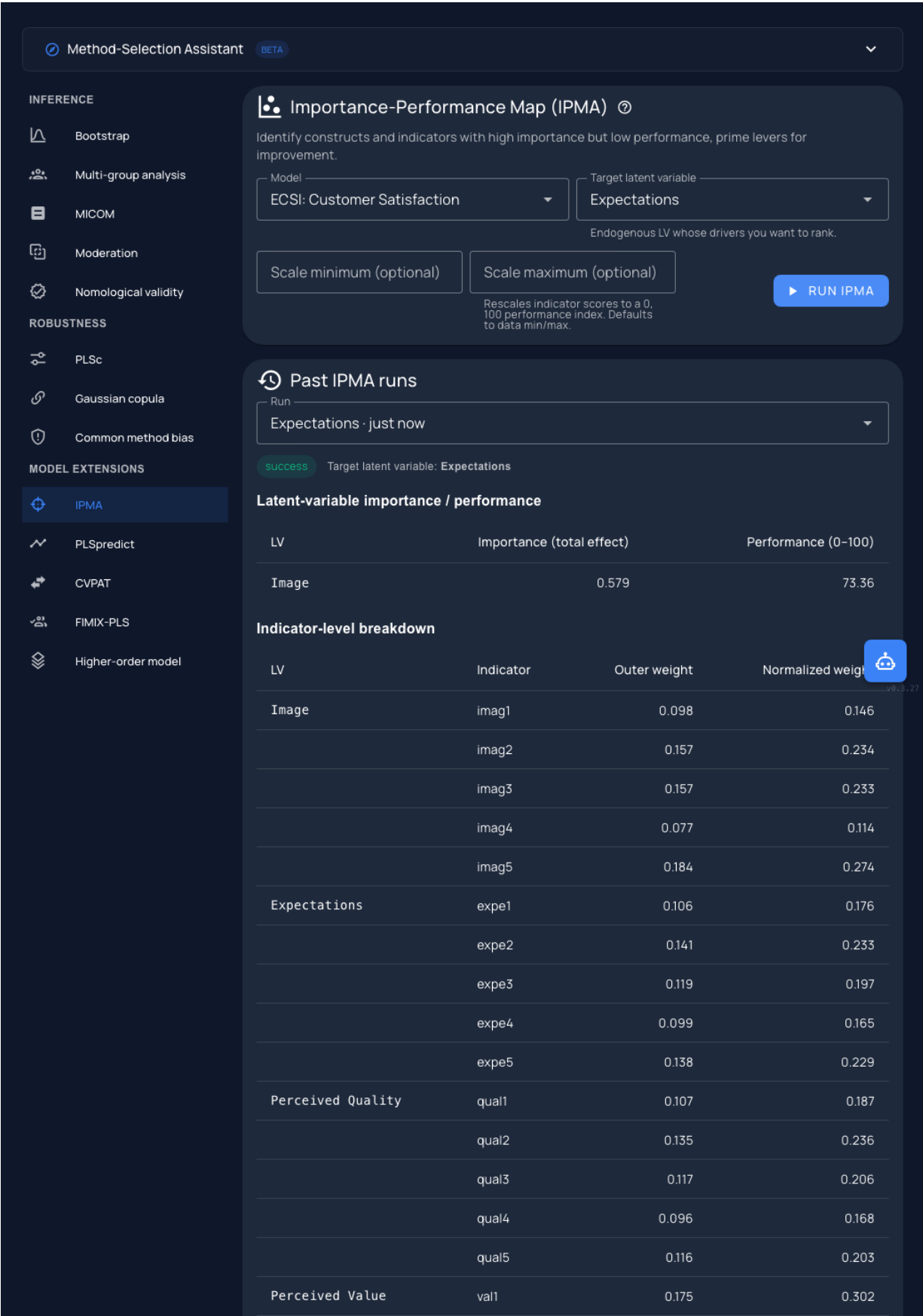


Figure 22 IPMA 面板：选择目标 LV，运行。

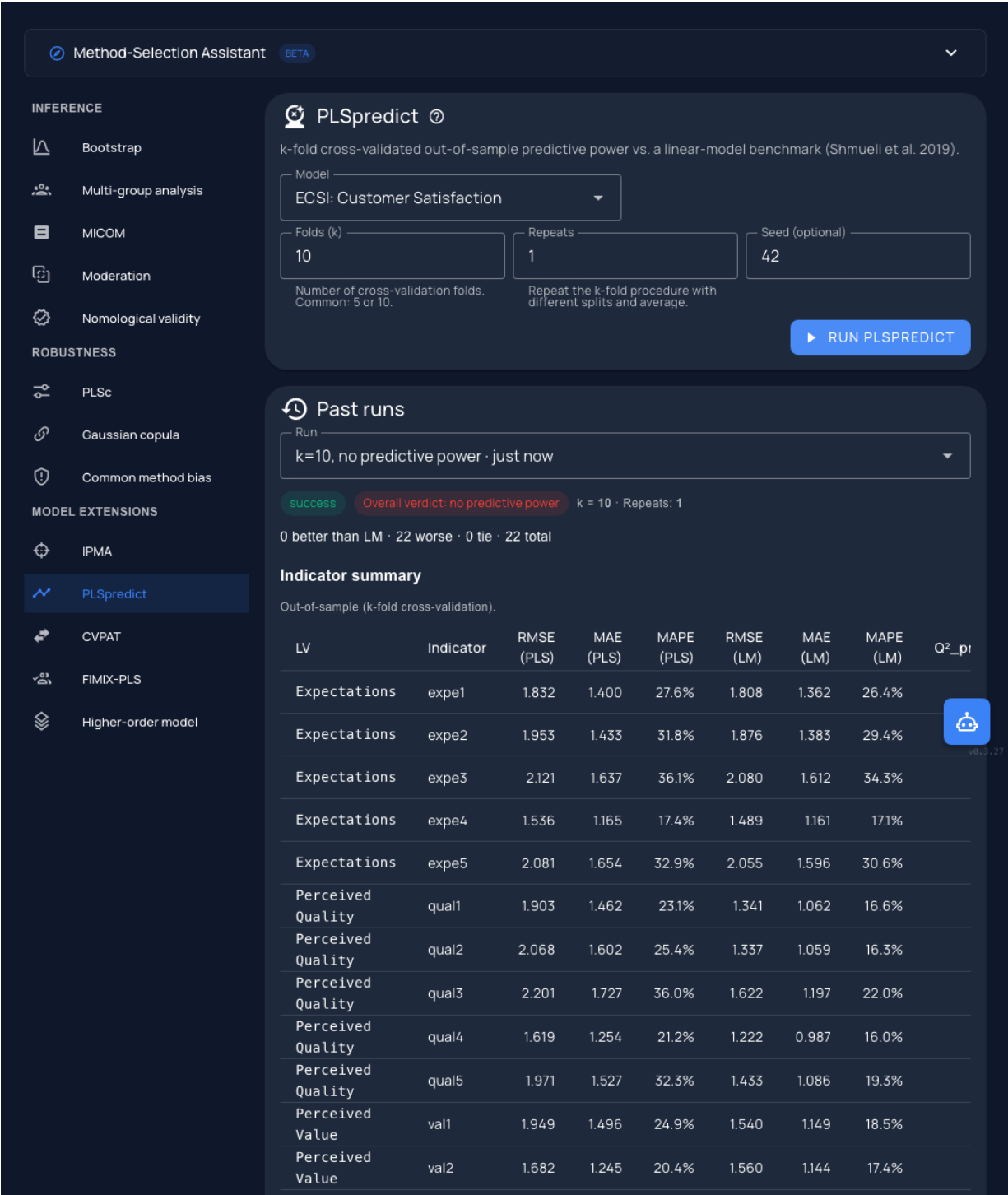


Figure 23 PLSpredict 面板: k 折、重复次数、种子。

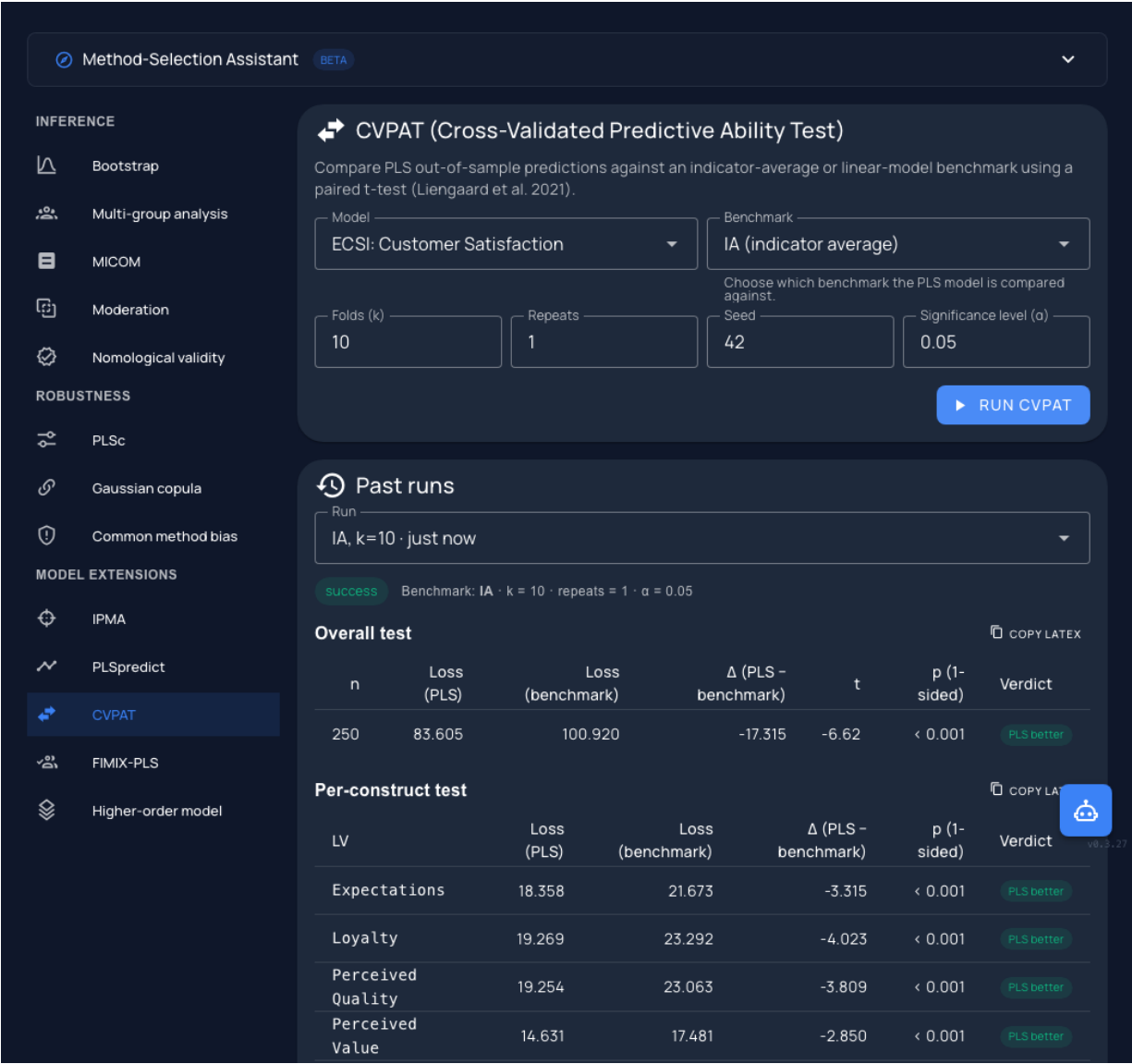


Figure 24 CVPAT 面板：基准选择器（IA = 指标平均，LM = 线性模型）、k 折、重复次数。

输入。基准 + k 折 + 重复次数 + 种子。默认值是安全的。

解释。对于每个内生 LV，OpenPLS 报告损失差异 (PLS - 基准) 及其 p 值。负的损失差异 + $p < 0.05$ = PLS 的预测优于基准，这是一个强有力的结论。

6.12 FIMIX-PLS

目的。有限混合 PLS。当你怀疑存在异质性但不知道分组时，发现被调查者的潜类别（图 25）。

输入。

1. 类别数 (K)：尝试 2..5，比较拟合准则。
2. EM 重启次数：n 次随机启动中的最佳。越多越稳健。
3. 最大迭代次数、容差：收敛控制；保留默认值。

解释。对于每个 K，OpenPLS 估计每类别的路径 + 每类别的混合比例 ρ 。跨 K 值比较 AIC、BIC、CAIC 和熵 (EN)，BIC 的拐点结合 $EN > 0.6$ 可识别出“正确的”K。

一旦决定 K，导出类别分配并按类别重新运行标准 PLS 进行详细推断。或者将其用作 MGA 中的分组。

6.13 高阶构念 (HOC)

目的。通过分离的两阶段方法 (Sarstedt et al. 2019) 将多个一阶 LV 组合成一个二阶构念（图 26）。

输入。

1. Name：HOC 的 LV 名称，例如 BrandExperience。
2. 一阶 LV：要组合的组件（至少 2 个）。
3. HOC 模式：A（反映型）或 B（形成型）。

解释。OpenPLS 运行阶段 1（普通 PLS 获取一阶 LV 分数），然后阶段 2（用 HOC 替换四个维度进行新的 PLS）。报告：

- Loadings：一阶 LV 对 HOC 的载荷 (Mode A) 或权重 (Mode B)。
- R^2 (阶段 2)：HOC 下游目标解释的方差。
- Paths (阶段 2)：重新估计模型中的结构系数。

