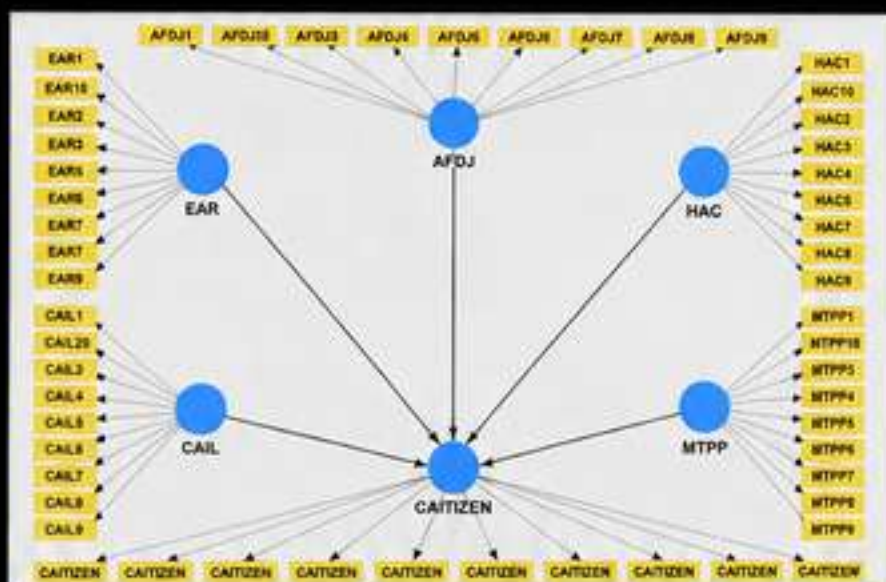


Análisis estructural aplicado:

Uso de SmartPLS en la investigación de las ciencias de la administración

JUAN MEJIA TREJO



INVESTIGACIÓN CUANTITATIVA
MODELOS DE ECUACIONES ESTRUCTURALES
MODELO CAITIZEN
METODOLOGÍA DE LA INVESTIGACIÓN

Modelo CAITIZEN

EAR
AFDJ
HAC
CAIL
MTPP
= CAITIZEN



PLS-SEM con SmartPLS

Evaluación del modelo de medición

- Confianza
- Validez convergente
- Validez discriminante

Evaluación del modelo estructural

- Colinealidad
- Significancia de caminos
- R^2 y Q^2
- Q^2 predictivo

Evaluación de la capacidad predictiva

- PLS predict
- Validación cruzada

Reporte académico

- Tablas
- Figuras
- Interpretación

AMIDI
Scientia et Praxis

**Análisis estructural
aplicado:
Uso de SmartPLS en la
investigación de las
ciencias de la
administración**



Análisis estructural aplicado: Uso de SmartPLS en la investigación de las ciencias de la administración

Juan Mejía Trejo



Este libro fue sometido a un proceso de dictamen por pares doble ciego de acuerdo con las normas establecidas por el Comité Editorial de la Academia Mexicana de investigación y Docencia en Innovación (AMIDI)

Esta obra se encuentra bajo la licencia Atribución-No Comercial-Sin Derivadas 4.0 (CC BY- NC-ND 4.0), de Creative Commons. Usted puede descargar esta obra y distribuir en cualquier medio o formato dando crédito a los autores, pero no se permite su uso comercial ni la generación de obras derivadas.



Primera edición, 2026

D.R. © Academia Mexicana de Investigación y Docencia en Innovación (AMIDI)
Av. Paseo de los Virreyes 920.
Col. Virreyes Residencial
C.P. 45110, Zapopan, Jalisco

Nota editorial.

La autoría, revisión final y responsabilidad intelectual de esta obra corresponden exclusivamente al autor.

ISBN: 978-970-96061-4-0

Hecho y editado en México
Made and edited in Mexico



AMIDI
Academia Mexicana
de Investigación y Docencia
en Innovación

Sobre el autor

Juan Mejía Trejo

Universidad de Guadalajara, CUCEA, Jalisco, México

ORCID: <https://orcid.org/0000-0003-0558-1943>

Correo: jmejia@cucea.udg.mx

Doctor en Ciencias Administrativas y profesor investigador del Centro Universitario de Ciencias Económico-Administrativas de la Universidad de Guadalajara. Es Investigador Nacional Nivel III del Sistema Nacional de Investigadoras e Investigadores (SNII, de México). Su trayectoria académica se centra en la administración de la innovación para el desarrollo sostenible y las aplicaciones de la inteligencia artificial en las ciencias de la administración.

Índice

Sobre el autor v	
INTRODUCCIÓN	1
CAPÍTULO 1. Bases del análisis estructural aplicado.....	5
PLS-SEM como ruta aplicada de modelación estructural.....	7
Modelo de medición y modelo estructural	11
Entorno básico de SmartPLS4	17
Controles básicos para construir el modelo	21
Controles básicos para estimar, revisar y reportar resultados.....	26
CAPÍTULO 2. Bases operativas del modelado PLS-SEM en SmartPLS4.....	39
Preparación del archivo para SmartPLS4	41
Revisión inicial del archivo de datos importado en SmartPLS4.....	42
Name.....	43
No.	43
Type.....	43
Missings	44
Mean	44
Median	44
Scale min	44
Scale max.....	45
Observed min.....	45
Observed max.....	45
Standard deviation	45
Excess kurtosis	46
Skewness	46
Cramér-von Mises p value	46
Decisiones metodológicas si algún criterio no se cumple	47
Diagnóstico de cumplimiento de la base CAITIZEN en SmartPLS4	47
Interpretación del diagnóstico.....	48
Decisión metodológica.....	49
Conclusión del diagnóstico	49
Creación del área de trabajo en SmartPLS	49
Procedimiento paso a paso	51
Pantalla posterior a la creación del proyecto e importación del CSV.....	55
Síntesis del procedimiento	56
CAPÍTULO 3. Construcción del modelo en SmartPLS4	57
Inicio de la construcción del modelo CAITIZEN.....	58
Creación del modelo PLS-SEM en SmartPLS4	59
Apertura del área gráfica de trabajo.....	60

Creación de constructos latentes.....	61
Asignación y alineación de indicadores.....	62
Organización del modelo de medición reflectivo	63
Revisión de propiedades de los constructos	64
Establecimiento de relaciones estructurales.....	65
Verificación final del modelo CAITIZEN	66
Cierre metodológico de la construcción del modelo	67
Síntesis del procedimiento	68
CAPÍTULO 4. PLS-SEM Algorithm y colinealidad	69
Ejecución inicial del algoritmo	71
Función del algoritmo PLS-SEM en el modelo CAITIZEN	71
Verificación previa del modelo construido	71
Selección del algoritmo PLS-SEM en SmartPLS4.....	72
Configuración del algoritmo PLS-SEM	72
Resultados iniciales de la estimación PLS-SEM.....	73
Revisión inicial del modelo de medición	74
Evaluación de colinealidad interna mediante VIF.....	74
Diferencia entre Outer model VIF e Inner model VIF	74
Revisión del Inner model VIF en CAITIZEN	75
Criterios de interpretación del VIF	75
Resultados VIF del modelo CAITIZEN	75
Decisión metodológica sobre colinealidad	76
Síntesis del procedimiento	76
CAPÍTULO 5. Evaluación del modelo de medición reflectivo	78
Naturaleza reflectiva del modelo CAITIZEN.....	80
Evaluación de cargas externas.....	81
Criterios para conservar o eliminar indicadores.....	83
Evaluación de confiabilidad interna.....	84
Evaluación de validez convergente mediante AVE	87
Evaluación de validez discriminante	88
Criterio Fornell-Larcker	88
Criterio Cross Loading	89
Criterio HTMT	90
Validez discriminante conjunta	92
Decisión metodológica sobre el modelo de medición	93
Copias del modelo original	95
Exportación de reportes a Excel.....	99
Síntesis del procedimiento	101
CAPÍTULO 6. Ajuste global del modelo saturado	103
Modelo saturado y función del ajuste global	105
Preparación de una copia del modelo para el análisis de fit.....	106
Apertura del modelo FIT y revisión previa de la estructura	108
Configuración operativa mediante “Invert measurement model”	109
Configuración como <i>composite</i> en Modo A	112

Verificación visual del modelo configurado en Modo A	115
Ejecución del algoritmo para obtener el ajuste global	118
Decisión metodológica sobre el ajuste del modelo CAITIZEN	119
Síntesis del procedimiento	119
CAPÍTULO 7. <i>Bootstrapping</i> e inferencia del modelo	124
Función del <i>Bootstrapping</i> en PLS-SEM	126
Diferencia entre algoritmo PLS-SEM y <i>Bootstrapping</i>	127
Configuración de Complete <i>Bootstrapping</i> en SmartPLS4	129
Ejecución del cálculo <i>Bootstrap</i> y control del proceso.....	130
Salida gráfica del <i>Bootstrapping</i>	132
Recolección de <i>t-values</i> y <i>p-values</i> para cargas externas.....	133
Interpretación de coeficientes de trayectoria	135
Índice de ajuste aproximado: SRMR	136
Ajuste exacto mediante d_{ULS}	137
Ajuste exacto mediante d_G	138
Interpretación de resultados parciales en los índices de ajuste global	140
Decisión metodológica sobre el ajuste global del modelo saturado	143
Criterio editorial para reportar el ajuste global en artículos científicos	145
Síntesis del procedimiento	146
CAPÍTULO 8. Evaluación del modelo estructural	149
Función de la evaluación del modelo estructural	151
Modelo reflectivo final depurado CAITIZEN_PLS_SIN CAIL2	153
Selección del algoritmo PLS-SEM en SmartPLS4	155
Configuración del PLS-SEM Algorithm.....	156
Resultados gráficos del PLS-SEM Algorithm	157
Resultado R-square del constructo CAITIZEN	159
Ejemplo didáctico de descomposición de la varianza explicada.....	160
Resultado f-square del modelo CAITIZEN.....	162
Síntesis del procedimiento	163
CAPÍTULO 9. De la explicación estructural a la predicción del modelo	166
Por qué explicar no equivale a predecir en PLS-SEM.....	168
Preparación del modelo final para PLSpredict	170
Acceso a PLSpredict/CVPAT en SmartPLS4.....	172
Configuración predictiva: folds, repeticiones y semilla fija.....	174
Ejecución del cálculo y generación del reporte	177
Criterio de evaluación predictiva	179
Interpretación de $Q^2_{predict}$, RMSE y MAE	182
Síntesis del procedimiento	184
CAPÍTULO 10. Integración del modelo PLS-SEM	186
Función del reporte académico en PLS-SEM	188
Modelo de medición CAITIZEN	189
Validez discriminante mediante Fornell–Larcker y HTMT del modelo CAITIZEN	195
Evaluación de colinealidad del modelo estructural mediante valores VIF internos...	197
Evaluación del modelo estructural	199

Evaluación predictiva del modelo CAITIZEN mediante PLSpredict	201
Interpretación de resultados para discusión	206
Síntesis de procedimiento	209
<u>REFERENCIAS</u>	212

INTRODUCCIÓN

El análisis estructural aplicado mediante **PLS-SEM** se ha consolidado como una ruta metodológica especialmente pertinente para investigaciones orientadas a explicar fenómenos complejos, estimar relaciones entre constructos latentes y evaluar modelos teórico-empíricos con datos observables. Sin embargo, en la práctica académica persiste un vacío relevante: muchos investigadores utilizan **SmartPLS4** como una herramienta operativa, pero no siempre articulan con claridad la relación entre **teoría, construcción del modelo, depuración de indicadores, evaluación de medición, análisis estructural, capacidad predictiva y reporte académico**. Este vacío genera reportes fragmentados, decisiones metodológicas poco trazables y conclusiones que, aunque pueden apoyarse en salidas estadísticas, no siempre muestran una defensa conceptual y procedimental suficiente.