比较 Mode A vs. Mode B：用另一种模式重新拟合，并检查哪一个产生更干净的信度和路径。

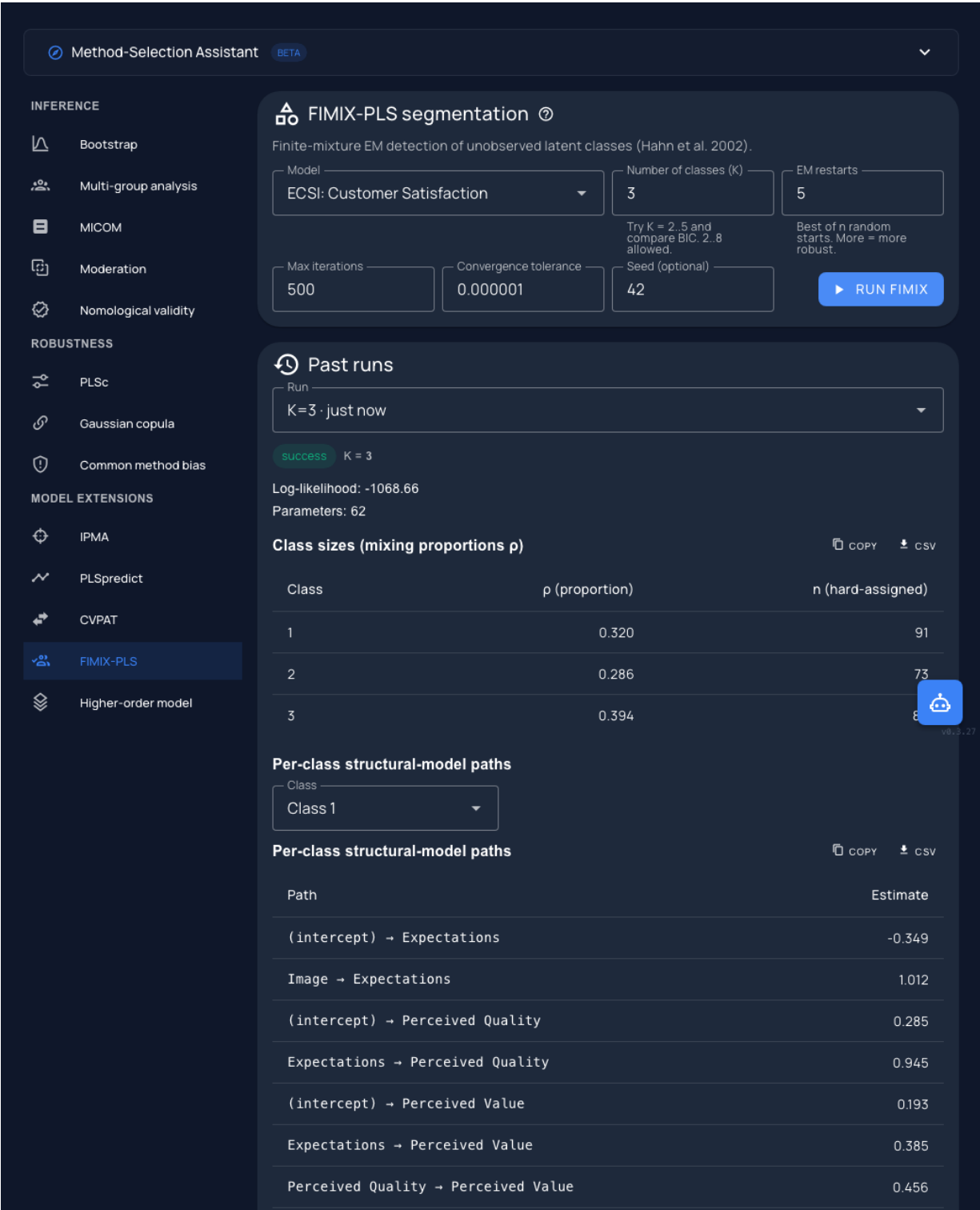


Figure 25 FIMIX 面板：类别数、EM 重启次数、容差。

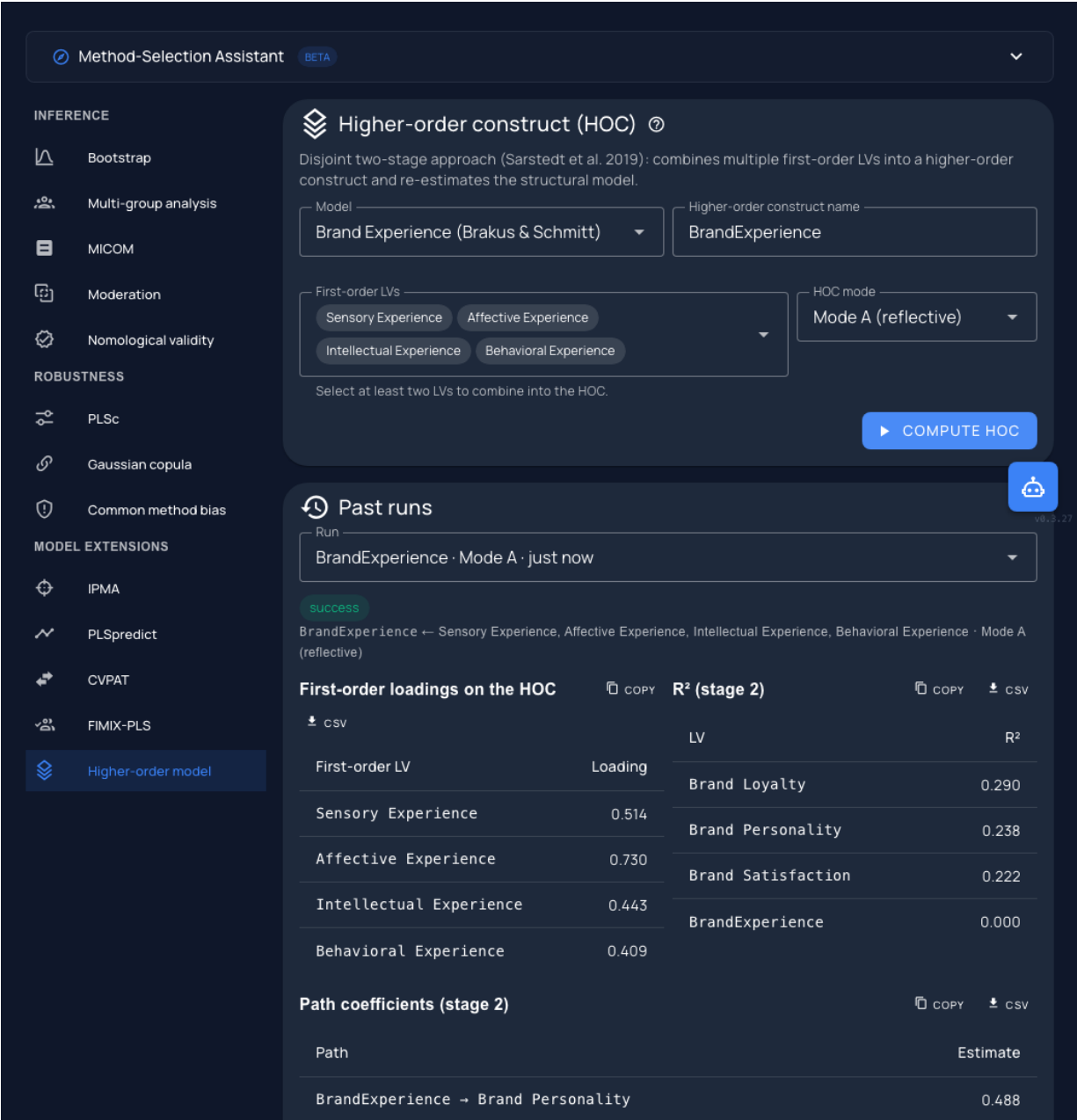


Figure 26 HOC 面板：为 HOC 命名，选择 2 个以上的一阶 LV，选择 Mode A 或 B。

7 导出

OpenPLS 中的每个结果都可以以四种格式之一离开工具。导出按钮位于 Results 标签页页头、KPI 卡片上方，也可从 Editor 工具栏中的模型导出下拉框访问。

7.1 PDF 报告

内容。一份可打印的 PDF，包含：

- 带有模型 + 数据集元数据的封面
- 路径图（根据当前 Editor 位置渲染）
- Results 标签页中的所有结果表，为打印进行了格式化
- OpenPLS 方法的参考文献（每个指标都引用定义它的论文）

使用场景。附于论文投稿、与非技术协作者共享、包含在幻灯片中。

文件名。{model_name}_report.pdf，经过清理去除空格和斜杠。

7.2 Excel 工作簿 (XLSX)

内容。一个工作簿，每个结果表一个工作表。工作表名称如文件中所显示：

- Overview: 运行元数据（计算时间、GoF、SRMR、LV/路径数）
- Inner Summary: 每个 LV 的 R^2 、调整后 R^2 、AVE、共同性
- Model Fit: SRMR + d_ULS（仅在已计算时）
- Path Coefficients: 每条内部模型路径的估计值
- Inner Model: 带有 Bootstrap SE、t、p、CI（如果可用）
- Outer Model: 每个指标的载荷（Mode A）或权重（Mode B）
- Effect Sizes f^2 : 每个预测路径的直接效应大小
- Predictive Relevance Q^2 : 每个 LV 的交叉验证冗余
- VIF, Outer 和 VIF, Inner: 共线性
- HTMT: 异质一同质性矩阵
- Reliability: Cronbach α 、组合信度、特征值
- Model, LVs 和 Model, Paths: 模型规范本身

条件工作表，仅在存在对应的进阶运行时添加：

- Nomological Validity: 每个假设的符号 + 判定
- CMB (Lindell-Whitney): 经标记校正的相关性
- CVPAT: 交叉验证预测能力检验

使用场景。审稿人要求原始数字，或者你希望按论文的排版风格重新格式化表格。

文件名。{model_name}_results.xlsx。

7.3 Report JSON 包

内容。一个自包含的 JSON，包含：

- 完整的模型规范（LV、指标、路径、位置、测量模式）
- 数据集架构（列名、类型、缺失计数）
- 每个结果表的原始内容
- 计算元数据（时间戳、时长、OpenPLS 版本）

使用场景。作为论文的补充材料。任何导入该 JSON 的人都可以在不运行 OpenPLS 的情况下检查你所估计的内部文件名。{model_name}.report.json。

7.4 模型导出（Editor 工具栏）

Editor 工具栏上有一个下载图标，会打开一个包含两种仅模型导出格式的菜单：

- JSON (**.openpls.json**)：包含位置信息的完整 OpenPLS 模型。可往返：你可以将其导入到另一个 OpenPLS 项目中，用不同的数据集重复使用该模型。
- SmartPLS (**.splsm**)：与 SmartPLS 3 和 4 兼容的 XML 文件。让你把模型交给使用该工具的协作者。

这些只是模型，不包含结果。请使用 Report JSON 包来处理结果。

7.5 LaTeX 片段

Results 标签页中的每个结果表在右上角都有一个小的 Copy LaTeX 图标。点击它，直接粘贴到你论文的 .tex 源文件中。生成的 LaTeX 使用 booktabs (`\toprule`、`\midrule`、`\bottomrule`)，与 JMS、MIS Quarterly 等期刊的传统表格风格相匹配。

LaTeX 片段在表格标题中嵌入了模型名称、数据集名称和 OpenPLS 版本，以便追溯。如果你希望更简洁的外观，请在论文中调整标题。

7.6 可重现性

对于论文来说，理想的可重现性包应包含：

1. Report JSON（结果和模型 + 数据集架构）。
2. 原始数据集（作为补充文件，如许可允许的话）。
3. 一份简短的说明，指向 `openpls.app` 作为估计工具。

任何读者都可以在不依赖你本地安装的情况下重新运行完全相同的计算。

导出的数值与你在 Results 标签页中看到的数字相同。在已发表的 PLS-SEM 数据集上与 SmartPLS 的对比正在进行中；最新结果链接自 openpls.app/validation。

8 协作

OpenPLS 将项目视为可共享的工作区。邀请合著者、研究助理或审稿人查看或编辑同一模型。每个操作都会留下记录。所有与协作相关的功能都在标签栏的 **Manage** 下。三个子视图：**Team**、**Activity**、**Settings**。

8.1 Team 面板

打开 **Manage** → **Team** 进入图 27 所示的面板。

三个部分依次排列：

1. **Team**：每个当前成员，带有姓名、邮箱、头像和角色胶囊。项目所有者（创建者）会被标记。每行右侧有操作按钮。
2. **Invite person**：邮箱输入框 + 角色下拉框 + **Invite** 按钮。通过 **Resend** 发送签名的邀请邮件。收件人可通过 `/invitations/{id}` 的链接接受。
3. **Pending invitations**：每个未决邀请一张卡片，显示被邀请的邮箱、时间戳和 **Revoke** 操作。

8.2 角色

三种角色：

- **Owner**：创建者，可以做一切，包括删除项目。
- **Editor**：可以编辑数据集、模型，运行计算，运行进阶方法，编辑设置。不能删除项目或邀请新成员（只读）。
- **Viewer**：只读。可以查看所有结果和导出，但不能更改任何内容。

角色更改立即生效并显示在活动日志中。

8.3 发送邀请

1. 在 **Email address** 字段中输入被邀请者的邮箱。
2. 选择他们的角色：**Editor, can edit** 或 **Viewer, read only**。
3. 点击 **Invite**。

OpenPLS 会在数据库中创建一个邀请条目，并向被邀请者发送带有接受链接的邮件。如果该邮箱已属于某个 OpenPLS 用户，则邀请会直接出现在他们自己的 `/invitations` 路由下。