El propósito de este libro es atender dicho vacío mediante una ruta aplicada, sistemática y didáctica para el análisis **PLS-SEM** con **SmartPLS4**, tomando como eje de trabajo el modelo **CAITIZEN**. La obra no se limita a mostrar comandos del software, sino que busca explicar **por qué se ejecuta cada procedimiento, qué decisión metodológica implica y cómo debe interpretarse académicamente cada resultado**. En este sentido, el libro aporta una guía estructurada para transitar desde la preparación inicial del archivo de datos hasta la construcción del reporte final de resultados. Su contribución principal consiste en integrar el uso técnico de **SmartPLS4** con una lógica de investigación orientada a la trazabilidad, la coherencia teórica, la validez de medición, la explicación estructural y la predicción del modelo.

La metodología de investigación que sustenta la obra corresponde a un enfoque **cuantitativo, explicativo-predictivo y aplicado**, apoyado en la modelación de ecuaciones estructurales por mínimos cuadrados parciales. El procedimiento se organiza a partir de una secuencia progresiva: primero se prepara y revisa la base de datos; después se construye el modelo en **SmartPLS4**; posteriormente se evalúa el modelo de medición reflectivo; enseguida se revisan criterios de ajuste, inferencia, colinealidad, varianza explicada, tamaños de efecto y capacidad predictiva; finalmente se estructura el reporte académico. Esta lógica permite que el investigador no interprete resultados de manera aislada, sino como parte de una ruta metodológica integrada.

El **Capítulo 1, Bases del análisis estructural aplicado**, introduce los fundamentos necesarios para comprender la lógica general del **PLS-SEM**. En este capítulo se explica la diferencia entre modelo de medición y modelo estructural, así como la función de **SmartPLS4** como entorno de trabajo para construir, estimar, revisar y reportar modelos. Su importancia radica en establecer el lenguaje metodológico básico que permitirá comprender los capítulos posteriores. Aquí se aclara que **PLS-SEM** no es solo una técnica estadística, sino una estrategia de análisis que exige correspondencia entre teoría, indicadores, constructos y relaciones estructurales.

Juan Mejía Tréjo

El **Capítulo 2, Preparación del archivo para SmartPLS**, aborda una etapa frecuentemente subestimada: la revisión inicial del archivo de datos. Este capítulo muestra cómo verificar nombres de variables, número de casos, tipo de dato, valores perdidos, medias, medianas, mínimos, máximos, desviación estándar, asimetría, curtosis y pruebas de distribución. Su aporte es metodológico porque permite anticipar problemas antes de construir el modelo. Además, incorpora decisiones sobre qué hacer si algún criterio no se cumple, lo cual fortalece la trazabilidad del análisis.

El **Capítulo 3, Construcción del modelo CAITIZEN**, desarrolla el proceso de creación del modelo **PLS-SEM** en **SmartPLS4**. Se explica la apertura del área gráfica, la creación de constructos latentes, la asignación de indicadores, la organización del modelo reflectivo y el establecimiento de relaciones estructurales. Este capítulo es central porque traduce el modelo teórico en una representación operativa dentro del software. Su función es asegurar que la construcción visual del modelo corresponda con la lógica conceptual propuesta por la investigación.

El **Capítulo 4, Ejecución inicial del algoritmo PLS-SEM**, presenta la primera estimación del modelo. Se revisa la selección del algoritmo, su configuración, los resultados iniciales y la evaluación preliminar del modelo de medición. También se incorpora la revisión de colinealidad interna mediante valores VIF. Este capítulo permite identificar si el modelo puede continuar hacia evaluaciones más profundas o si requiere ajustes preliminares. Su valor está en mostrar que la estimación inicial no debe interpretarse como resultado final, sino como diagnóstico técnico-metodológico.

El **Capítulo 5, Evaluación del modelo de medición reflectivo**, se concentra en la calidad de los indicadores y constructos. Se revisan cargas externas, criterios para conservar o eliminar indicadores, confiabilidad interna, validez convergente mediante AVE y validez discriminante. También se abordan decisiones metodológicas sobre el modelo de medición, copias del modelo original y exportación de reportes a Excel. Este capítulo es fundamental porque sin un modelo de medición válido no puede sostenerse con rigor la interpretación del modelo estructural.

El **Capítulo 6, Ajuste global y modelo saturado**, introduce la preparación de una copia específica para el análisis de fit. Se explican la función del modelo saturado, la configuración mediante “Invert measurement model”, el uso de composite en Modo A y la ejecución del algoritmo para obtener el ajuste global. Este capítulo permite comprender que el ajuste del modelo en **PLS-SEM** requiere decisiones operativas específicas y no debe confundirse con la evaluación ordinaria del modelo estructural.

El **Capítulo 7, Bootstrapping e inferencia estadística**, aborda la diferencia entre el algoritmo **PLS-SEM** y el Bootstrapping. Se explica cómo configurar Complete Bootstrapping, ejecutar el cálculo, revisar salidas gráficas, recolectar t-values y p-values, interpretar coeficientes de trayectoria y revisar índices como SRMR, d_ULS y d_G. Este capítulo permite pasar de la estimación del modelo a la evaluación de estabilidad inferencial, lo cual es indispensable para reportar resultados con respaldo estadístico.

Juan Mejía Tréjo

El **Capítulo 8, Evaluación del modelo estructural**, analiza la función explicativa del modelo final depurado. Se revisan los resultados gráficos del algoritmo, el R-square del constructo **CAITIZEN** y los valores f-square del modelo. Este capítulo permite interpretar la magnitud de la varianza explicada y la relevancia de los predictores dentro de la estructura propuesta. Su aporte consiste en conectar los resultados estadísticos con la explicación teórica del fenómeno.

El **Capítulo 9, Evaluación predictiva mediante PLSpredict/CVPAT**, introduce una distinción clave: explicar no equivale necesariamente a predecir. Este capítulo muestra cómo configurar folds, repeticiones y semilla fija, ejecutar el cálculo predictivo, interpretar Q^2_{predict} , RMSE y MAE, y comparar el modelo **PLS-SEM** con el modelo lineal de referencia. Su valor está en ampliar la lectura del modelo **CAITIZEN** desde la explicación estructural hacia la utilidad anticipatoria.

Finalmente, el **Capítulo 10, Reporte académico de resultados PLS-SEM**, integra los resultados en una narrativa científica. Se abordan el modelo de medición, la validez discriminante mediante Fornell–Larcker y HTMT, la colinealidad estructural, la evaluación estructural, la evaluación predictiva y la interpretación de resultados para discusión. Este capítulo cierra la obra al mostrar cómo convertir las salidas de **SmartPLS4** en un reporte académico coherente, defendible y publicable. Así, el libro ofrece una ruta completa para que el investigador no solo ejecute procedimientos, sino que construya conocimiento metodológicamente sólido a partir del modelo **CAITIZEN**.

La obra se organiza en una **ruta metodológica cuantitativa, explicativa-predictiva y aplicada**, estructurada en **10 capítulos** que conducen desde la comprensión inicial del **PLS-SEM** hasta el reporte académico final del modelo **CAITIZEN**, como se observa en la **Tabla I**.

Tabla I. Ruta Metodológica del libro

Capítulo	Etapla metodológica	Propósito
Capítulo 1. Bases del análisis estructural aplicado	Fundamentación metodológica	Comprender la lógica del PLS-SEM , el modelo de medición, el modelo estructural y el entorno básico de SmartPLS4 .
Capítulo 2. Preparación del archivo para SmartPLS	Preparación de datos	Revisar la base importada, sus variables, valores perdidos, estadísticos descriptivos y condiciones mínimas para iniciar el análisis.
Capítulo 3. Construcción del modelo CAITIZEN	Diseño operativo del modelo	Crear el modelo en SmartPLS4 , asignar constructos, indicadores y relaciones estructurales.
Capítulo 4. Ejecución inicial del algoritmo PLS-SEM	Estimación preliminar	Ejecutar el algoritmo PLS-SEM , revisar resultados iniciales, cargas y colinealidad interna mediante VIF.
Capítulo 5. Evaluación del modelo de medición reflectivo	Validación de medición	Evaluar cargas externas, confiabilidad interna, AVE, validez discriminante y decisiones sobre indicadores.
Capítulo 6. Ajuste global del modelo saturado	Evaluación de ajuste	Preparar una copia del modelo, configurar el modelo FIT y revisar criterios de ajuste global como SRMR, d _{ULS} y d _G .