复制邀请链接。每张待处理邀请卡片上都有一个“Copy link”图标，将接受 URL 放到剪贴板，当收件人的邮件到达时，邀请会自动接受。在待处理卡片上点击 **Revoke**。邀请变为无效；接受链接不再有效。

尚未拥有 OpenPLS 账户的被邀请者可以通过接受流程创建一个账户（邮箱会预填）。邀请为他们保留位置，

8.4 活动日志

切换到 **Activity** 子视图（图 28）查看审计跟踪。

后端为每个事件记录一个条目。当前发出的事件：

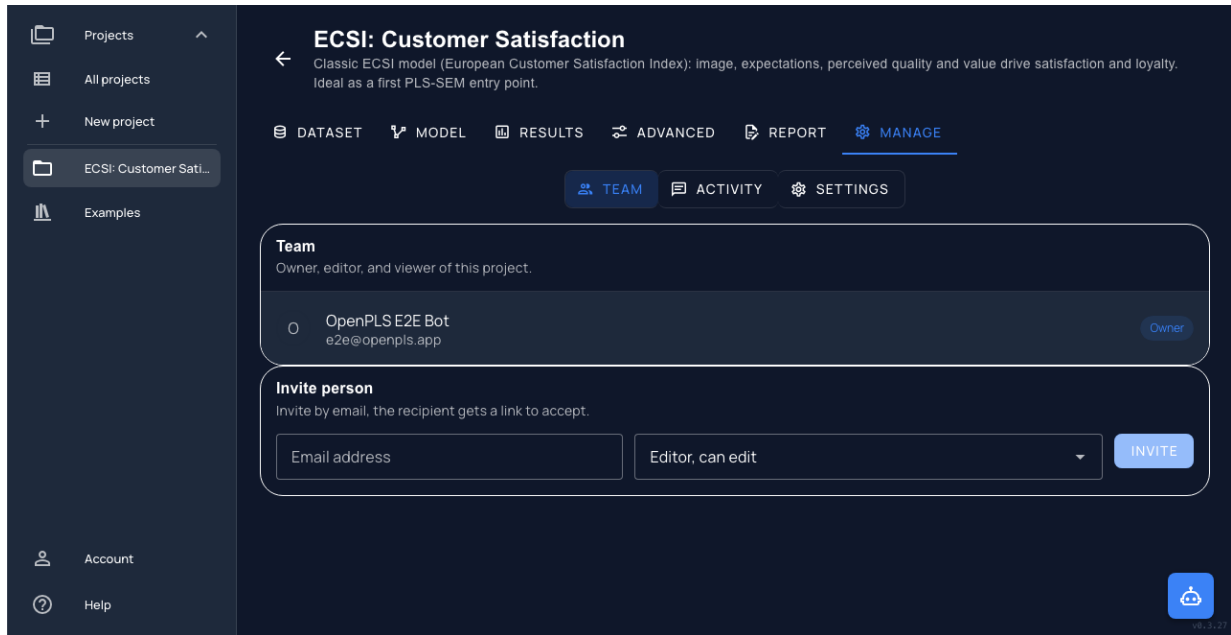


Figure 27 Team 子视图：顶部是成员列表，下方是邀请表单，底部是待处理的邀请。

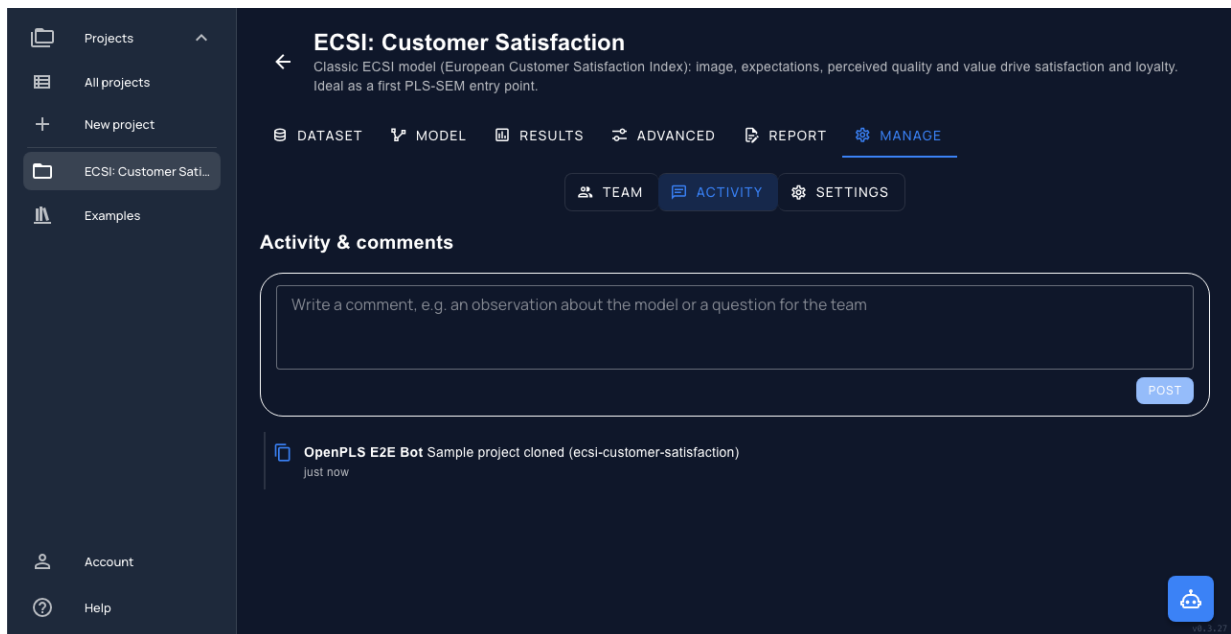


Figure 28 Activity 子视图：按时间顺序显示每个项目事件，带有图标和操作者。

- **团队事件：**成员被邀请、邀请被接受、邀请被撤销、邀请邮件失败、成员角色变更、成员离开、样例项目
- **模型计算：**一次完整的主计算完成。
- **进阶方法运行：**MICOM、Moderation、PLSc、Copula、IPMA、PLSpredict、FIMIX、HOC、Nomol

(尚) 未记录：Bootstrap 运行、MGA 运行、单个模型保存、数据集上传和数据集重新解析。Bootstrap 和 MGA 在 Cloud Run 服务上运行，直接写入其运行文档而不接触活动日志，请查看面板自身的“Past runs”历史记录。

评论。活动流还包含一个评论撰写器。评论作为顶层条目存储，并按时间顺序合并到流中，不会以线程形式挂

8.5 离开项目

Editor 和 Viewer 在团队列表中自己的行上看到 **Leave** 按钮。点击后即可从工作区删除该项目并释放席位。所有

所有权转移尚未在 UI 中公开。在当前版本中，所有者是点击 “New project” 的人，并在项目的生命周期内一直是所有者。所有权转移在路线图上，在此之前，如果你需要移交项目，请

8.6 协作者能看到什么

Editor 拥有与所有者相同的 UI，只是没有危险区操作。Viewer 可看到：

- 所有标签页：Dataset、Model、Results、Advanced、Manage
- 所有只读内容
- 导出按钮：Viewer 可以下载 PDF、XLSX、JSON
- 表格上的 Copy-LaTeX 图标

Viewer 不能：

- 编辑模型或数据集
- 计算或触发任何进阶方法
- 邀请任何人或更改角色
- 更改设置

如果 Viewer 尝试点击编辑操作，OpenPLS 会显示一个不显眼的”权限被拒绝”提示，而不是默默失败。

9 设置与语言

Manage 下的 Settings 子视图处理项目级元数据和危险区。顶栏中的语言切换器处理当前会话的 UI 区域设置，并将其持久化到你的用户配置文件中，以便在多个设备上跟随你。

9.1 项目设置

打开 Manage → Settings 进入图 29 所示的面板。

两个部分：

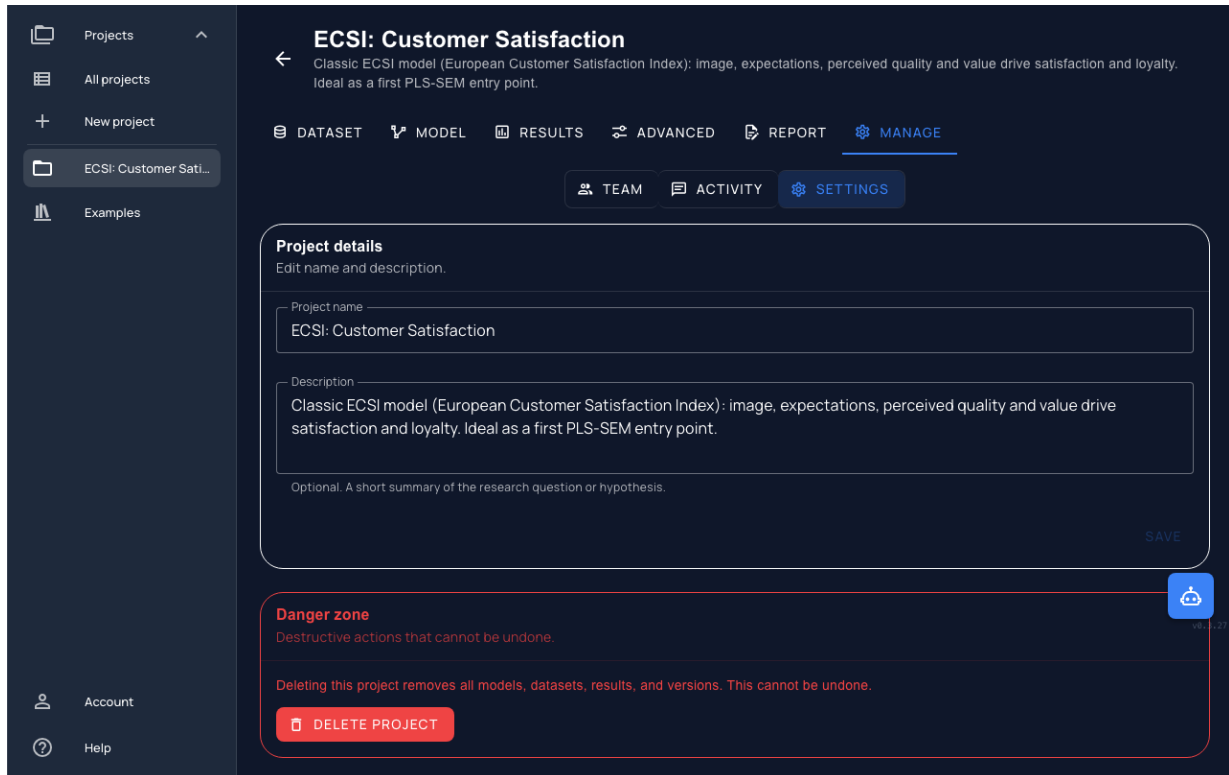


Figure 29 Settings 子视图：顶部是名称 + 描述表单，底部是危险区。

1. Project details: 与新建项目对话框相同的名称 + 描述字段，可随时编辑。点击 Save 持久化。成功提示确
2. Dangerzone: Delete project 按钮。点击它会打开一个拼出项目名称以便再次确认的对话框。删除会级

只有所有者能看到危险区卡片。Editor 和 Viewer 只能看到项目详情部分。

删除项目也会删除邀请历史和活动日志。如果你想在删除前保留审计跟踪的存档，请先导出 Activity 视图（可通过 Activity 工具栏中的下载图标）。

9.2 语言

OpenPLS 提供八种 UI 区域设置：

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

点击顶栏中的旗帜打开语言菜单（图 30）。

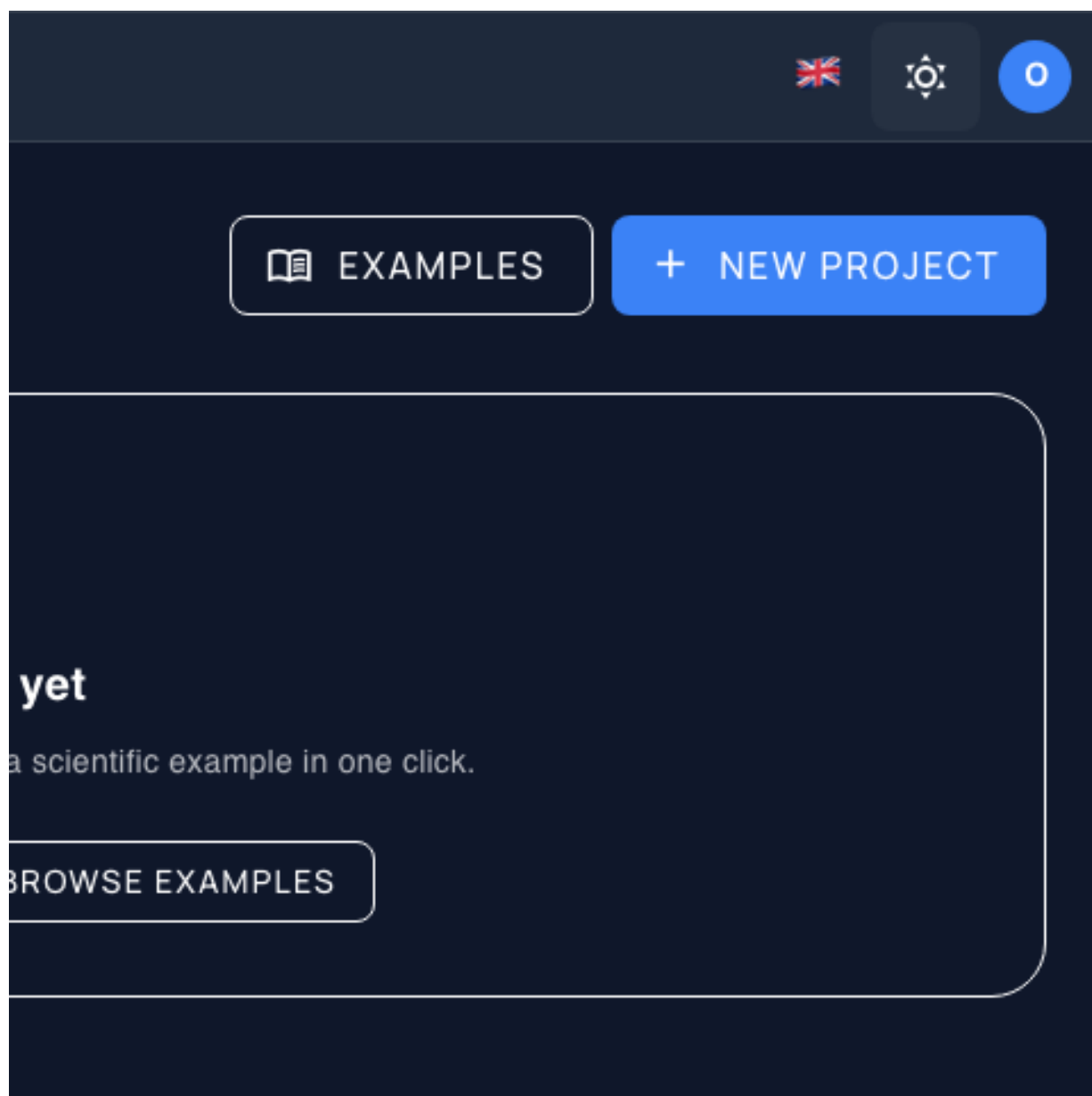


Figure 30 来自顶栏的语言下拉菜单。当前区域设置会突出显示。

选择一种区域设置。OpenPLS 会获取消息包（较小，约 30-60 KB）并重新渲染。你的选择将持久化到你的用户。哪些会被翻译。

- 所有 UI 标签、按钮、对话框、提示
- Help 页面 (/help)
- Report PDF 标题和章节标题
- 用于邀请、密码重置、验证的邮件模板

哪些不会被翻译。

- 指标值和列名：统计信息保持科学标准（基于拉丁字母的数字标签，英文列头）。这是有意为之，以便数学。
- 示例项目引用和详细描述：科学引用保持在原始出版语言。

邮件模板（邀请、密码重置、验证）目前仅提供英文，即使收件人的配置文件区域设置为其他语言。完整的

9.3 账户设置

你自己的个人资料位于 /account（侧边栏中的链接）。可编辑：

- 显示名称：显示在团队面板和活动日志中。
- 头像：点击上传；PNG 或 JPEG，最大 5 MB，裁剪成圆形。
- 首选区域设置：与顶栏切换器效果相同，但直接持久化。
- 密码：在此更改。输入一次当前密码，两次新密码；不需要确认邮件。
- 邮件偏好：一个“通过邮件接收产品更新”的开关。事务性邮件（邀请、密码重置、验证）无论此设置如何。

通过顶栏右上角的圆形头像按钮退出（在主题切换器和语言切换器旁边）。

9.4 主题

顶栏的主题切换器（太阳/月亮图标）在当前会话中在浅色和深色模式之间切换。深色模式是默认值。主题选择存储在 `localStorage` 中，因此不会跨设备同步。

10 帮助与延伸阅读

10.1 应用内帮助

OpenPLS 在 /help 提供了应用内帮助中心，可从每个项目的侧边栏或从 `cloud.openpls.app/help` 访问。图 31 显示了概览页面。

六个部分：

- How-to：面向任务的分步指南，与本指南各章节对应，但更精确。是在不加载完整 PDF 的情况下查询”
- Metrics：OpenPLS 报告的每一个指标，包含公式、解释阈值以及定义它的论文。如果你不确定为什么调 R^2 与 R^2 不同，请从这里开始。

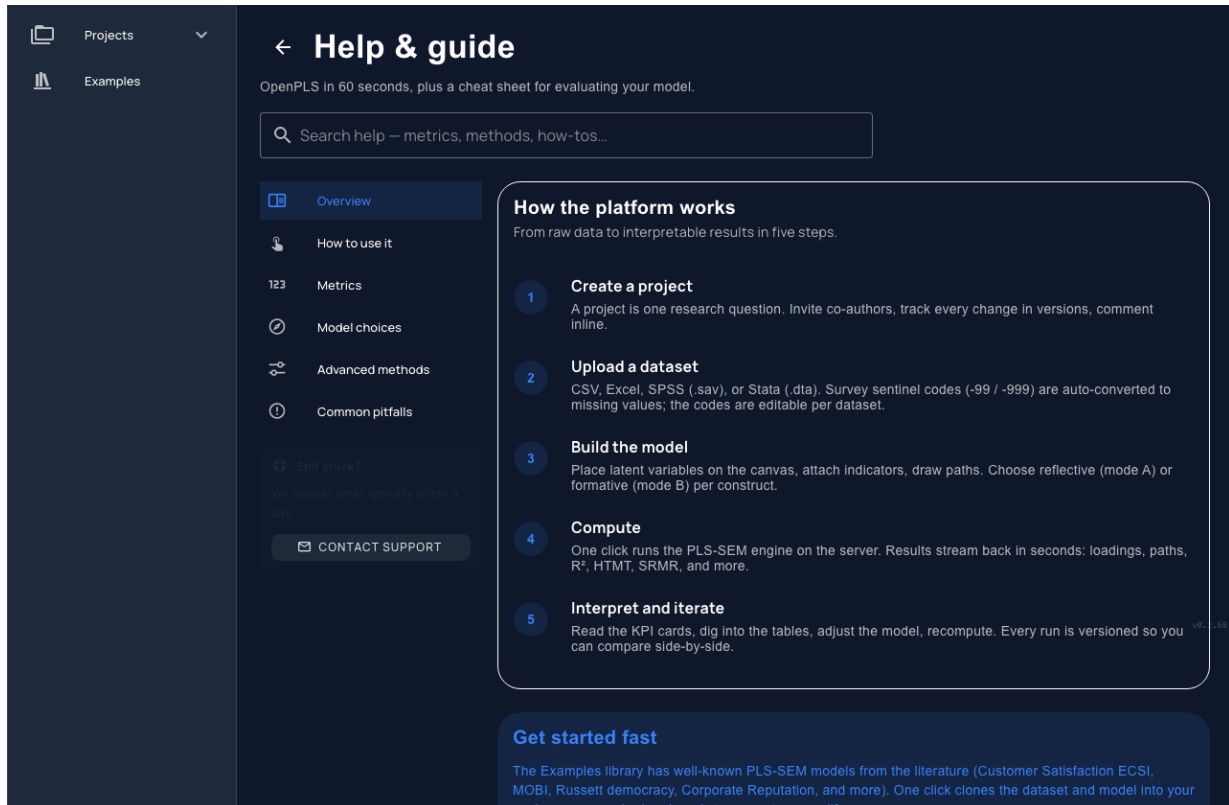


Figure 31 Help 概览：进入操作指南、指标、方法选择、进阶和常见陷阱的入口。

- Choose：常见问题的决策树：“我的构念何时是反映型 vs. 形成型？”、“哪种进阶方法适用于我的假设应该有多大？”。
- Advanced：每种进阶方法一个页面，内容与本指南第 6 章相同，附加每个方法的示例和解释演练。从 Advanced 面板上的小 ? 图标即可跳转到匹配的帮助用户页面锚点。
- Pitfalls：我们在支持工单中反复看到的建模错误：哨兵值吞没有效观测、错误的测量模式、比较不可比较

帮助用户页面已在所有八种 UI 语言中本地化。

10.2 OpenPLS 引擎（开源）

支持 OpenPLS 的估计代码是一个名为 `openpls-engine` 的开源 Python 包，以 GPL-3.0 许可发布，地址是 <https://github.com/jojacobsen/openpls-engine>。你在云端计算的内容也可以本地运行：

```
pip install openpls-engine
```

该包包含：

- 核心 PLS-SEM 估计器（在已验证的案例上与 SmartPLS 4 的输出匹配到 4 位小数）
- 第 6 章描述的每种进阶方法
- 用于可重现批处理运行的 CLI + Python API
- 用于开箱即用的自托管的 Docker 镜像

如果你想自托管整个 OpenPLS UI，引擎是计算组件；有关 Docker Compose 编排的当前状态，请参阅仓库上

10.3 支持、Bug 和功能请求

对于云端应用（cloud.openpls.app）相关的任何事宜，请写信至 hello@openpls.app。请附上你的项目 ID（在 URL 中可见），如果是报告 Bug，请附上简短的期望结果与实际结果说明。

引擎特定的 Bug 和功能请求也可以在 <https://github.com/jojacobsen/openpls-engine/issues> 作为 GitHub issue 提交，或者如果你喜欢也可以通过邮件。

10.4 引用 OpenPLS

如果 OpenPLS 在你的论文中产出了结果，请引用：

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

此引用涵盖估计引擎和云端应用，因为它们共享同一个经过验证的实现。

11 AI Companion

OpenPLS 内置了 AI Companion：一个应用内助手，它读取你的项目状态，引用同行评审的方法学论文，并帮助完成 PLS-SEM 中那些点击按钮无法完成的部分：选择下一步运行什么、起草论文文稿、预判审稿人会提出的质疑。

Companion 是默认关闭的。在你接受同意条款之前，什么都不会被调用；你也可以随时在 Account 设置中禁用它。本章按你通常遇到它们的顺序介绍每个界面。

每个 AI 响应都基于两个来源：(1) 你项目中实际的数字（数据集形状、模型规范、最新计算结果），助手通过一个精心策划的开放获取 PLS-SEM 论文库，其中的摘要通过语义相似度检索匹配你的问题。助手对你的项目

11.1 启用

首次用户登录后会看到一个一步式的同意对话框。它总结了 Companion 会读取什么（你的项目数据）、会向谁分享（100 轮的软性上限，防止失控使用）。

如果你之前关闭了该对话框，请打开 Account - AI Companion 并在那里切换。同一页面还带有 AI hints 的子切换（见下文）：Companion 开启但 hints 关闭时，只有主动特性保持静默——聊天、模型设计器

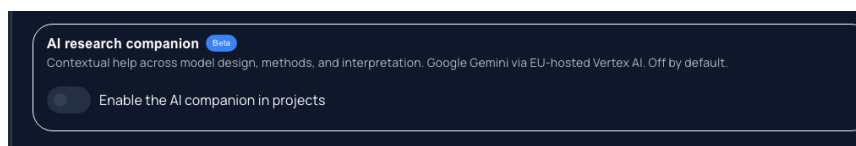


Figure 32 AI Companion 同意对话框。阅读四条要点，然后点击 Enable AI Companion。跳过则保持关闭；你稍后可以在 Account 设置中启用它。

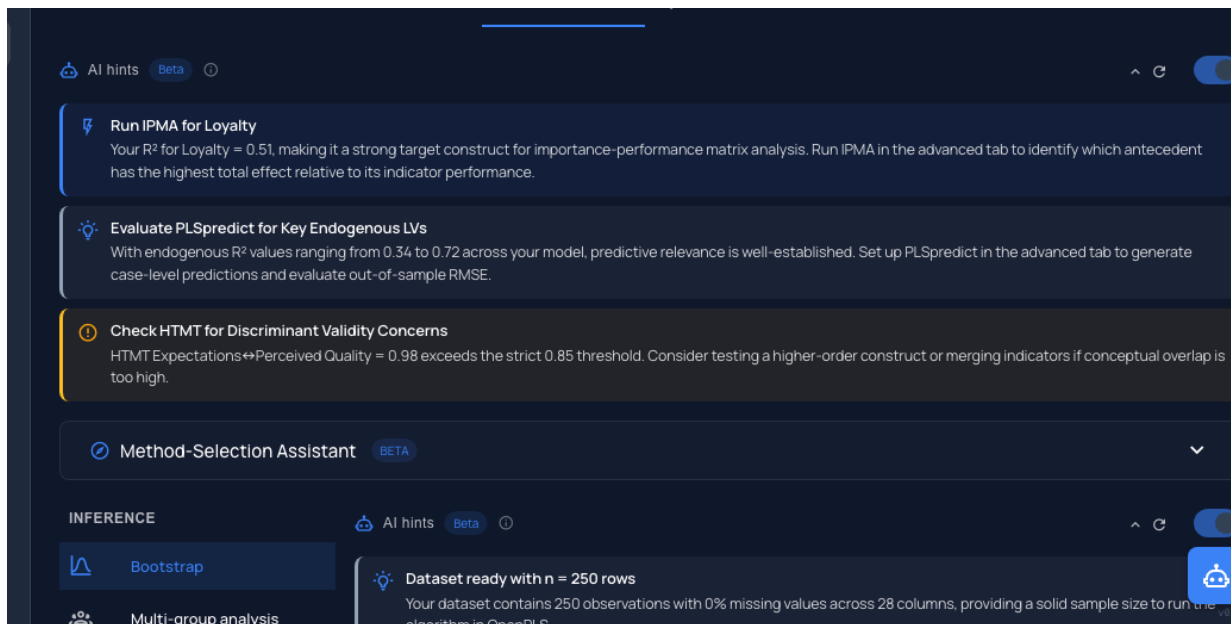


Figure 33 Advanced 标签页上的 AI hints 条带。页头带有 Beta 标记、信息提示、刷新、折叠以及开关切换。每条 hint 都有严重程度图标：灯泡表示信息，警告圈表示警告，

11.2 AI hints: 常驻可见的条带

每个有合理下一步的标签页顶部都会显示一个小小的 AI hints 条带：根据你在此项目中当前的情况，提供一到三个 Dataset 标签页上，它可能会标记一个有 22% 缺失值的列；在 Results 上，它可能会注意到你的 Loyalty R^2 为 0.42 并建议运行 IPMA；在 Advanced 上，它会指向根据当前状态最有说服力的方法。