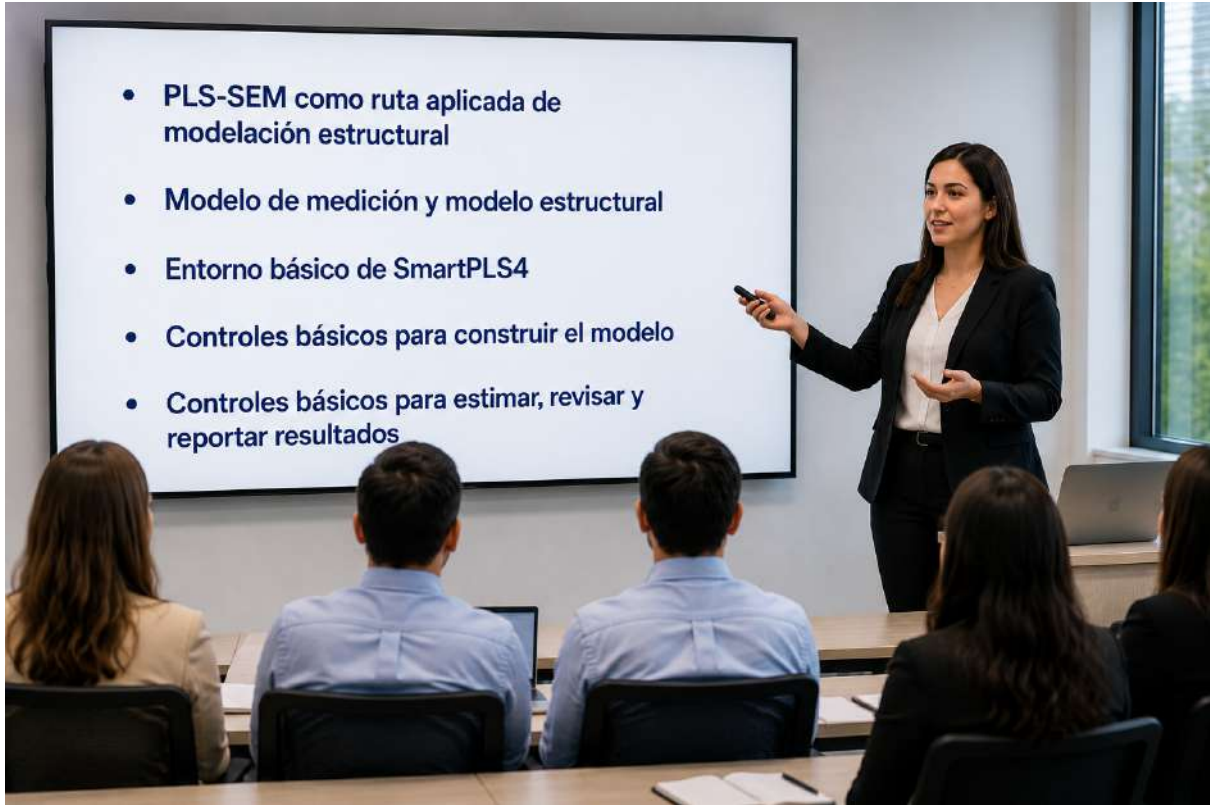
Juan Mejía Tréjo

Capítulo	Etapa metodológica	Propósito
Capítulo 7. Bootstrapping e inferencia estadística	Validación inferencial	Ejecutar Bootstrapping, obtener t-values, p-values, cargas externas, coeficientes de trayectoria e índices de ajuste.
Capítulo 8. Evaluación del modelo estructural	Explicación estructural	Analizar R-square, f-square, colinealidad estructural y capacidad explicativa del modelo CAITIZEN .
Capítulo 9. Evaluación predictiva mediante PLSpredict/CVPAT	Predicción del modelo	Evaluar Q^2_{predict} , RMSE, MAE y comparar el desempeño PLS-SEM frente al modelo lineal de referencia.
Capítulo 10. Reporte académico de resultados PLS-SEM	Integración y comunicación científica	Organizar los resultados del modelo de medición, estructural, predictivo y su interpretación para discusión académica.

Fuente: Elaboración propia

Juan Mejía Tréjo

CAPÍTULO 1. Bases del análisis estructural aplicado



El capítulo 1. Bases del análisis estructural aplicado, constituye una ruta metodológica para transformar fenómenos complejos en modelos empíricamente evaluables. En las ciencias de la administración, muchas variables relevantes no se observan de forma directa, como la innovación, la sostenibilidad, la confianza, la preparación tecnológica o la ciudadanía asistida por inteligencia artificial. Por ello, el investigador necesita convertir conceptos abstractos en constructos medibles, relacionarlos mediante hipótesis y evaluarlos con criterios estadísticos consistentes. Este capítulo introduce esa lógica a partir del modelo **CAITIZEN** (Mejía-Trejo, 2025) a manera de ejemplo, entendido como una arquitectura aplicada para analizar la ciudadanía sostenible asistida por inteligencia artificial desde una perspectiva crítica, ética, colaborativa y metacognitiva. Su importancia radica en ofrecer al lector la base conceptual necesaria para no reducir **SmartPLS4** a una operación mecánica de software, sino comprenderlo como apoyo a una estrategia metodológica de investigación (Mejía-Trejo, 2026).

El capítulo abona, en primer lugar, a la comprensión de **PLS-SEM** como una ruta de modelación estructural orientada a la explicación y la predicción. Su valor no consiste únicamente en estimar coeficientes, sino en articular problema de investigación,

Juan Mejía Trejo

constructos, indicadores, modelo de medición, modelo estructural e interpretación de resultados. Esta lógica resulta especialmente útil cuando se trabaja con modelos aplicados, fenómenos emergentes o estructuras conceptuales en desarrollo. En el caso de **CAITIZEN**, el interés no se limita a saber si existen relaciones estadísticas entre variables, sino a identificar cómo determinadas capacidades vinculadas con IA contribuyen a explicar una ciudadanía sostenible, ética y en red. Por tanto, el capítulo aporta una entrada ordenada para que el lector entienda qué se estima, por qué se estima y cómo debe interpretarse (Hair et al., 2019).

En segundo lugar, el capítulo permite distinguir entre **modelo de medición** y **modelo estructural**, diferencia indispensable para realizar un análisis **PLS-SEM** riguroso. El modelo de medición explica cómo los indicadores observables representan a los constructos latentes; el modelo estructural explica cómo esos constructos se relacionan entre sí. Esta distinción evita errores frecuentes, como interpretar hipótesis antes de verificar la calidad de los indicadores o asumir que una flecha en **SmartPLS4** equivale automáticamente a una relación válida. En el modelo **CAITIZEN**, esta separación resulta central porque constructos como **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** deben medirse adecuadamente antes de analizar sus relaciones estructurales. Así, el capítulo abona a una lectura metodológica progresiva: primero medir bien, después interpretar relaciones (Hair et al., 2021).

En tercer lugar, el capítulo introduce la lógica de las relaciones estructurales como hipótesis que deben estar teóricamente justificadas. Una ruta como **CAIL** → **CAITIZEN** o **HAIC** → **CAITIZEN** no debe entenderse como una conexión gráfica arbitraria, sino como una afirmación conceptual que requiere sustento. En este sentido, el capítulo enseña que el análisis estructural no comienza cuando se ejecuta el algoritmo, sino cuando el investigador define qué constructos intervienen, qué indicadores los representan y qué relaciones tienen sentido dentro del fenómeno estudiado. Esta aportación es relevante porque fortalece la responsabilidad teórica del investigador frente al uso del software. **SmartPLS4** facilita la construcción visual del modelo, pero la pertinencia de las rutas depende de la coherencia conceptual que las sostiene (Chin, 1998).

En cuarto lugar, el capítulo abona a la comprensión de los constructos reflectivos y su relación con los indicadores. En un modelo reflectivo, los indicadores son manifestaciones observables de una variable latente; por ello, deben guardar coherencia conceptual y estadística con el constructo al que pertenecen. Esta precisión es importante para el modelo **CAITIZEN**, donde cada dimensión se expresa mediante ítems asociados a alfabetización crítica en IA, responsabilidad ética, justicia de datos, colaboración humano-IA, transparencia metacognitiva y ciudadanía asistida por IA. El capítulo prepara al lector para comprender por qué, en capítulos posteriores, será necesario revisar cargas externas, confiabilidad interna, validez convergente y validez discriminante. Con ello, se evita tratar los indicadores como simples columnas de una base de datos y se les reconoce como evidencia empírica de una definición conceptual (Gefen et al., 2000).

Juan Mejía Tréjo

Finalmente, el capítulo introduce el entorno básico de **SmartPLS4** como espacio de trabajo donde se organiza, construye, estima, revisa y reporta el modelo. Su contribución no es solo operativa, sino pedagógica: familiariza al lector con elementos como **Workspace**, proyecto, archivo de datos, panel de indicadores, área gráfica, controles de construcción, controles de cálculo y reporte completo. Esta base inicial permite que los capítulos posteriores avancen con mayor claridad hacia la preparación del archivo, la construcción del modelo, la ejecución del algoritmo, la evaluación de medición, el Bootstrapping, la predicción y el reporte académico. En consecuencia, el capítulo abona una comprensión fundamental: cada acción en **SmartPLS4** debe responder a una decisión metodológica previa, y cada salida del software debe convertirse en interpretación académica defendible (SmartPLS GmbH, s. f.-e).

PLS-SEM como ruta aplicada de modelación estructural

El **modelado de ecuaciones estructurales por mínimos cuadrados parciales**, conocido como **PLS-SEM**, constituye una ruta metodológica aplicada para analizar modelos integrados por **constructos, indicadores y relaciones estructurales**. En ciencias de la administración, muchos fenómenos relevantes no pueden observarse directamente, como la confianza, la innovación, la sostenibilidad, la intención de uso, la preparación tecnológica o la ciudadanía asistida por inteligencia artificial. Por ello, el investigador requiere transformar esos fenómenos en constructos medibles y en relaciones empíricamente evaluables. En este libro, **PLS-SEM** no se presenta únicamente como una técnica estadística, sino como una lógica progresiva de investigación que conecta el **problema**, la **pregunta de investigación**, los **constructos**, los **indicadores**, el **modelo de medición** y el **modelo estructural**. Esta perspectiva permite evitar que el uso de **SmartPLS4** sea entendido como una operación mecánica de software. Antes de ejecutar un algoritmo, el investigador debe comprender qué desea explicar o predecir, qué constructos intervienen, cómo se miden y qué relaciones teóricas justifican el modelo. Por tanto, **PLS-SEM** debe asumirse como una ruta de **construcción, estimación, evaluación e interpretación** de modelos aplicados. Esta orientación resulta especialmente pertinente para el modelo **CAITIZEN**, porque permite vincular teoría, medición y análisis estructural dentro de una secuencia metodológica ordenada (Mejía-Trejo, 2026).

La pertinencia de **PLS-SEM** se relaciona con su orientación **explicativa y predictiva**. A diferencia de enfoques centrados principalmente en el ajuste global del modelo, **PLS-SEM** busca estimar relaciones entre constructos y maximizar la varianza explicada de las variables dependientes. Esta característica resulta útil en investigaciones aplicadas donde se analizan fenómenos complejos, modelos emergentes o estructuras con varios constructos latentes. Sin embargo, su orientación predictiva no elimina la necesidad de sustento teórico. Un modelo puede mostrar coeficientes relevantes y, aun así, ser débil si las relaciones no están justificadas conceptualmente. Por eso, el investigador debe equilibrar dos exigencias: la **capacidad del modelo para explicar o predecir** el fenómeno y la **coherencia teórica**

Juan Mejía Trejo

de las hipótesis. En el caso del modelo **CAITIZEN**, el interés no se limita a estimar coeficientes entre variables, sino a analizar cómo determinadas capacidades críticas, éticas, colaborativas y metacognitivas se relacionan con la **ciudadanía sostenible asistida por inteligencia artificial**. Esta lógica hace que **PLS-SEM** sea adecuado para un libro aplicado, siempre que el análisis estadístico se mantenga conectado con la teoría y la interpretación sustantiva. El software puede facilitar el cálculo, pero no reemplaza la responsabilidad metodológica del investigador ni la necesidad de explicar por qué cada relación forma parte del modelo (Hair et al., 2019).