Hints 每个标签页每个会话获取一次，然后缓存。点击刷新可强制重新读取，例如刚导入了新数据集或编辑了模型。你可以单独关闭 hints。并非每位用户都想让条带出现在眼前。条带右侧的切换会翻转每用户的 `aiHintsEnabled` 标志：hints 停止获取，但 Companion 的其余部分（聊天、设计器、报告）在你打开时仍能工作。关闭 → 条带折叠为一个不显眼的占位符。

11.3 Chat: 浮动抽屉

一个小机器人按钮永久位于右下角。点击它打开聊天抽屉。

空状态提供三个预制提示：

1. Explain my results in plain language. 读取最新的计算输出，写出针对你具体数字校准过的一段话总结。
2. Which advanced method should I run next? 应用与 Method-Selection Assistant（见下文）相同的规则。
3. Are my constructs valid? 遍历信度 + HTMT + AVE 表并按名称标记任何阈值违规。

或者输入你自己的问题。关于项目的任何事情 — 模型设计、数据集质量、方法选择、路径解释、审稿人质疑 — 都在范围内。Companion 有四个只读工具：

- `getProject`: 名称、描述、成员数
- `getModel`: LV 及其模式、指标、结构路径
- `getDataset`: 行列数、每列统计信息（缺失情况、min/max/mean），永远不包含实际行的值

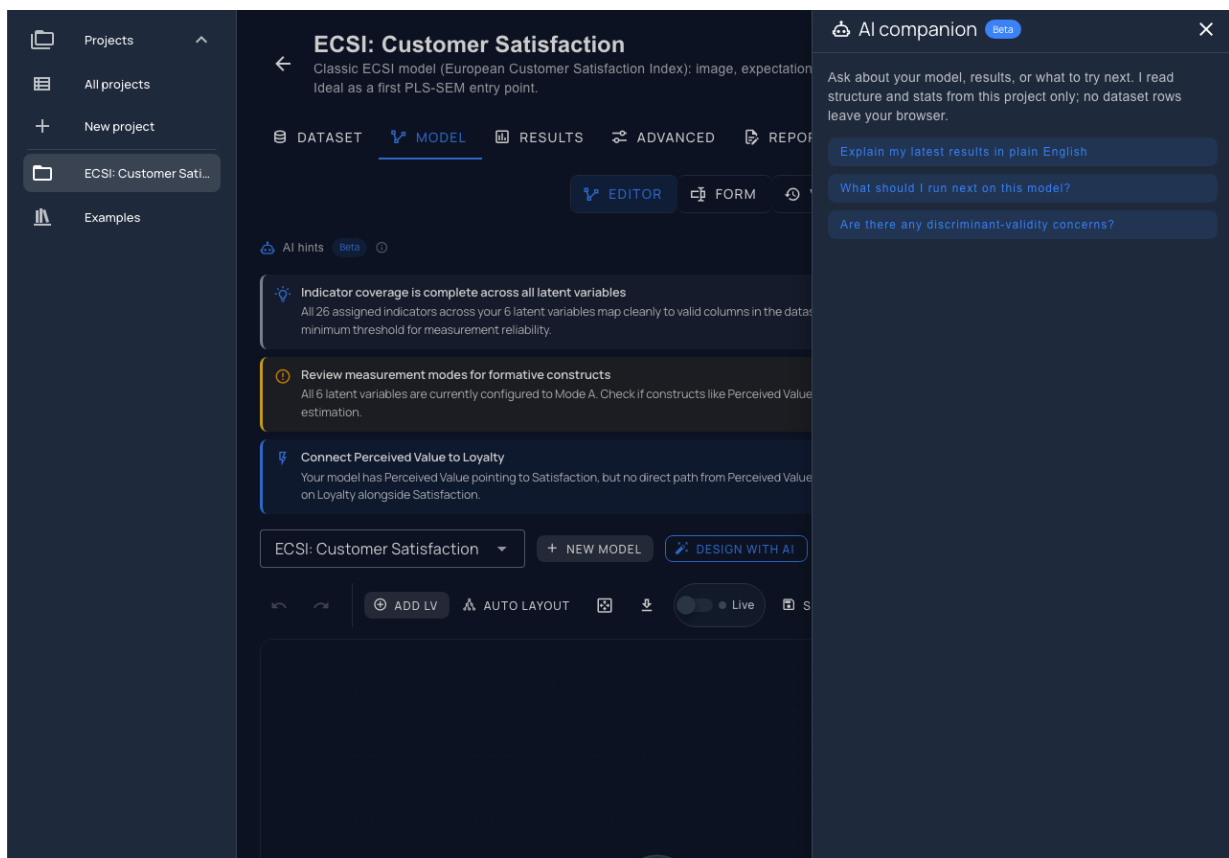


Figure 34 打开在 Results 标签页之上的 AI Companion 抽屉。页头显示标题和 beta 标记，供首次使用者的建议卡片提供三个预制问题，底部撰写器带有实时配额计数器。

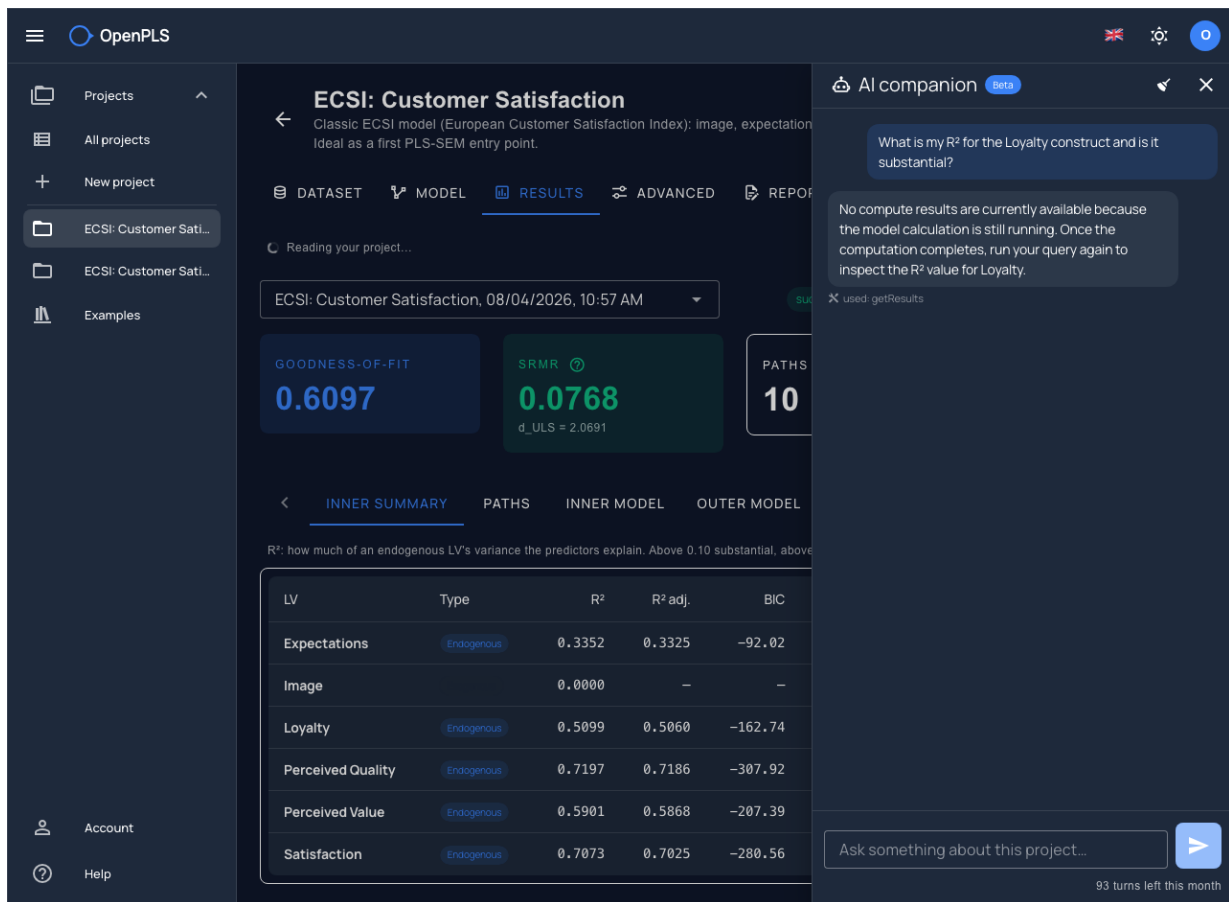


Figure 35 带有 `used-tools` 行的聊天回合，可见于回答下方。Companion 在回答前调用了 `getResults` 以获取实际的 R^2 、 β 和 p 值。

- `getResults`: 每个 LV 的 R^2 、带 p 值的路径系数、HTMT 矩阵、信度、Fornell-Larcker 判定、SRMR、GoF

每个回答下方的 `used-tools` 行告诉你调用了哪些工具。没有其他内容离开项目。

引用。当回答中的某项主张借鉴了方法学文献时，Companion 会写一个带方括号的角色名称（例如 `[henseler-2015-htmt]`）。前端将其渲染为可点击的胶囊标记，显示作者和年份（例如 Henseler+2015）。悬停会显示完整标题；点击会在新标签页中打开 DOI 或出版商页面。

这些胶囊背后的论文库是一个精心策划的约 100 篇开放获取 PLS-SEM 参考文献集合（见下文的论文库）。

历史 + 持久化。聊天回合按用户保存在账户设置存储中。重新加载页面，打开抽屉，你的对话仍然在那里。页头 `Clear conversation` 图标会清空它。

配额。Beta 期间每用户每月 100 次聊天回合。撰写器底部的计数器显示剩余额度。模型设计和报告起草回合共

11.4 从研究问题设计模型

如果你的模型还不存在一全新项目，未绘制 LV 一打开 Editor 并点击 `Design with AI`。对话框会要求填写你的研究问题。你会得到一份结构化的建议：

- 3-6 个潜变量，带有模式（A = 反映型，B = 形成型）、简短的论证以及每个 LV 3-5 个指标提示（描述该构

The screenshot displays the OpenPLS software interface for configuring a PLS-SEM model. The main workspace is titled "ECSI: Customer Satisfaction" and shows the "MODEL" tab. The "Latent variables" section lists six variables: Image, Expectations, Perceived Quality, Perceived Value, Satisfaction, and Loyalty. Each variable is configured with a name, mode (A, reflective), and indicators. The "Structural model (paths)" section shows a flow from Image to Expectations, Satisfaction, and Loyalty, and from Expectations to Perceived Quality, Perceived Value, and Satisfaction. On the right, the "AI companion" panel is open, displaying a chat interface with a question about R-squared for the Loyalty construct and a detailed answer about common method bias (CMB) and its remedies, including procedural and statistical assessments. The chat interface includes a text input field and a send button.

Figure 36 聊天回答中内联渲染的胶囊引用。每个胶囊都链接到原始论文的 DOI 页面。

Figure 37 AI 模型设计器对话框。输入研究问题（必填），列出已想到的构念（可选），选择学科。点击 Submit 接收结构化建议。

— 不是完整的问卷措辞）。

- 5–12 条结构路径，每条带有明文假设和简短的文献注释。
- 2–3 个蓝图模型，来自具有相同结构的已发表 PLS-SEM 文献。
- 注释：需要考虑的注意事项或替代规范。

在你点击 Apply 之前，该建议不会被写入你的项目。预览它，如有需要请调整你的提示（Refine 会重新生成），或者丢弃并重新开始。

应用后，模型会作为一个普通的可拖动画布出现在 Editor 中，你可以重命名 LV、添加路径、重新分配指标。建

11.5 设计问卷

一旦你有了模型（LV 已分配好模式），OpenPLS 就可以生成完整的条目集：每个 LV 4–6 个调查条目，包含措辞、量表、反向编码标志，以及当 LV 映射到文献中已知构念时对已验证量表的引用（信任、满意度、忠诚度、TAM 的感知有用性等）。

打开 Model 标签页，然后打开 Form 子视图。如果任何 LV 缺少条目，你会看到一个 Design items with AI 按钮。点击它打开问卷设计器对话框。

结果是每个 LV 的条目表：条目 ID、措辞、反向编码标志、来源引用。浏览一遍，微调措辞，删除或添加条目，一条目会作为指标落到每个 LV 上。

从这里，你有两个真正收集数据的选项：构建一份问卷让受访者在公开 URL 上填写（见下一章，Public surveys），或者将条目集带到外部调查工具并将回复作为 CSV 导入。