PLS-SEM permite evaluar simultáneamente dos niveles del modelo: el **modelo de medición** y el **modelo estructural**. El modelo de medición analiza cómo los indicadores observables representan a los constructos latentes. Por ejemplo, los indicadores **CAIL1 a CAIL10** deben representar adecuadamente el constructo **CAIL**, relacionado con la alfabetización crítica en inteligencia artificial. El modelo estructural, en cambio, analiza las relaciones entre constructos. Por ejemplo, la ruta **CAIL → CAITIZEN** representa una hipótesis sobre la relación esperada entre alfabetización crítica en **IA** y ciudadanía sostenible asistida por **IA**. Esta distinción es fundamental porque no puede interpretarse con rigor una relación estructural si antes no se ha evaluado la calidad de la medición. Un coeficiente significativo entre dos constructos pierde valor si los indicadores de esos constructos son débiles, ambiguos o inconsistentes. Por ello, el análisis **PLS-SEM** debe avanzar por etapas: preparar los datos, construir el modelo, ejecutar el algoritmo, revisar el modelo de medición, evaluar la validez discriminante, analizar el modelo estructural y contrastar hipótesis mediante **Bootstrapping**. Esta secuencia evita conclusiones prematuras y fortalece la interpretación metodológica del modelo. En este sentido, el lector debe comprender desde el inicio que **PLS-SEM** no consiste solo en obtener resultados, sino en validar progresivamente la calidad de una estructura conceptual y empírica (Hair et al., 2021).

En una investigación aplicada, **PLS-SEM** no debe confundirse con una simple técnica para obtener resultados numéricos. **SmartPLS4** puede generar cargas externas, coeficientes de trayectoria, valores **R^2 , f^2 , VIF, AVE**, confiabilidad compuesta, **t-values, p-values** e intervalos de confianza. Sin embargo, esos resultados requieren interpretación metodológica y sustantiva. Una carga externa alta no significa necesariamente que el constructo esté bien conceptualizado; del mismo modo, un coeficiente significativo no implica automáticamente una contribución teórica fuerte. La utilidad de **PLS-SEM** depende de que el investigador sepa conectar resultados con decisiones. Por ejemplo, si un indicador presenta una carga baja, no debe eliminarse de manera automática; antes debe revisarse su pertinencia conceptual, su efecto sobre la confiabilidad y su contribución al contenido del constructo. En el modelo **CAITIZEN**, esta lógica es importante porque los indicadores no son solo datos importados a **SmartPLS4**, sino expresiones observables de constructos vinculados con **alfabetización crítica, ética, justicia de datos, colaboración humano-IA y prácticas metacognitivas**. Por tanto, **PLS-SEM** exige criterio investigador, no solo dominio operativo del software. El investigador debe saber interpretar umbrales, justificar decisiones y reconocer que cada resultado estadístico tiene sentido únicamente dentro del marco conceptual que lo originó (Hair et al., 2019).

Juan Mejía Tréjo

Una característica central de **PLS-SEM** es la necesidad de especificar adecuadamente la naturaleza de los constructos (Diamantopoulos & Winklhofer (2001). Los constructos pueden ser **reflectivos**, **formativos**, compuestos o jerárquicos, y cada tipo exige criterios de evaluación distintos. En un constructo reflectivo, los indicadores son manifestaciones del constructo; por tanto, se espera que estén correlacionados y que reflejen una misma variable latente. En un constructo formativo, los indicadores contribuyen conjuntamente a constituir el significado de la variable latente, por lo que no son necesariamente intercambiables ni se espera que presenten correlaciones elevadas entre sí; en consecuencia, su evaluación debe centrarse en la relevancia de los pesos, la colinealidad y la validez de contenido (Diamantopoulos & Winklhofer, 2001; Jarvis et al., 2003). Esta distinción no debe decidirse después de ver los resultados, sino antes de construir el modelo. En el caso **CAITIZEN**, los constructos se trabajan como **reflectivos** porque los indicadores se entienden como manifestaciones de capacidades latentes. Por ejemplo, si aumenta la responsabilidad ética ante la **IA**, se esperaría que aumenten también las respuestas asociadas a los ítems de **EAR**. Esta decisión afecta la lectura de cargas externas, confiabilidad interna, validez convergente y validez discriminante. Se introduce esta diferencia de forma breve, para que el lector comprenda que construir un modelo en **SmartPLS4** implica decisiones conceptuales previas a la estimación. En consecuencia, la especificación del modelo no es un detalle técnico menor, sino una condición metodológica que determina cómo se evaluarán los resultados posteriores (Hair et al., 2021).

La ruta aplicada de **PLS-SEM** inicia antes de abrir **SmartPLS4**. Primero se requiere una arquitectura conceptual mínima: identificar el fenómeno, formular el problema, precisar la pregunta estructural, definir constructos, establecer indicadores e identificar relaciones hipotéticas. Solo después tiene sentido preparar la base de datos, importarla al software y construir el modelo gráfico. Esta secuencia es relevante porque muchos errores en **PLS-SEM** no se originan en el algoritmo, sino en decisiones previas mal formuladas. Si dos constructos son conceptualmente indistinguibles, es probable que después aparezcan problemas de validez discriminante. Si los indicadores no representan adecuadamente el constructo, pueden presentarse cargas bajas o confiabilidad débil. Si se incorporan variables sin función teórica clara, el modelo puede volverse innecesariamente complejo. Por ello, el análisis estructural debe verse como una continuidad entre razonamiento conceptual y evaluación empírica. En este libro, **SmartPLS4** será usado como herramienta para construir, estimar y revisar el modelo **CAITIZEN**, pero la validez del análisis dependerá de la coherencia acumulada entre **teoría, medición, modelo y resultados**. Esta lógica permite que el software funcione como apoyo metodológico y no como sustituto del pensamiento investigador (Mejía-Trejo, 2026).

El análisis **PLS-SEM** también exige reconocer que los criterios de evaluación se aplican de manera ordenada. En primer lugar, se revisa que la base de datos esté limpia, codificada e importada correctamente. En segundo lugar, se construye el modelo con sus constructos, indicadores y relaciones. En tercer lugar, se ejecuta el algoritmo **PLS-SEM** para obtener una estimación inicial. Después se evalúa el modelo

Juan Mejía Trejo

de medición mediante **cargas externas, confiabilidad interna y validez convergente**. Más adelante se revisa la **validez discriminante** para comprobar que los constructos sean diferentes entre sí. Posteriormente se evalúa el modelo estructural mediante **colinealidad, coeficientes de trayectoria, R^2 y f^2** . Finalmente, se aplica **Bootstrapping** para analizar significancia estadística, **t-values, p-values** e intervalos de confianza. Esta ruta evita saltar directamente a la aceptación o rechazo de hipótesis sin haber revisado primero la calidad de la medición. En el modelo **CAITIZEN**, seguir esta secuencia permitirá interpretar con mayor rigor la influencia de **CAIL, EAR, AFDJ, HAIC y MTPP** sobre **CAITIZEN**. La fuerza de una hipótesis no se sostiene únicamente en su valor p, sino en la calidad integral del proceso metodológico que la precede (Hair et al., 2019).

En un libro aplicado, es conveniente que el lector comprenda desde el inicio que **PLS-SEM** combina razonamiento metodológico y operación técnica. La operación técnica se realizará en **SmartPLS4** mediante controles como **Workspace**, proyecto, importación de datos, creación del modelo, asignación de indicadores, conexión de constructos, ejecución del algoritmo, **Bootstrapping** y apertura del reporte completo. Sin embargo, cada acción del software debe corresponder a una decisión metodológica. Crear un constructo en el lienzo gráfico equivale a representar una variable latente definida conceptualmente. Asignar indicadores significa vincular ítems observables con un constructo. Dibujar una flecha implica formular una relación estructural hipotética. Ejecutar el algoritmo supone estimar un sistema de medición y relaciones. Consultar el reporte implica evaluar criterios de calidad y tomar decisiones. Por ello, el aprendizaje de **SmartPLS4** debe ir acompañado de comprensión **PLS-SEM**. El software facilita la visualización y cálculo del modelo, pero no sustituye la justificación teórica ni la interpretación académica. Esta distinción debe quedar clara para evitar una enseñanza meramente instrumental del programa. El objetivo no es formar usuarios que solo opriman botones, sino investigadores capaces de comprender qué significa cada botón dentro de una ruta de análisis estructural (Hair et al., 2021).

PLS-SEM resulta especialmente útil cuando se analizan modelos orientados a explicar fenómenos emergentes o aplicados. En campos como **inteligencia artificial, sostenibilidad, innovación educativa, emprendimiento o ciudadanía digital**, los constructos suelen estar en desarrollo y requieren validación empírica progresiva. Esta condición hace necesario utilizar herramientas que permitan evaluar simultáneamente la calidad de los indicadores y las relaciones entre variables latentes. Sin embargo, la flexibilidad del método no debe confundirse con permisividad metodológica. Aunque **PLS-SEM** puede ser adecuado para modelos complejos, investigación exploratoria o enfoques predictivos, el investigador debe justificar por qué lo utiliza, cómo especifica los constructos y qué criterios aplicará para evaluar resultados. En el caso **CAITIZEN**, el modelo busca analizar una ciudadanía sostenible asistida por **IA** a partir de dimensiones críticas, éticas, algorítmicas, colaborativas y metacognitivas. Esta estructura exige un método capaz de examinar tanto la medición de cada constructo como sus relaciones estructurales. Por ello, **PLS-SEM** ofrece una ruta pertinente, siempre que el análisis sea transparente, ordenado y metodológicamente defendible.

Juan Mejía Trejo

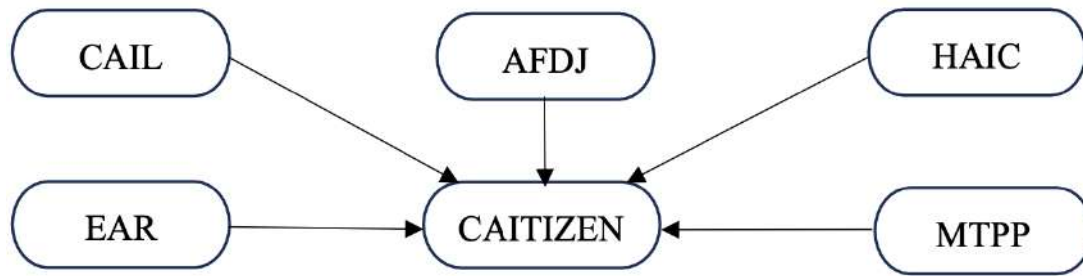
La contribución del modelo no dependerá únicamente del software empleado, sino de la forma en que el investigador explique, mida, evalúe e interprete el fenómeno estudiado (Hair et al., 2019).