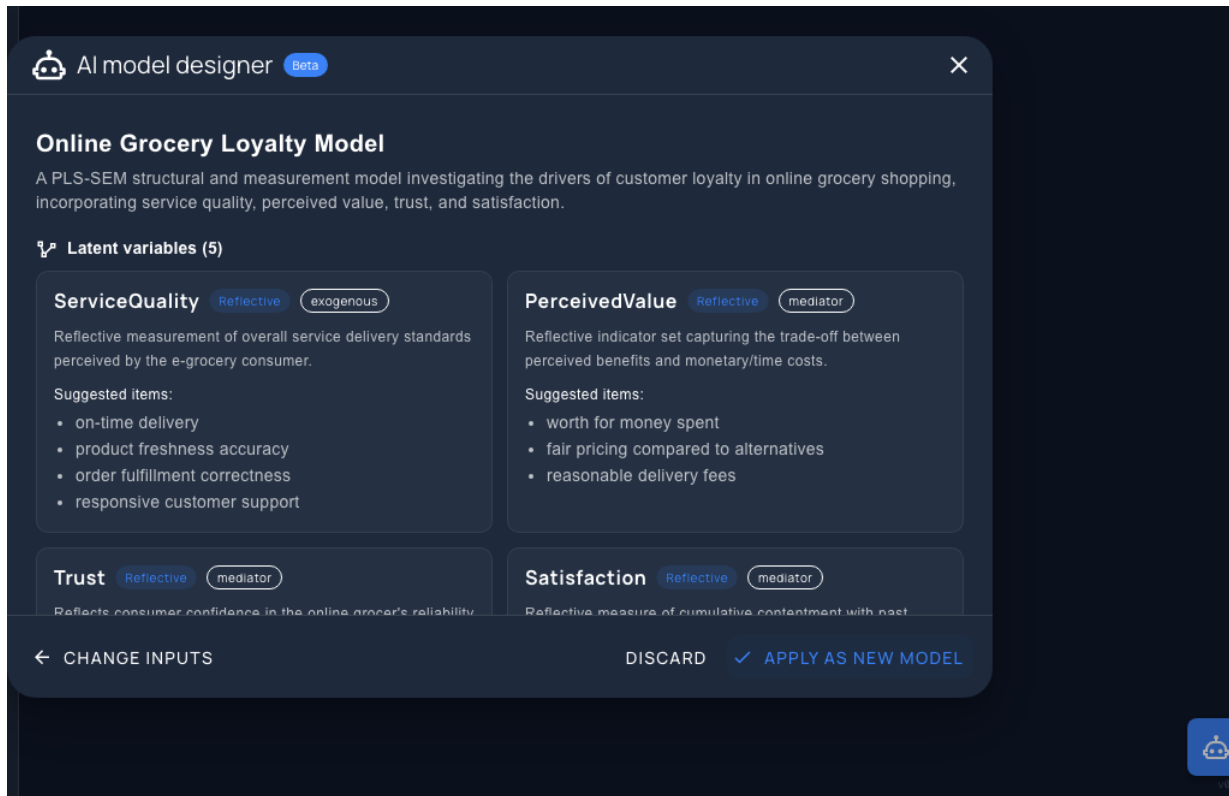


Figure 38 建议卡片。每个 LV 显示模式、指标提示和论证；路径列出源 → 目标以及假设。Apply 将模型作为新版本写入你的项目。

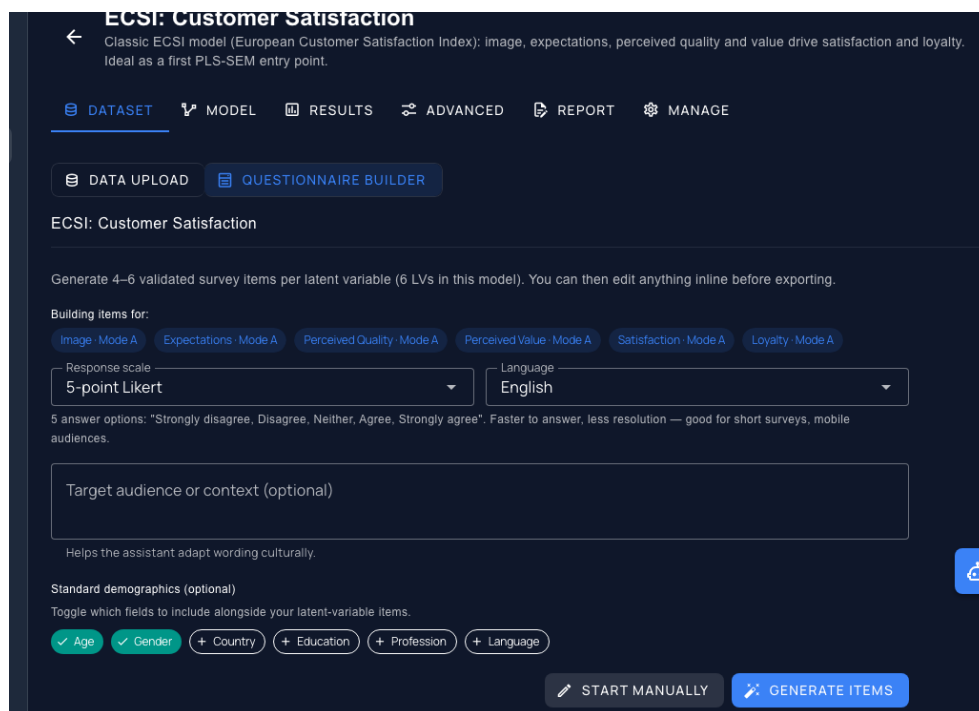


Figure 39 AI 问卷设计器对话框。选择响应量表（5 点或 7 点 Likert）、受众语言，并点击 Submit。对话框显示它将为哪些 LV 构建条目；已经有条目的 LV 会被跳过，除非勾选 Overwrite existing。

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.

DATA SET 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.

t1 ×

imag2

I consider this organization to be stable and trustworthy.

t1 ×

imag3

This organization stands for quality and customer focus.

t1 ×

imag4

The organization contributes positively to the community.

t1 ×

imag5

Overall, this organization has a poor public image.

t1 ×

+ 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.

Figure 40 生成的问卷卡片。每个 LV 块列出条目及措辞 + 反向编码标记 + 引用。可就地编辑，然后点击 Apply 写入你的模型。

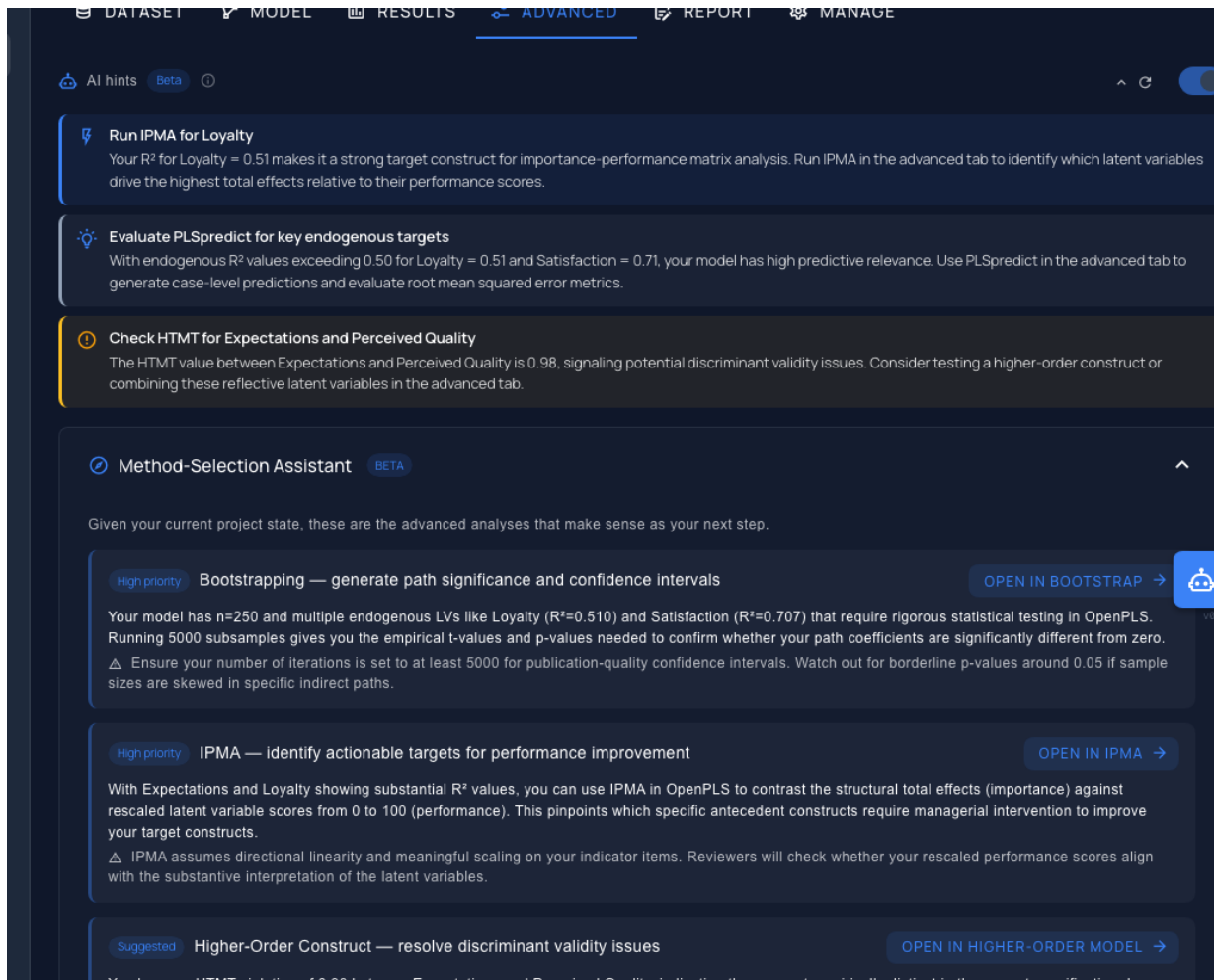


Figure 41 Method-Selection Assistant 展开。优先级标记 (High / Suggested / Optional)、每个方法一句话的理由、风险说明以及一个 “Open in ...” 按钮，用于导航到对应的 Advanced 子标签页。

11.6 Method-Selection Assistant

一旦你有了计算结果，Advanced 标签页顶部会显示一个 Method-Selection Assistant 面板。点击展开。该助手会查看你的项目状态 — 样本量、数据集中的分组候选、每个内生 LV 的 R^2 、HTMT 违规、信度阈值 — 并返回一个排好序的候选列表，列出最合理的下一步进阶方法：

- Bootstrap (High)，当你有计算结果但还没有置信区间时。
- MGA + MICOM (High)，当数据集有一个分类列，有 2–5 个层级且每组样本量足够时。
- IPMA (High)，当目标 LV 的 $R^2 \geq 0.33$ 时 — 有足够的实质性方差供前置变量排名产生可行的行动建议。
- PLSpredict (Suggested)，至少有一个内生 LV 的 $R^2 \geq 0.19$ 且 $n \geq 100$ 时。
- FIMIX-PLS (Suggested)， $n \geq 200$ 时作为异质性检查。
- HOC (Suggested)，当两个 LV 的 HTMT 接近 1 时 — 它们可能是同一二阶构念的一阶维度。
- PLSc (Optional)，当所有 LV 都是反映型时 — 一种共因子稳健性检查。
- Copula / CMB / Nomological / CVPAT (Optional)，作为标准的预防审稿人质疑的诊断。

每个建议都带有基于你实际数字的理由（例如 “Target LV Loyalty has $R^2 = 0.42$ — IPMA identi-

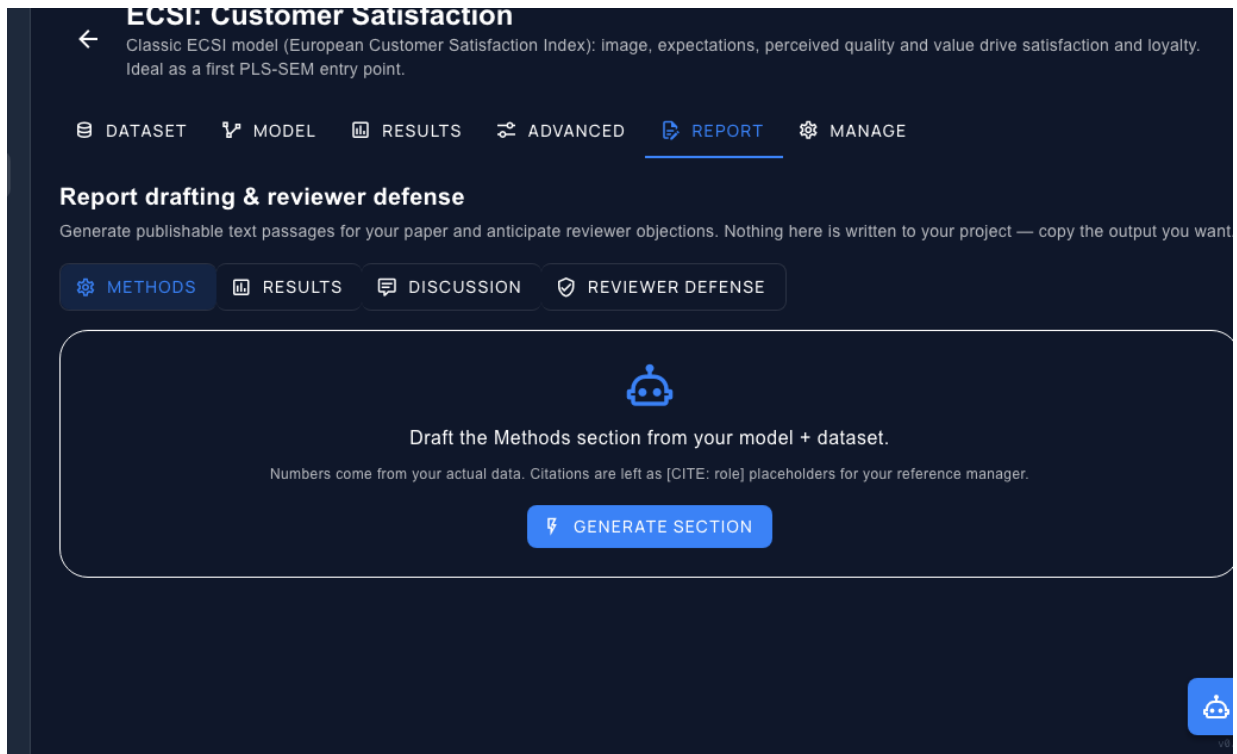


Figure 42 Report 标签页概览。四个子按钮 — Methods、Results、Discussion、Reviewer defense — 以及一个大卡片，用于打开相应的起草器。

fies which antecedents matter most for practical action”)。Open in ... 按钮（标注具体方法名称）会路由到 Advanced 子标签页，以便你配置和运行。

该面板是按需获取的；除非你打开它，否则什么都不会运行。点击 Regenerate suggestions 可重新读取你的项目数据。

11.7 报告起草

一个新的标签页 — Report — 位于 Advanced 和 Manage 之间。它的四个子按钮生成从你的实际数据中提炼出的报告草稿。

11.7.1 Methods / Results / Discussion

选择一个章节并点击 Generate section。你会得到：

- 左侧的 Markdown，可复制到草稿或博客文章。使用 Unicode 统计符号 (R^2 、 β 、 χ^2 、 f^2) — Markdown 输出中没有 LaTeX。
- 右侧的 LaTeX，可用于论文源码。使用 R^2 、 β 、 $\text{cite}\{role\}$ 等适合期刊模板的惯例。
- 最右侧的引用槽：本节引用的每个方括号角色一行，每一行链接到 DOI 或出版商页面。
- Notes：起草器认为人类作者应该知道的一到两句注意事项（例如 “your SRMR of 0.11 is above the 0.08 threshold - consider commenting on fit” ）。

起草器遵循 Hair/Ringle/Sarstedt 的惯例：

- Methods 报告样本、测量模型、估计选择、信度 + 效度标准、模型拟合指标，并预告进阶程序。

←

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

REGENERATE

Markdown

COPY

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 [Streukens+ 2016](#). Measurement model reliability and convergent validity were assessed via Cronbach's alpha, composite reliability (pc), 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

COPY

```
\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{streukens-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-htmt}. 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}.
```

Citation slots

Hair+ 2019

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

Streukens+ 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.