En síntesis, **PLS-SEM** debe entenderse como una ruta aplicada que integra construcción conceptual, medición, estimación, evaluación e interpretación. Su uso en **SmartPLS4** no inicia con el cálculo del algoritmo, sino con una comprensión previa del modelo que se desea analizar. Para el lector, el primer aprendizaje consiste en reconocer que cada etapa cumple una función específica: los datos permiten observar indicadores; los indicadores representan constructos; los constructos integran el modelo de medición; las flechas expresan hipótesis; el algoritmo estima relaciones; los criterios de calidad diagnostican la medición y la estructura; y el **Bootstrapping** permite evaluar significancia estadística. En este libro, esa ruta se aplicará progresivamente al modelo **CAITIZEN**, de modo que el lector no solo aprenda a usar comandos de **SmartPLS4**, sino también a comprender qué significa cada decisión dentro del análisis **PLS-SEM**. Así, el capítulo inicial debe ofrecer una base clara, breve y orientadora: qué es **PLS-SEM**, por qué se utiliza, cómo se articula con **SmartPLS4** y qué papel tendrá en la evaluación del modelo **CAITIZEN**. Esta comprensión permitirá que los capítulos posteriores avancen con mayor claridad desde la preparación del archivo hasta el reporte académico de resultados (Hair et al., 2021).

Modelo de medición y modelo estructural

Antes de distinguir entre **modelo de medición** y **modelo estructural**, es necesario presentar brevemente el caso que servirá como hilo conductor del libro. El modelo **CAITIZEN** fue publicado inicialmente como un **modelo cualitativo** para comprender la **ciudadanía asistida por inteligencia artificial en formación** en estudiantes universitarios. En esa primera formulación, apoyada en **ATLAS.ti**, el modelo se orientó a analizar cómo la inteligencia artificial se relaciona con la **alfabetización crítica en IA**, la **responsabilidad ética**, la **justicia de datos**, la **colaboración humano-IA** y la **transparencia metacognitiva** en el uso de prompts. En dicho antecedente, **CAITIZEN** se planteó como un **marco multidisciplinario de innovación social** vinculado con la formación ciudadana, la educación superior y los **Objetivos de Desarrollo Sostenible**. A partir de esa base cualitativa, este libro adopta la precisión conceptual **CAITIZEN —Citizenship Assisted by Artificial Intelligence for Sustainable, Ethical, and Networked formation—**, entendida como **ciudadanía asistida por inteligencia artificial para una formación sostenible, ética y en red**. Esta denominación no sustituye el antecedente cualitativo, sino que **precisa su alcance para la fase cuantitativa PLS-SEM**, al hacer explícitas sus dimensiones formativas, sostenibles, éticas y relacionales. Ver **Figura 1.1** (Mejía-Trejo, 2025).

Figura 1.1. Modelo CAITIZEN



Notas: **CAIL**. Critical Artificial Intelligence Literacy; **EAR**. Ethical Awareness and Responsibility; **AFDJ**. Awareness of Fairness and Data Justice; **HAIC**. Human–AI Creative Collaboration; **MTPP**. Metacognitive Transparency in Prompting Practices.

Fuente. Mejía-Trejo (2025)

El modelo **CAITIZEN** se integra por **cinco constructos antecedentes** y un **constructo dependiente**. Los constructos antecedentes son: **CAIL**, *Critical Artificial Intelligence Literacy* o **alfabetización crítica en inteligencia artificial**; **EAR**, *Ethical Awareness and Responsibility* o **conciencia ética y responsabilidad**; **AFDJ**, *Awareness of Fairness and Data Justice* o **conciencia de equidad y justicia de datos**; **HAIC**, *Human–AI Creative Collaboration* o **colaboración creativa humano–IA**; y **MTPP**, *Metacognitive Transparency in Prompting Practices* o **transparencia metacognitiva en prácticas de prompting**. El constructo dependiente es **CAITIZEN**, entendido en este libro como **ciudadanía sostenible asistida por inteligencia artificial**. En la formulación cualitativa previa, estas dimensiones funcionaron como **categorías analíticas** para interpretar significados, discursos y relaciones conceptuales; no se plantearon todavía como relaciones causales estimadas estadísticamente. En este libro, esas mismas dimensiones se traducen en **constructos latentes medibles** mediante indicadores y en **relaciones estructurales evaluables** con **SmartPLS4**. Por tanto, el paso de lo cualitativo a lo cuantitativo no cambia la esencia del modelo; la convierte en una arquitectura empírica susceptible de estimación mediante **PLS-SEM** (Mejía-Trejo, 2025).

El análisis **PLS-SEM** se organiza a partir de dos componentes fundamentales: el **modelo de medición** y el **modelo estructural**. El **modelo de medición** establece la relación entre los **indicadores observables** y los **constructos latentes**; el **modelo estructural** representa las relaciones hipotéticas entre los constructos. Esta distinción es esencial porque muchos fenómenos estudiados en las ciencias de la administración no pueden medirse de forma directa. Conceptos como **alfabetización crítica en inteligencia artificial**, **responsabilidad ética**, **justicia de datos**, **colaboración humano–IA**, **transparencia metacognitiva** o **ciudadanía sostenible asistida por inteligencia artificial** son abstracciones que requieren ser traducidas en **ítems**, **escalas** o **indicadores empíricos**. En el modelo **CAITIZEN**, esta lógica permite convertir dimensiones conceptuales previamente sustentadas en variables analizables mediante **SmartPLS4**. Así, **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** no son simples nombres dentro del software, sino **constructos latentes** que deben estar

Juan Mejía Trejo

respaldados por indicadores coherentes con su definición teórica. Comprender esta diferencia evita que el lector confunda la **representación gráfica del modelo** con su **justificación metodológica**. En **PLS-SEM**, dibujar un modelo no equivale a validarlo; primero debe comprobarse que los indicadores representan adecuadamente a sus constructos y, después, que las relaciones entre constructos tienen sentido teórico y empírico. Por tanto, el análisis comienza con una pregunta básica: **qué se mide y cómo se conecta con lo que se desea explicar** (Hair et al., 2021).

El **modelo de medición** responde a la pregunta sobre la calidad con la que los indicadores representan a un constructo. En **SmartPLS4**, esta relación se observa cuando el investigador asigna indicadores a una **variable latente**. Por ejemplo, los indicadores **CAIL1–CAIL10** se asignan al constructo **CAIL**, mientras que **EAR1–EAR10** se asignan a **EAR**. Esta asignación no debe hacerse únicamente porque los nombres de las variables coincidan, sino porque existe una **correspondencia conceptual** entre cada ítem y el constructo al que pertenece. Si un indicador no representa adecuadamente el contenido teórico del constructo, puede afectar la interpretación de las **cargas externas**, la **confiabilidad interna** y la **validez convergente**. Por ello, la evaluación del modelo de medición no es un trámite técnico, sino una **condición previa para cualquier interpretación estructural**. En investigaciones de administración, donde los constructos suelen representar percepciones, capacidades, actitudes o comportamientos, esta etapa adquiere especial relevancia. Una escala mal especificada puede producir resultados aparentemente significativos, pero metodológicamente débiles. En el caso de **CAITIZEN**, cada bloque de indicadores debe mostrar que captura una dimensión particular del fenómeno: **alfabetización crítica, responsabilidad ética, justicia de datos, colaboración humano–IA, transparencia metacognitiva y ciudadanía asistida por IA**. La validez del análisis posterior dependerá de que esta primera correspondencia entre indicador y constructo sea defendible (Gefen et al., 2000).

El **modelo estructural** responde a una pregunta diferente: **cómo se relacionan los constructos entre sí**. Mientras el modelo de medición conecta indicadores con constructos, el modelo estructural conecta constructos con otros constructos mediante **relaciones hipotéticas**. En el caso de **CAITIZEN**, las rutas **CAIL → CAITIZEN**, **EAR → CAITIZEN**, **AFDJ → CAITIZEN**, **HAIC → CAITIZEN** y **MTPP → CAITIZEN** representan hipótesis estructurales. Cada flecha indica que un **constructo antecedente** puede contribuir a explicar la **variable dependiente**. Sin embargo, una flecha en **SmartPLS4** no debe entenderse como una simple conexión visual; representa una **afirmación teórica** que debe justificarse con literatura, lógica conceptual o evidencia previa. El antecedente cualitativo del modelo **CAITIZEN** ayuda precisamente a sostener esta transición, porque mostró que las cinco dimensiones se articulan conceptualmente alrededor de una ciudadanía asistida por **IA** entendida como **sistema ético, cognitivo y social**. El modelo estructural, entonces, traduce esa lógica conceptual previa en relaciones evaluables cuantitativamente. **PLS-SEM** permite estimar la **magnitud, dirección y significancia** de esas rutas; pero su validez depende de la coherencia conceptual que las sostiene. Por ello, antes de interpretar

coeficientes de trayectoria, el investigador debe explicar por qué cada relación forma parte del modelo (Chin, 1998).

La diferencia entre **modelo de medición** y **modelo estructural** permite ordenar la secuencia metodológica del análisis. Primero se identifican los **constructos** que integran el modelo; después se definen los **indicadores** que los representan; posteriormente se establecen las **relaciones esperadas** entre constructos. Esta secuencia impide que el investigador inicie directamente en el software sin haber definido una **arquitectura conceptual previa**. En **SmartPLS4**, el lector observará que los indicadores provienen del archivo de datos importado, mientras que los constructos se crean en el lienzo gráfico. Asignar indicadores a un constructo corresponde al **modelo de medición**; dibujar flechas entre constructos corresponde al **modelo estructural**. Por ejemplo, vincular **CAIL1–CAIL10** con **CAIL** forma parte de la medición; trazar la ruta **CAIL** → **CAITIZEN** forma parte de la estructura. Esta separación es pedagógicamente importante porque cada acción del software tiene una función metodológica distinta. No es lo mismo revisar si los indicadores miden correctamente que analizar si un constructo explica a otro. En el libro, esta distinción permitirá que el lector comprenda por qué no debe saltar inmediatamente a las hipótesis después de ejecutar el algoritmo. Antes de interpretar rutas, **t-values** o **p-values**, debe verificarse la **calidad de la medición**. Solo cuando los constructos están adecuadamente medidos, el modelo estructural puede interpretarse con mayor rigor (Hair et al., 2021).