Figure 43 生成的 Methods 章节。左侧是 Markdown 面板，下方是 LaTeX 面板，右侧是引用栏。文本中的每一个数字都可追溯到项目。

- **Results** 报告测量模型质量 (α 、 ρ_C 、AVE)、区分效度 (HTMT, 以及任何违规都会点名)、带显著性的路径 LV 的 R^2 以及 Chin 1998 的口头描述符, 还有 SRMR。
- **Discussion** 实质性地解释显著路径, 处理非显著路径, 为具有最高 R^2 的目标 LV 阐明实际含义, 列出局限

用面板页头中的复制图标按钮复制任一面板。用章节页头中的 **Regenerate** 按钮重新生成一起草器会重新读取 Bootstrap 或模型调整后很有用。

11.7.2 Reviewer defense

第四个子按钮 — **Reviewer defense** — 是起草器的对抗模式。点击 **Generate reviewer defense**。

你会得到同行审稿人可能提出的四到七个预期反对意见, 每一个包含:

- **Concern**: 用你快照中的确切数字表达的一到两句话。
- **Reviewer persona**: 短的画像 (“CB-SEM traditionalist”、“Statistical purist”、“Senior methods editor”), 让你知道期望什么样的框架。
- **Severity**: Critical、Moderate 或 Nitpick。
- **Defense**: 实际的反驳论点, 引用你的数字并在其中援引论文库。
- **Evidence**: 支撑辩护的具体快照数值。
- **Mitigation**: 0–2 句关于在提交之前还应运行什么内容 (如有)。
- **Suggested citations**: 可点击的胶囊, 指向辩护者依赖的论文。

卡片顶部的短整体策略段落列出了固定你修改信的两三个行动 — 通常是 “lead with X, concede Y as a limitation, commit to Z in a follow-up”。

辩护者是一个模拟器, 不是先知。真正的审稿人会带来你无法预见的背景 — 与你特定文献相关的理论异议、你目标期刊的方法论特殊性、个人偏好。将输出用作演练, 而不是保证。

11.8 论文库

每一个产生引用的 AI 界面都从相同的策划库中提取 — 大约 100 篇开放获取的 PLS-SEM 论文, 涵盖方法学、批评 (R)。每个条目是标题、作者列表、年份、DOI 以及出版商页面上的摘要。没有全文 — 只有审稿人自己找到论文。论文库存储在与你项目相同的数据库中, 通过语义向量搜索查询: 当你问一个问题时, 助手用一个多语言嵌入 LLM。然后模型使用带方括号的角色名称进行引用 (例如 [henseler-2015-htmt]), 前端将其替换为胶囊标记。如果一个胶囊显示为灰色/扁平 (不可点击), 则表示助手引用了库中还没有的角色 — 要么是超出我们策划窗口。

11.9 配额与隐私

配额。Beta 期间每用户每月 100 轮。分别适用于解释性回合 (聊天、模型设计器、报告起草器) 和 hints (每月 100 次 hint 获取)。回合不结转。

隐私。Companion 读取项目结构和汇总 (行数、列统计信息、 R^2 、路径系数), 从不读取实际的调查行。所有

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.

Figure 44 一张 reviewer-defense 卡片。严重程度标记 (Critical / Moderate / Nit-pick)、审稿人角色、担忧、辩护、证据、缓解措施以及建议的引用胶囊。

11.10 关闭

Account → AI Companion → 切换关闭会在整个应用中禁用所有 AI 界面：抽屉消失，hints 条带折叠，设计器 + 报告标签页仍然存在，但其生成按钮变为不活动状态，并显示 “not enabled” 提示。

存储的对话会保留（什么都不会被删除）；稍后重新启用时，你的聊天历史依然停留在离开时的位置。

12 公开调查

OpenPLS 可以将你的问卷托管在公开 URL 上，收集匿名回复，并将它们直接拼接回你的项目作为数据集。对于 PLS-SEM 工作流程，你不需要第三方调查工具。

典型流程：

1. 构建 Model 标签页中的条目集（手动构建，或通过 AI 问卷设计器 — 见上一章）。
2. 发布问卷到公开 /q/{slug} URL。
3. 分享该 URL 给受访者。新回答几乎实时进入构建器。
4. 导入累积的回复，一键作为数据集。
5. 计算模型在你的新鲜数据集上。

本章介绍每一步。

12.1 构建问卷

打开项目的 Model 标签页，然后打开 Form 子视图。如果你的 LV 还没有指标，面板会以空状态开始；使用 AI 问卷设计器（见第 12 章）或手动添加条目。每个条目对应未来响应 CSV 的一列。

每条目字段：

- Wording: 受访者看到的句子。反向编码的条目带有 R 标记。
- Item ID: 自动生成（例如 TRUST1、SAT2），用作导出数据集中的列名。
- Reverse coded: 如果条目的极性反转（受访者需要心理否定），请切换此开关。反向编码在导入时自动应用。
- LV 绑定: 继承自父 LV 块；在块之间拖动条目以重新分配。

全局字段（右列）：

- Response scale: 5 点或 7 点 Likert。适用于所有 Likert 类型条目；开放式和人口学条目保持自己的类型。
- Audience language: 向受访者显示的语言（English、German、Spanish、French、Italian、Japanese）。
- Demographics: 从六个内建可选字段（年龄、性别、国家、教育、语言、职业）中选择零个或多个。每个条目最多只能有一个 demographic。
- Public intro: 受访者在第一个问题上方看到的段落。与你的内部备注分开（内部备注永远不会出现在公开版本中）。

12.2 发布

面板的 Publish 部分切换调查上线。点击 Publish: OpenPLS 生成一个随机 slug，将问卷的公开快照写入 responses 集合，并返回给你 URL：

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

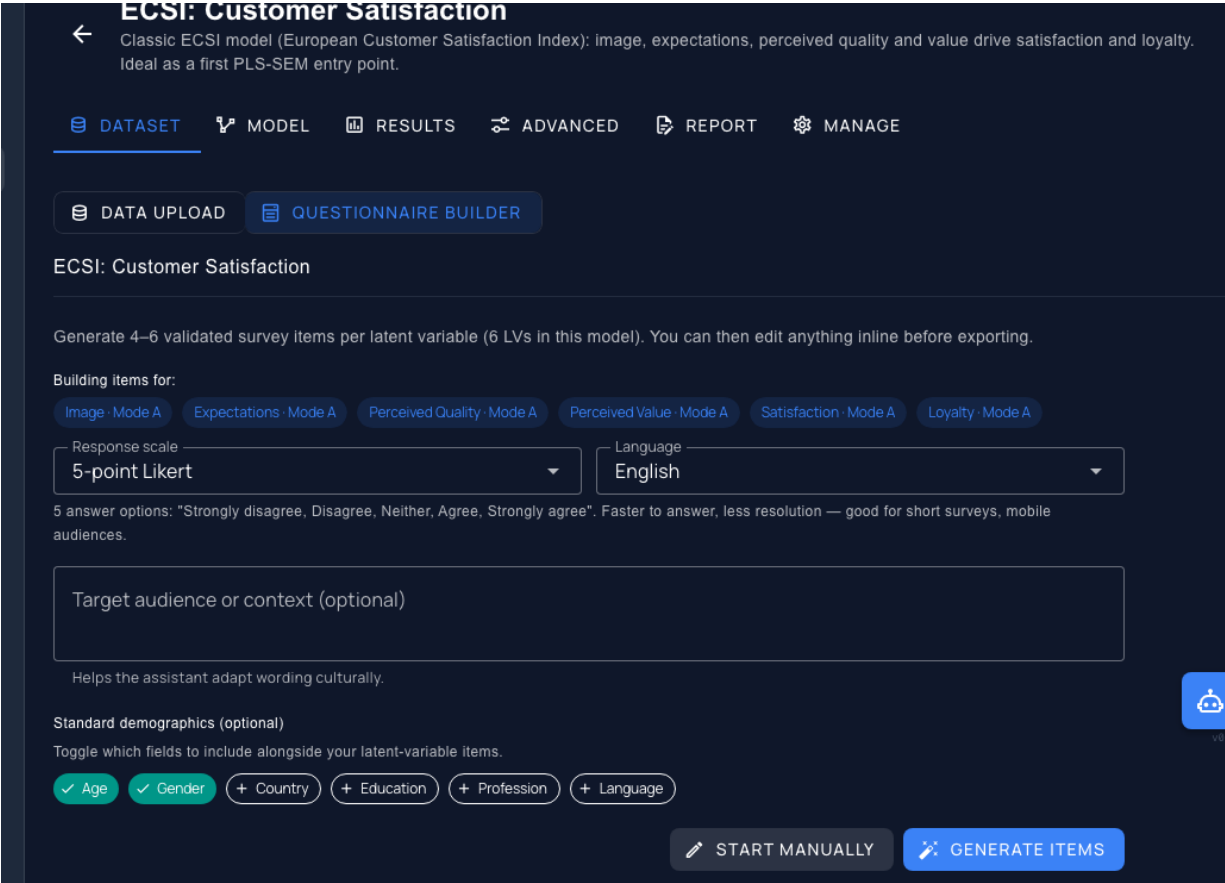


Figure 45 问卷构建器面板。左侧：LV 及其条目（拖动重新排序，点击移除，点击措辞进行编辑）。右侧：发布卡片，包含响应量表、语言、受众和人口学切换。

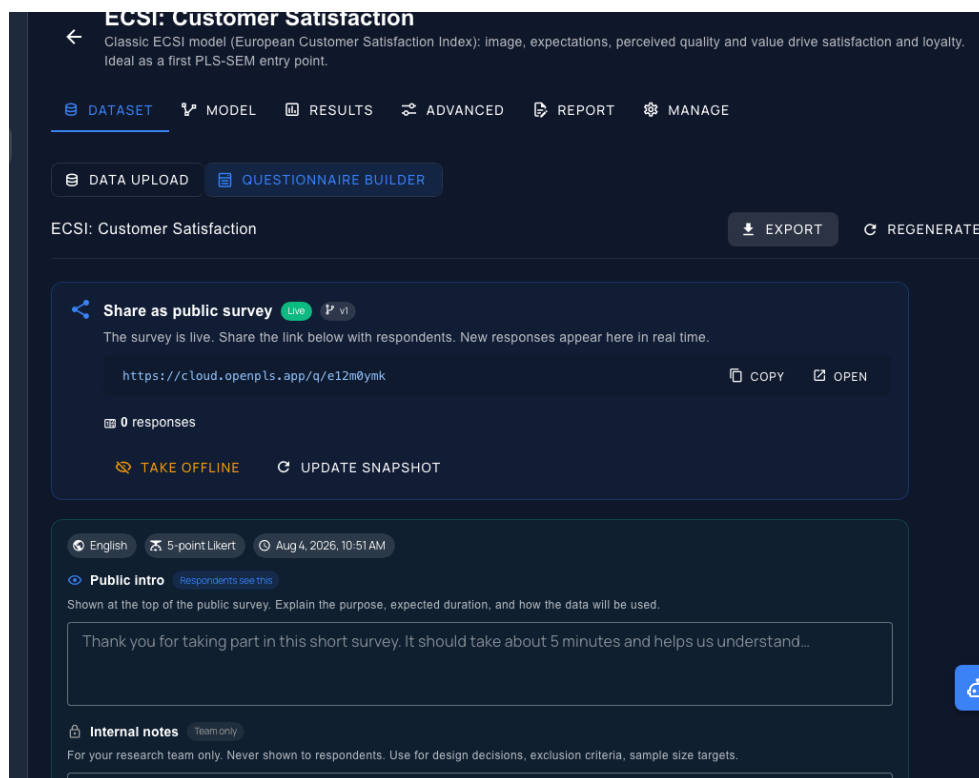


Figure 46 已发布的调查卡片。URL 显示带复制到剪贴板功能、实时响应计数器、三个操作（Preview / Unpublish / Publish as new version），以及版本标记。