En el caso de modelos **reflectivos**, como se plantea para **CAITIZEN**, el modelo de medición supone que los indicadores son **manifestaciones del constructo latente**. Esto significa que el constructo explica la variación observada en sus indicadores. Si una persona tiene mayor **alfabetización crítica en inteligencia artificial**, se esperaría que responda de manera consistente a los distintos ítems asociados con **CAIL**. Si tiene mayor **responsabilidad ética ante la IA**, se esperaría un patrón coherente en los indicadores de **EAR**. Esta lógica reflectiva implica que los indicadores deben compartir cierta consistencia conceptual y estadística. Por ello, en los capítulos posteriores se revisarán criterios como **cargas externas**, **confiabilidad interna**, **validez convergente** y **validez discriminante**. Las cargas externas permiten observar qué tan bien cada indicador representa al constructo. La confiabilidad interna revisa si los indicadores funcionan de manera consistente. La validez convergente analiza si los indicadores comparten suficiente varianza. La validez discriminante permite verificar si cada constructo es empíricamente distinto de los demás. En **CAITIZEN**, estos criterios son necesarios para demostrar que **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** no son etiquetas arbitrarias, sino constructos medidos con suficiente calidad. Por tanto, el modelo reflectivo requiere una lectura cuidadosa de la relación entre constructo e indicadores, antes de avanzar hacia la interpretación de relaciones estructurales (Gefen et al., 2000).

El **modelo estructural** se interpreta después de evaluar la calidad del **modelo de medición**. Esta regla metodológica es fundamental porque las relaciones entre constructos solo son defendibles si los constructos fueron medidos adecuadamente.

Juan Mejía Trejo

En el modelo **CAITIZEN**, no sería prudente afirmar que **AFDJ** influye en **CAITIZEN** si antes no se confirma que los indicadores de **AFDJ** representan de forma consistente la conciencia sobre **justicia de datos** y **equidad algorítmica**. Tampoco sería adecuado interpretar la influencia de **HAIC** si sus indicadores no reflejan realmente **colaboración humano-IA**. Por ello, **PLS-SEM** sigue una lógica progresiva: primero se revisa **cómo se mide** y después **cómo se relaciona**. En el modelo estructural se analizan criterios como **colinealidad entre predictores**, **coeficientes de trayectoria**, **R^2** , **f^2** y **significancia estadística mediante *Bootstrapping***. Estos elementos permiten estimar qué tanto los constructos antecedentes contribuyen a explicar el constructo dependiente. Sin embargo, los resultados estructurales no sustituyen la evaluación de la medición; dependen de ella. Un modelo puede mostrar rutas significativas, pero si los constructos no tienen validez suficiente, la interpretación será débil. Por eso, el análisis estructural debe entenderse como una **segunda etapa evaluativa**. La primera pregunta es si los constructos están bien medidos; la segunda, si las relaciones propuestas son empíricamente relevantes y teóricamente defendibles (Chin, 1998).

La distinción entre ambos modelos también ayuda a evitar errores comunes en el uso de **SmartPLS4**. Un error frecuente consiste en observar los **coeficientes de las flechas** y asumir que el modelo ya fue validado. Otro error consiste en revisar únicamente **p-values** sin considerar **cargas**, **confiabilidad**, **AVE** o **validez discriminante**. También puede ocurrir que el investigador elimine indicadores con cargas bajas sin examinar su relevancia conceptual. Estos errores surgen cuando se interpreta el modelo como una salida automática del software y no como una estructura metodológica compuesta por decisiones sucesivas. En **CAITIZEN**, por ejemplo, un indicador puede tener una carga menor a la esperada, pero su eliminación debe evaluarse con prudencia: puede afectar el contenido conceptual del constructo o reducir la cobertura de una dimensión relevante. De igual forma, una ruta no significativa no necesariamente invalida todo el modelo; puede indicar que una relación específica no tiene suficiente apoyo empírico en la muestra analizada. Por eso, separar **medición** y **estructura** permite diagnosticar con mayor precisión dónde se encuentra el problema: en los indicadores, en los constructos, en la relación hipotética o en la especificación general del modelo. Esta claridad es indispensable para que el lector aprenda a tomar **decisiones metodológicas**, no solo a leer resultados generados por **SmartPLS4** (Hair et al., 2021).

SmartPLS4 facilita la visualización del **modelo de medición** y del **modelo estructural** en un mismo entorno gráfico. En el lienzo de trabajo, los constructos aparecen como nodos, los indicadores se vinculan a cada constructo y las flechas representan relaciones entre variables latentes. Esta visualización ayuda al lector a comprender la estructura general del modelo. Sin embargo, la facilidad visual puede producir una falsa impresión de simplicidad. Crear un constructo no significa haberlo definido correctamente; asignar indicadores no garantiza validez de medición; conectar constructos no demuestra una relación causal; ejecutar el algoritmo no convierte automáticamente el modelo en válido. Por ello, cada elemento gráfico debe interpretarse como representación de una **decisión metodológica**. En **CAITIZEN**, el

constructo **CAIL** representa alfabetización crítica en **IA**; sus indicadores representan manifestaciones observables de esa alfabetización; la flecha hacia **CAITIZEN** expresa una hipótesis sobre su relación con la ciudadanía asistida por IA. Lo mismo ocurre con **EAR**, **AFDJ**, **HAIC** y **MTPP**. Esta lectura permite que **SmartPLS4** no sea usado como una **caja negra**, sino como una plataforma de modelación estructural que exige razonamiento previo y evaluación posterior. El lector debe aprender que la imagen del modelo es solo el punto de partida visual; la validez del modelo depende de su **justificación conceptual**, su **calidad de medición** y su **evaluación empírica** (Gefen et al., 2000).

La relación entre **modelo de medición** y **modelo estructural** también determina la forma en que deben reportarse los resultados. Un informe **PLS-SEM** no debe limitarse a presentar coeficientes de trayectoria y **p-values**. Primero debe reportar la calidad del modelo de medición: **cargas externas**, **confiabilidad interna**, **AVE** y **validez discriminante**. Después debe reportar el modelo estructural: **VIF interno**, **coeficientes de trayectoria**, **R^2** , **f^2** , relevancia predictiva cuando corresponda y resultados de **Bootstrapping**. Esta secuencia permite que el lector del estudio comprenda que las hipótesis estructurales se interpretaron sobre una base de medición previamente evaluada. En el caso **CAITIZEN**, el reporte deberá mostrar primero si **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** son constructos confiables y válidos. Solo después podrá indicarse qué predictores explican mejor a **CAITIZEN**, cuáles rutas son significativas y cuál es la capacidad explicativa del modelo. Esta lógica evita conclusiones apresuradas y fortalece la **transparencia metodológica**. Además, permite distinguir entre **significancia estadística** y **relevancia sustantiva**. Una ruta puede ser significativa, pero tener un efecto pequeño; otra puede no ser significativa, pero sugerir una relación conceptual que conviene discutir. Por ello, el reporte debe articular medición, estructura e interpretación académica, no solo reproducir tablas del software (Chin, 1998).

En síntesis, el **modelo de medición** y el **modelo estructural** son los dos pilares del análisis **PLS-SEM**. El primero permite transformar conceptos abstractos en **indicadores observables**; el segundo permite evaluar **relaciones hipotéticas entre constructos**. En el modelo **CAITIZEN**, esta distinción organiza toda la ruta metodológica del libro. La preparación del archivo asegura que los indicadores estén codificados correctamente. La construcción del modelo permite asignar indicadores a constructos y trazar relaciones estructurales. La ejecución del algoritmo proporciona la estimación inicial. La evaluación del modelo de medición permite revisar cargas externas, confiabilidad y validez. La evaluación del modelo estructural permite analizar colinealidad, coeficientes de trayectoria y capacidad explicativa. Finalmente, el **Bootstrapping** permite contrastar hipótesis. Comprender la diferencia entre **medir** y **relacionar** es indispensable para interpretar correctamente el análisis. El lector debe tener claro que **PLS-SEM** no consiste solo en dibujar flechas ni en obtener reportes automáticos; consiste en demostrar que los constructos están bien medidos y que las relaciones entre ellos son teórica y empíricamente defendibles. Por ello, antes de avanzar a los controles específicos de **SmartPLS4**, resulta necesario dominar esta

distinción conceptual básica: el **modelo de medición responde cómo se mide**; el **modelo estructural responde cómo se relaciona lo medido** (Hair et al., 2021).

Entorno básico de SmartPLS4

El **entorno básico de SmartPLS4** constituye el espacio operativo donde el investigador organiza, construye, estima, revisa e interpreta un modelo **PLS-SEM**. En este libro, dicho entorno se utilizará para desarrollar el modelo **CAITIZEN**, integrado por los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITIZEN**. Antes de avanzar hacia la preparación del archivo, la construcción del modelo o la ejecución del algoritmo, es necesario que el lector reconozca los elementos principales de la plataforma. **SmartPLS4** no debe entenderse únicamente como un software para “calcular resultados”, sino como un entorno de modelación que permite representar visualmente una arquitectura metodológica. En él se integran el **Workspace**, el **proyecto**, el **archivo de datos**, el **panel de indicadores**, el **área gráfica**, los **controles de construcción**, los **controles de cálculo** y el **reporte completo**. Cada elemento cumple una función dentro de la ruta **PLS-SEM**: organizar el análisis, importar los indicadores, crear constructos, asignar ítems, dibujar relaciones, ejecutar estimaciones y revisar criterios de calidad. Por ello, el primer acercamiento a **SmartPLS4** debe ser introductorio y metodológico, no solamente instrumental. El propósito de este apartado es que el lector comprenda el entorno básico antes de iniciar las operaciones detalladas que se desarrollarán en los capítulos posteriores (SmartPLS GmbH, s. f.-e), con la siguiente consideración de elemento:

El **Workspace**, entendido como el espacio general donde se organizan proyectos, modelos, datos y resultados. En términos metodológicos, el **Workspace** funciona como una carpeta de trabajo que permite mantener ordenada la investigación. Para el modelo **CAITIZEN**, crear un **Workspace** con ese nombre ayuda a conservar una trazabilidad clara del análisis. Dentro de este espacio se podrán almacenar el archivo de datos importado, el modelo gráfico, las estimaciones del algoritmo **PLS-SEM**, los resultados de **Bootstrapping** y los reportes generados. Esta organización es importante porque un proyecto **PLS-SEM** puede producir diferentes versiones, ajustes y salidas estadísticas. Si el investigador no ordena desde el inicio sus archivos, puede confundir versiones del modelo, bases de datos o resultados. Por eso, el **Workspace** no es un detalle administrativo menor, sino un componente de orden metodológico. En una investigación reproducible, debe saberse con claridad qué datos se usaron, qué modelo se estimó, bajo qué configuración se ejecutó el algoritmo y qué reporte se interpretó. En este sentido, el uso adecuado del entorno digital se vincula con la calidad del procedimiento analítico. La lógica aplicada de **PLS-SEM** exige que la organización técnica del proyecto acompañe la organización conceptual del modelo (Hair et al., 2021).

El **proyecto**, que representa la unidad específica de análisis dentro del **Workspace**. Mientras el **Workspace** funciona como el espacio general, el proyecto agrupa los elementos concretos de una investigación. En el caso **CAITIZEN**, el proyecto deberá contener el archivo de datos, el modelo **PLS-SEM** y los reportes generados. Esta

Juan Mejía Tréjo

diferencia entre **Workspace** y proyecto ayuda a evitar confusiones: el **Workspace** puede contener varios proyectos, pero cada proyecto corresponde a un análisis particular. Por ejemplo, un investigador podría tener en un mismo **Workspace** proyectos distintos como **CAITIZEN**, **SUSTAINN-4E** o **AI4DMS**, cada uno con sus propios datos y modelos. No se busca explicar todavía el procedimiento paso a paso para crear el proyecto; eso corresponderá al Capítulo 2. Sin embargo, sí es conveniente que el lector comprenda su función. El proyecto permite concentrar los materiales del análisis y facilita la continuidad del trabajo. También permite que **SmartPLS4** relacione el archivo de datos con el modelo que será construido. Por ello, nombrar adecuadamente el proyecto tiene importancia práctica y metodológica. Un nombre claro evita ambigüedades en capturas, reportes y resultados. Para este libro, el nombre **CAITIZEN** identifica el caso aplicado que será desarrollado progresivamente mediante **SmartPLS4** (SmartPLS GmbH, s. f.-d).

El **archivo de datos**, que contiene las variables observables utilizadas como indicadores del modelo. En **SmartPLS4**, el archivo de datos se importa al proyecto y sirve como base empírica para estimar el modelo **PLS-SEM**. En el caso **CAITIZEN**, cada columna del archivo representa un indicador: **CAIL1–CAIL10**, **EAR1–EAR10**, **AFDJ1–AFDJ10**, **HAIC1–HAIC10**, **MTPP1–MTPP10** y **CAITIZEN1–CAITIZEN10**. Cada fila representa un caso o participante. Esta estructura es fundamental porque el software no interpreta el sentido conceptual de los indicadores por sí mismo; solo reconoce nombres de variables y valores numéricos. Por ello, la preparación del archivo requiere cuidado: los encabezados deben ser claros, los valores deben corresponder a la escala utilizada, no deben existir columnas innecesarias y los datos perdidos deben estar controlados. Aunque el tratamiento detallado de la base se desarrollará en el Capítulo 2, es necesario anticipar que una base mal preparada puede afectar la estimación posterior. La calidad del análisis estructural no depende únicamente del algoritmo, sino también de la calidad de los datos que se incorporan al modelo. En investigaciones **SEM**, el rigor metodológico exige coherencia entre conceptualización, medición, codificación de datos y análisis empírico (Gefen et al., 2000).

El **panel de indicadores**, donde **SmartPLS4** muestra las variables observables disponibles después de importar el archivo de datos. Este panel permite seleccionar los indicadores y asignarlos a los constructos correspondientes. En el modelo **CAITIZEN**, el panel mostrará nombres como **CAIL1**, **CAIL2**, **EAR1**, **AFDJ1**, **HAIC1**, **MTPP1** o **CAITIZEN1**. Su utilidad es práctica, pero también metodológica: permite vincular los ítems observados con los constructos latentes definidos previamente. Esta asignación no debe realizarse de manera automática ni solo por semejanza de nombres. Cada indicador debe corresponder conceptualmente al constructo que pretende medir. Por ejemplo, un ítem de **EAR** no debe asignarse a **AFDJ** solo porque ambos se relacionen con temas éticos o normativos. Si se asignan indicadores al constructo equivocado, se alteran las cargas externas, la confiabilidad, la validez convergente y la validez discriminante. Por ello, el panel de indicadores es el puente entre la base empírica y el modelo de medición. En términos didácticos, permite que el lector observe cómo las respuestas del cuestionario se convierten en elementos del

modelo **PLS-SEM**. Esta operación será desarrollada con detalle en el Capítulo 3, pero desde ahora debe quedar claro que asignar indicadores es una decisión de medición, no solo una acción gráfica (Hair et al., 2021).

El **área gráfica de modelación**, también llamada lienzo o espacio de construcción del modelo. En esta área se representan los constructos latentes, se asignan los indicadores y se dibujan las relaciones estructurales. Para el modelo **CAITIZEN**, el área gráfica permitirá visualizar los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, así como las rutas que conectan a los cinco predictores con la variable dependiente. Esta visualización es una de las ventajas pedagógicas de **SmartPLS4**, porque permite observar simultáneamente el modelo de medición y el modelo estructural. Sin embargo, la facilidad visual no debe confundirse con validación metodológica. Dibujar un constructo no significa que esté bien definido; asignar indicadores no significa que sean válidos; trazar una flecha no significa que exista influencia empírica; ejecutar el algoritmo no significa que el modelo sea adecuado. El área gráfica solo representa la arquitectura propuesta. La calidad del modelo se evaluará después mediante cargas, confiabilidad, **AVE**, validez discriminante, **VIF**, coeficientes de trayectoria, **R²**, **f²** y **Bootstrapping**. Por ello, el lector debe aprender a observar el lienzo como una representación metodológica: los nodos representan constructos, los indicadores representan medición y las flechas representan hipótesis estructurales (SmartPLS GmbH, s. f.-e).

Los **controles básicos de construcción**, es decir, las herramientas que permiten crear, seleccionar, mover, conectar y revisar elementos dentro del modelo. Entre estos controles se encuentran funciones como **Create model**, **Select**, **Latent variable**, **Connect** y **Construct properties**. En el caso **CAITIZEN**, estas herramientas permitirán construir progresivamente el modelo: primero se crea el modelo **PLS-SEM**; después se incorporan los constructos latentes; luego se asignan los indicadores; posteriormente se trazan las relaciones estructurales. **Select** permite organizar visualmente los elementos; **Latent variable** permite crear constructos; **Connect** permite dibujar relaciones entre variables latentes; y **Construct properties** ayuda a revisar nombre, tipo de medición y configuración básica del constructo. Estos controles no deben explicarse únicamente como comandos del programa. Cada control expresa una decisión metodológica. Crear un constructo equivale a representar una variable latente; asignar indicadores equivale a configurar el modelo de medición; conectar constructos equivale a formular una hipótesis estructural. En consecuencia, **SmartPLS4** traduce gráficamente una lógica metodológica previa. Desde una perspectiva clásica del **enfoque PLS**, el modelo de rutas debe representar relaciones teóricamente justificadas y no conexiones arbitrarias dibujadas en el software (Chin, 1998).

Los **controles de cálculo y estimación**. Una vez construido el modelo, **SmartPLS4** permite ejecutar procedimientos como **PLS-SEM Algorithm** y **Bootstrapping**. El algoritmo **PLS-SEM** genera la estimación inicial del modelo y permite obtener resultados como cargas externas, coeficientes de trayectoria, confiabilidad, **AVE**, **VIF**, **R²** y otros criterios de evaluación. **Bootstrapping**, por su parte, permite evaluar la

significancia estadística de los resultados mediante remuestreo, generando **t-values**, **p-values** e intervalos de confianza. Aquí, solo se presenta una visión general de estos controles; su ejecución paso a paso se desarrollará en capítulos posteriores. No obstante, conviene que el lector comprenda desde ahora que **Calculate** no es simplemente un botón para obtener resultados, sino la entrada a procedimientos distintos de estimación y contraste. En el caso **CAITIZEN**, el algoritmo permitirá obtener la primera evaluación del modelo y el **Bootstrapping** permitirá contrastar las hipótesis estructurales: **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN**. La documentación oficial de SmartPLS describe el algoritmo **PLS-SEM** como procedimiento para estimar el modelo de rutas a partir de los datos de indicadores y del modelo especificado (SmartPLS GmbH, s. f.-c).

El **reporte completo**, también identificado como **Full report** u opción equivalente de consulta de resultados. Este reporte organiza los resultados generados por **SmartPLS4** y permite revisar criterios del modelo de medición y del modelo estructural. En el caso **CAITIZEN**, será necesario consultar cargas externas, confiabilidad interna, **AVE**, validez discriminante, **VIF**, coeficientes estructurales, **R²**, **f²**, **t-values**, **p-values** e intervalos de confianza. El reporte no debe leerse de manera desordenada. Primero se revisa el modelo de medición; después, el modelo estructural; finalmente, la significancia y la capacidad explicativa o predictiva. Esta secuencia evita interpretar hipótesis antes de comprobar la calidad de los constructos. El reporte completo no es solo una salida automática del software, sino una herramienta de diagnóstico metodológico. Permite tomar decisiones sobre conservación de indicadores, confiabilidad de constructos, validez convergente, validez discriminante, colinealidad y relevancia de las rutas estructurales. En un reporte académico, el investigador no debe copiar tablas sin interpretación; debe explicar qué significan los resultados y cómo respaldan o limitan el modelo propuesto. Por ello, aprender a leer el reporte es tan importante como aprender a construir el modelo. El uso adecuado de **PLS-SEM** implica saber cuándo utilizar el método y cómo reportar sus resultados con rigor (Hair et al., 2019).

La **relación entre navegación, vistas y secuencia de trabajo** dentro de **SmartPLS4**. El usuario interactúa con diferentes espacios: una vista para proyectos, una vista para datos, una vista para modelación y una vista para reportes. Esta organización permite que el análisis avance por etapas. Primero se ubica el proyecto; luego se revisan los datos; posteriormente se construye el modelo; finalmente se consultan los resultados. En el modelo **CAITIZEN**, esta lógica será evidente: el lector comenzará identificando el proyecto, verificará que el archivo de datos esté importado, construirá el modelo gráfico y después abrirá los reportes del algoritmo y **Bootstrapping**. Esta secuencia debe aprenderse desde el inicio porque ayuda a evitar errores frecuentes. Por ejemplo, no se puede construir correctamente el modelo si el archivo de datos no fue importado; no se pueden asignar indicadores si las variables no aparecen en el panel; no se puede interpretar **Bootstrapping** si antes no se ejecutó el modelo adecuado. **SmartPLS4** organiza visualmente el proceso, pero el investigador debe comprender qué etapa está realizando. Por eso, este apartado debe

presentar el entorno como una ruta ordenada de trabajo. La documentación oficial agrupa recursos, tutoriales, algoritmos y técnicas que apoyan esta progresión desde la creación del proyecto hasta la consulta de resultados (SmartPLS GmbH, s. f.-a).