拥有该 URL 的任何人都可以响应。无需账户—匿名是默认设置。公开页面只显示公开的介绍和条目本身，格式 Preview 在新标签页中打开调查，让你像受访者一样浏览一遍（预览模式下你的响应不计入）。

12.2.1 版本快照

一旦回复开始进入，更改问卷会造成数据完整性困境：如果你在中途重命名一个条目或更改其措辞，累积的回复会让你选择处理方式：

- Publish as new version（推荐）。停用当前 URL，生成新的 slug（/q/...），并将计数器移到零。旧回复以 v1 归档，仍可下载。
- Download responses, then save in place. 同一 URL，但你需要先下载更改前批次的 CSV 存档。
- Save in place. 同一 URL；新旧响应混合。仅在错别字修正时使用。

停用的 slug 仍然在线，但停止接受新响应—受访者会看到“不再接受响应”的消息。发布卡片列出所有先前版

12.3 公开受访者视图

受访者看到一个整洁的单列表单，带有进度条。条目按 LV 分组，以便认知负荷可控；LV 名称本身被隐藏（受访者看不到）。值得注意的设计决策：

- 语言切换器。受访者可以在你设置的受众语言之外，切换八种支持的语言—对国际部署很有用。

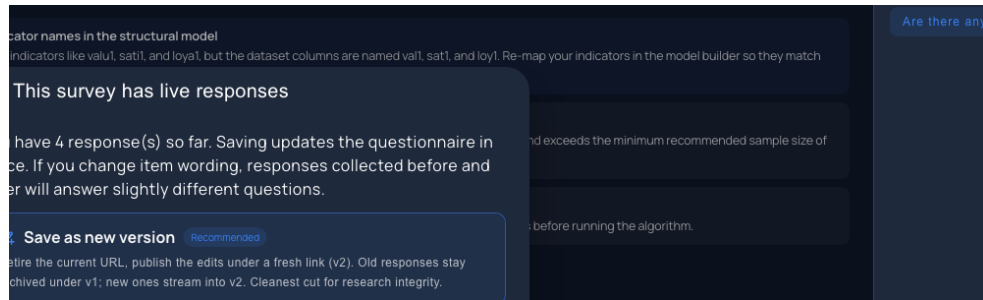


Figure 47 已发布调查有回复时的保存警告对话框。三条路径，每条都明确列出具体的后果。

- 同一浏览器一次响应的守卫。提交后，浏览器本地存储中的一个小标志会防止同一设备意外重新提交。并
- 感谢页面。提交后，受访者会看到一段简短的确认。没有账户创建提示，没有营销。
- 移动优先。该表单可在约 320 px 的视口下工作；条目卡片堆叠，语言切换器变为紧凑菜单。

12.4 观察响应到达

构建器中的发布卡片显示当前响应计数，几乎实时（每次提交后几秒内更新）。计数器下方的 Recent 小节列出最后五次提交的时间戳和语言。

12.5 将响应导入为数据集

当你有足够的响应来计算模型 — 对于典型的 PLS-SEM 研究，至少 100 个用于探索性工作；200+ 用于可发表级推断 — 点击发布卡片上的 Import as dataset。

OpenPLS 在你的项目中创建一个新的数据集。列布局：

- **item_id** 列：每个调查条目一列，使用你定义的 TRUST1、TRUST2、… 条目 ID。单元格值是整数形式的 Likert 响应（1..5 或 1..7）。反向编码的条目在导入时被反转，因此下游代码无需再考虑。
- **人口学**列：你启用的每个人口学字段一列（age、gender、country、…）。根据字段类型是字符串或整数。
- **submitted_at**：响应提交时的 ISO 8601 时间戳，对时间过滤或波次分析很有用。
- **language**：受访者使用的语言，对跨语言 MICOM/MGA 比较很有用。

打开 Data 标签页；新数据集已 Ready。像往常一样计算模型。如果你稍后发布新版本，再次导入会写入一个新数据集 — 旧的会保留，以便你可以在问卷更改前后比较估计。

12.6 下载原始 CSV

如果你更愿意将响应带入外部工具（R、Stata、SPSS、Python），点击发布卡片上的 Download responses (CSV)。CSV 的列与导入的数据集相同。每行是一个提交；空单元格表示受访者跳过了该可选字段。

下载只计算已完成的提交。部分响应（受访者在提交前关闭标签页）不会到达数据库。

ECSI: Customer Satisfaction

About you

Age *

e.g. 34

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 →

Figure 48 桌面上渲染的公开调查页面。右上角是语言切换器，进度条，带单选按钮的条目卡片，Previous / Next 导航，以及最终的 Submit 操作。

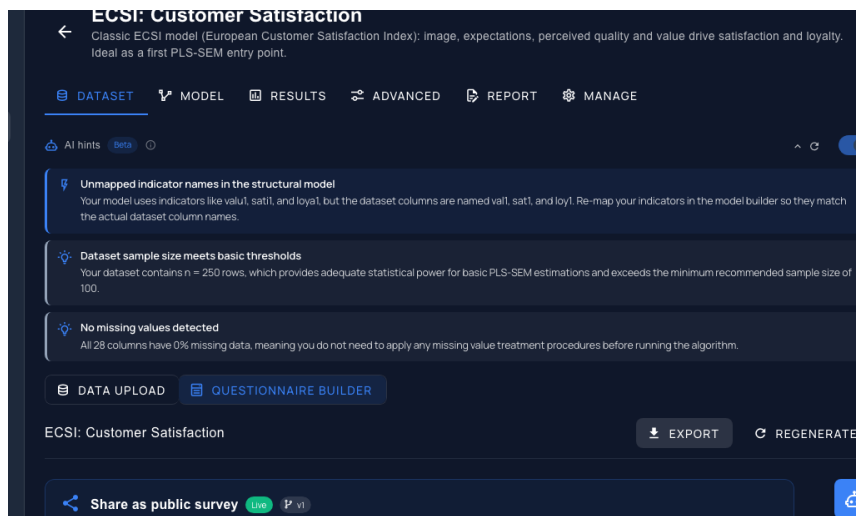


Figure 49 响应计数器和最近列表。计数器实时更新；最近列表显示 submitted_at 加语言标记。

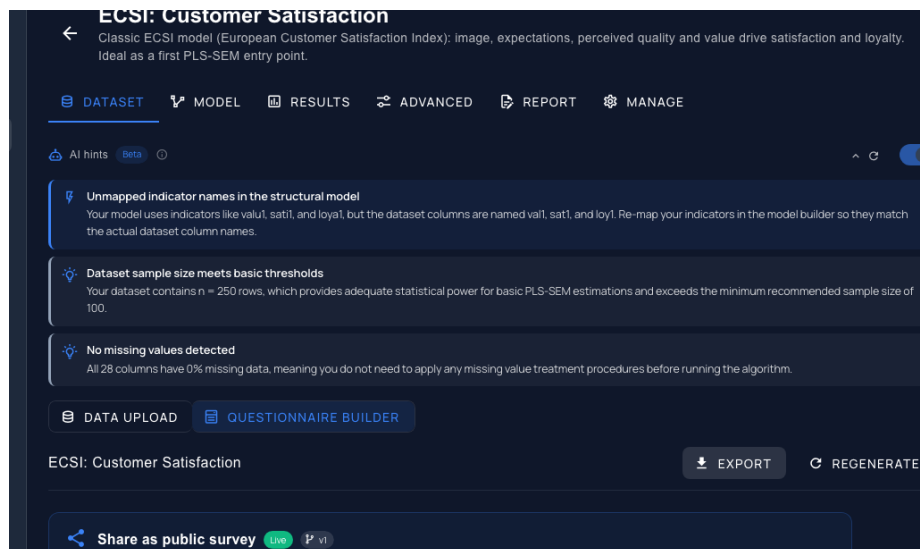


Figure 50 发布卡片上的 import-as-dataset 按钮。确认对话框显示：条目列、人口学列、样本量和版本。点击 Create dataset 写入。

12.7 取消发布

Unpublish 让 URL 下线。现有响应在构建器中仍然可访问，仍可下载或导入；新访问者会看到“不接受响应”的不会释放 — 在同一问卷上再次点击 Publish 时，会尽可能重用相同的 slug。

注意：如果你发布 → 收集 → 取消发布 → 更改问卷 → 重新发布 → 收集，你将得到绑定到同一 URL slug 的混合响应。上方的保存警告对话框在正常编辑流程中会防止这种情况，但取消发布后编辑再重新发布是 Publish as new version。

12.8 限制与惯例

- **每个调查的响应上限：**Beta 期间无硬性上限。现实的软性上限是每个项目每个调查几千个响应，超出后约 10 000+ 响应，请拆分为多个调查或使用外部面板提供商。
- **响应保留：**只要父项目存在就无限期保留。删除项目会级联删除响应。
- **匿名：**只存储受访者输入的内容。IP 地址、User-Agent 以及任何其他标识性头部都不会被持久化。如果有 ID 字段作为调查条目。

13 附录

13.1 术语表

AVE：平均方差抽取量 (Average Variance Extracted)。反映型 LV 指标的平方载荷均值。高于 0.5 = LV 对指标的解释多于误差。

BIC：贝叶斯信息准则。对额外参数施加惩罚的模型拟合指数；越低越好；仅在嵌套规范之间可比。

Bootstrap：有放回重抽样，用于为缺乏解析分布的统计量推导标准误、p 值和置信区间。

组合信度 (ρ_c 、 ρ_c)：PLS-SEM 中现代默认的信度指标。高于 0.7 可接受，高于 0.95 表明指标冗余。

CMB：共同方法偏差 (Common Method Bias)。由测量工具而非所测构念引起的系统性误差。通过 Lindell-Whitney 标记法或 Harman 单因子检验进行诊断。

CVPAT：交叉验证预测能力检验。将 PLS 的样本外损失与基准（指标平均或线性模型）进行比较。

内生性 (Endogeneity)：预测变量与误差项相关。违反 OLS 假设；在 PLS 中通过高斯 Copula 方法进行校正。

内生 LV (Endogenous LV)：至少有一个进入路径的构念。与外生对比。

外生 LV (Exogenous LV)：仅有出去路径的构念。其方差不由模型解释。

FIMIX-PLS：有限混合 PLS。当你怀疑存在未观测的异质性时发现潜类别。

形成型 (Mode B)：指标引起 LV 的测量模型（收入 + 教育 + 职业 → SES）。

Fornell-Larcker：经典的区分效度检验：对于每对 A、B， $\sqrt{AVE(A)} > correlation(A, B)$ 。

GoF：拟合优度 (Goodness-of-Fit)。 $\sqrt{(\text{mean}(\text{communality}) \times \text{mean}(R^2))}$ 。整体模型质量代理指标。

HTMT：异质一同质性比率。现代区分效度检验。高于 0.85（保守）或 0.90（宽松）表明构念过于相似。

HOC: 高阶构念 (Higher-Order Construct) 。由多个一阶 LV 组成的构念 (例如 Brand Experience = Sensory + Affective + Intellectual + Behavioral) 。

指标 / 显变量 (MV) : 测量潜变量的观测数据集列。

内部模型: LV 之间的结构路径。与外部模型对比。

IPMA: 重要性-绩效地图分析。识别目标 LV 的哪些前置变量结合了低重要性与低绩效 (最大的改进杠杆) 。

潜变量 (LV) : 通过指标间接测量的未观测构念。

载荷 (Loading) : 反映型测量模型中指标与其 LV 之间的相关性。

测量模型: LV 与其指标之间的箭头。反映型 (Mode A) 或形成型 (Mode B) 。

MGA: 多组分析。检验路径系数在组间是否有差异。

MICOM: 组合模型的测量不变性。跨组进行 MGA/调节效应分析的前提条件。

Mode A: 反映型测量 (LV → 指标) 。

Mode B: 形成型测量 (指标 → LV) 。

Moderation: X 对 Y 的效应取决于第三个变量 Z。通过交互项 $X \times Z$ 进行检验。

外部模型: LV 与指标之间的箭头。与内部模型对比。

路径系数 (β) : 内部模型箭头的标准化强度。像 OLS 回归中的 β 一样解释。

PLSc: 一致 PLS。当构念真正是共因子而非组合体时纠正衰减偏差。

PLSpredict: 通过 k 折交叉验证进行的样本外预测评估。

Q^2 : Stone-Geisser 交叉验证冗余。大于 0 = 预测相关性。

$Q^2_{predict}$: 使用 k 折 CV 的 Q^2 的 PLSpredict 变体。

反映型 (Mode A) : LV 引起其指标。改变构念, 所有指标一起变化。

R^2 : 由前置变量解释的内生 LV 方差。样本内。

SRMR: 标准化残差均方根。低于 0.08 拟合良好, 低于 0.10 可接受。

结构模型: 与内部模型相同。

t 统计量: 路径估计除以其 Bootstrap 标准误。高于约 1.96 在 $\alpha = 0.05$ 时显著。

VIF: 方差膨胀因子。低于 5 共线性可接受, 低于 3 更严格。既应用于形成型指标 (外部 VIF) , 也应用于每个内 LV (内部 VIF) 。

13.2 键盘快捷键

Editor

- Cmd+Z / Ctrl+Z: 撤销
- Cmd+Shift+Z / Ctrl+Y: 重做
- 选中 LV 后按 Delete 或 Backspace: 删除 LV 及其路径

- 双击一条边：删除该路径

Editor 拖动操作

- 从 LV 源手柄（右侧蓝点）拖到 LV 目标手柄（左侧绿点）：绘制新路径
- 从右侧侧边栏拖动指标标记到 LV 圆圈：将指标分配到该 LV

全局

- Cmd+K / Ctrl+K：未来的命令面板（路线图）

13.3 参考文献

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.