El **entorno básico de SmartPLS4** debe entenderse como una estructura integrada de análisis **PLS-SEM**. El **Workspace** organiza el espacio general; el **proyecto** concentra el análisis específico; el **archivo de datos** contiene los indicadores observables; el **panel de indicadores** permite vincular variables con constructos; el **área gráfica** representa el modelo; los **controles de construcción** permiten crear y conectar elementos; los **controles de cálculo** permiten estimar resultados; y el **reporte completo** permite revisar criterios de calidad. En el modelo **CAITIZEN**, cada componente cumplirá una función precisa a lo largo del libro. El Capítulo 2 se enfocará en preparar e importar la base de datos; el Capítulo 3 desarrollará la construcción del modelo; el Capítulo 4 ejecutará el algoritmo y revisará criterios iniciales; los capítulos posteriores evaluarán medición, validez discriminante, modelo estructural, **Bootstrapping** y reporte académico. Así, este apartado cumple una función introductoria: familiarizar al lector con el entorno en el que se aplicará la ruta metodológica. **SmartPLS4** facilita el análisis, pero no sustituye el criterio del investigador. La calidad del modelo dependerá de que cada acción operativa esté respaldada por una decisión metodológica clara. Por ello, el software debe usarse como una plataforma para expresar, estimar y evaluar modelos estructurales, no como una caja negra de resultados automáticos (Gefen et al., 2000).

Controles básicos para construir el modelo

Los **controles básicos para construir el modelo** en **SmartPLS4** permiten transformar una propuesta conceptual en una representación gráfica evaluable mediante **PLS-SEM**. En el caso del modelo **CAITIZEN**, estos controles se utilizarán para crear los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**; asignar sus indicadores; configurar su tipo de medición; y establecer las rutas estructurales correspondientes. Es importante comprender que estos controles no son simples herramientas visuales, sino recursos operativos que expresan decisiones metodológicas. Crear un constructo significa representar una variable latente; asignar indicadores significa configurar el modelo de medición; conectar constructos significa formular relaciones estructurales; y revisar propiedades significa verificar que el modelo esté correctamente especificado antes de estimarlo. Por ello, este apartado introduce los controles principales que el lector utilizará posteriormente en el Capítulo 3, sin desarrollar todavía el procedimiento completo. La intención es familiarizar al lector con la lógica de construcción del modelo antes de trabajar directamente con capturas, menús y configuraciones específicas. **SmartPLS4** facilita la construcción visual del modelo, pero el investigador debe comprender que cada acción dentro del software debe corresponder a una decisión conceptual previa. Así, construir un modelo no significa únicamente dibujar elementos, sino traducir una estructura teórica en una arquitectura empíricamente estimable (SmartPLS GmbH, s. f.-e), con los siguientes controles a considerar:

Juan Mejía Tréjo

Create model, que permite iniciar la construcción de un nuevo modelo dentro de un proyecto previamente creado. En términos metodológicos, este control marca el paso entre tener una base de datos importada y comenzar a representar el modelo **PLS-SEM**. Para el caso **CAITIZEN**, después de contar con el proyecto y el archivo de datos, el investigador debe crear un modelo que permita trabajar gráficamente con los constructos y sus relaciones. En esta etapa se define el tipo de modelo, normalmente **PLS-SEM**, y se asigna un nombre al modelo. El nombre debe ser claro, breve y congruente con la investigación. En este libro, el modelo puede denominarse **CAITIZEN**, mientras que etiquetas operativas como **CAITIZEN_PLS** pueden utilizarse únicamente para identificar archivos o versiones dentro del software. Esta distinción es importante porque el nombre conceptual del modelo no debe confundirse con la etiqueta técnica usada para organizarlo. Crear el modelo implica abrir un espacio donde se representará la relación entre indicadores y constructos, así como las hipótesis estructurales. Desde la lógica **PLS-SEM**, el modelo gráfico debe expresar un sistema de medición y relaciones previamente justificado, no una combinación improvisada de variables (Hair et al., 2021).

Latent variable, utilizado para crear constructos latentes dentro del área gráfica. Un constructo latente representa una variable que no se observa directamente, pero que puede medirse mediante indicadores. En el modelo **CAITIZEN**, los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** son variables latentes porque expresan dimensiones conceptuales relacionadas con alfabetización crítica en **IA**, responsabilidad ética, justicia de datos, colaboración humano-**IA**, transparencia metacognitiva y ciudadanía sostenible asistida por **IA**. Al crear cada constructo en **SmartPLS4**, el investigador no está generando una variable estadística nueva sin fundamento, sino representando una categoría conceptual previamente definida. Por esta razón, el control **Latent variable** debe usarse con prudencia metodológica. No se deben crear constructos únicamente porque existen ítems disponibles en la base de datos; cada constructo debe tener una función clara dentro del modelo. Además, su nombre debe ser consistente con la teoría, el cuestionario y el reporte final. Una denominación confusa puede afectar la comprensión del lector y generar errores en la interpretación de resultados. En el enfoque clásico **PLS**, los constructos y sus relaciones deben reflejar una lógica de rutas justificada por teoría o por objetivos analíticos claros (Chin, 1998).

El **panel de indicadores**, desde el cual se seleccionan las variables observables importadas desde la base de datos. Este panel permite asignar indicadores a los constructos latentes creados en el modelo. En **CAITIZEN**, por ejemplo, los indicadores **CAIL1–CAIL10** deben vincularse al constructo **CAIL**; **EAR1–EAR10** al constructo **EAR**; **AFDJ1–AFDJ10** al constructo **AFDJ**; **HAIC1–HAIC10** al constructo **HAIC**; **MTPP1–MTPP10** al constructo **MTPP**; y **CAITIZEN1–CAITIZEN10** al constructo dependiente **CAITIZEN**. Esta asignación constituye la construcción del modelo de medición. Aunque la operación puede realizarse mediante selección y arrastre, su sentido metodológico es más profundo: cada indicador debe representar una manifestación observable del constructo al que se asigna. Si un indicador se coloca en el constructo equivocado, se altera la medición y pueden afectarse las cargas externas,

la confiabilidad, la validez convergente y la validez discriminante. Por ello, antes de asignar indicadores, el investigador debe revisar que los nombres de las columnas sean claros y que exista correspondencia entre cuestionario, archivo de datos y modelo teórico. La práctica rigurosa de **SEM** exige coherencia entre conceptualización, medición y análisis empírico (Gefen et al., 2000).

Select, que permite seleccionar, mover, ordenar y ajustar visualmente los elementos dentro del lienzo de modelación. Aunque parece una herramienta menor, su uso adecuado facilita la lectura del modelo. En un modelo con varios constructos e indicadores, la disposición visual influye en la comprensión de la estructura. Para **CAITIZEN**, conviene colocar los constructos predictores —**CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP**— alrededor o a la izquierda de la variable dependiente **CAITIZEN**, de modo que las rutas estructurales sean claras. El control Select ayuda a evitar cruces innecesarios de flechas, acumulación visual de indicadores o desorden en la presentación gráfica. Esta organización no cambia los resultados estadísticos, pero sí mejora la comunicación del modelo. En un libro aplicado, la claridad visual es importante porque el lector aprende tanto de la explicación textual como de la representación gráfica. Un modelo visualmente desordenado puede dificultar la identificación de constructos, indicadores y relaciones. Por ello, Select no debe entenderse solo como herramienta de edición estética, sino como un control que contribuye a la comunicación metodológica. La representación clara del modelo permite distinguir el modelo de medición del modelo estructural, lo cual es esencial para interpretar correctamente los resultados posteriores (Hair et al., 2021).

Connect, que permite dibujar relaciones estructurales entre constructos latentes. En el modelo **CAITIZEN**, este control se utilizará para trazar las rutas **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN**. Cada flecha representa una hipótesis estructural y, por tanto, no debe dibujarse de manera arbitraria. La dirección de la flecha expresa una relación esperada entre un constructo predictor y un constructo dependiente. En este caso, se plantea que las capacidades críticas, éticas, algorítmicas, colaborativas y metacognitivas pueden explicar la ciudadanía sostenible asistida por IA. Por ello, Connect no es simplemente una herramienta gráfica para unir nodos; es el control que traduce la lógica teórica del modelo en relaciones estimables. Una flecha mal orientada puede cambiar completamente la interpretación del modelo. Por ejemplo, no es lo mismo plantear **CAIL** → **CAITIZEN** que **CAITIZEN** → **CAIL**. La primera ruta supone que la alfabetización crítica contribuye a explicar la ciudadanía asistida por IA; la segunda supondría lo contrario. Por tanto, al usar **Connect**, el investigador debe verificar dirección, justificación y coherencia de cada relación. La documentación de **SmartPLS4** presenta la construcción del modelo de rutas como una etapa central del modelado visual (SmartPLS GmbH, s. f.-e).

Construct properties, que permite revisar o modificar las propiedades de cada constructo. Esta opción resulta importante porque ayuda a verificar el nombre del constructo, su tipo de medición, su presentación visual y, en algunos casos, su configuración dentro del modelo. Para el caso **CAITIZEN**, esta revisión es

Juan Mejía Trejo

indispensable porque todos los constructos serán tratados como **reflectivos**. Esto significa que los indicadores se entienden como manifestaciones del constructo latente. Por ejemplo, los indicadores de **EAR** reflejan la responsabilidad ética ante la IA; los indicadores de **HAIC** reflejan la colaboración humano-IA; y los indicadores de **MTPP** reflejan la transparencia metacognitiva en prácticas de prompting. Si un constructo se configurara por error como formativo, los criterios de evaluación cambiarían y la interpretación metodológica sería incorrecta. Por ello, revisar las propiedades del constructo no debe considerarse un paso opcional. En **PLS-SEM**, la especificación del tipo de medición afecta la forma en que se interpretan cargas, confiabilidad, **AVE**, **VIF** y validez. La diferencia entre medición reflectiva y formativa debe decidirse conceptualmente antes de estimar el modelo, no después de observar los resultados. Por tanto, **Construct** properties permite verificar que la representación operativa en **SmartPLS4** sea congruente con la definición teórica de cada constructo (Hair et al., 2021).

La **alineación y organización de indicadores** alrededor de los constructos. Aunque **SmartPLS4** puede colocar los indicadores automáticamente, el investigador debe revisar que la distribución visual sea clara y que cada indicador esté vinculado al constructo correcto. En el modelo **CAITIZEN**, cada constructo cuenta con diez indicadores, por lo que una mala organización puede generar saturación visual. Alinear indicadores permite identificar rápidamente qué ítems pertenecen a cada bloque de medición. Esta claridad será útil cuando se revisen resultados como cargas externas o indicadores con valores bajos. Si el modelo está bien organizado, el lector podrá localizar con facilidad los indicadores problemáticos, revisar su constructo correspondiente y comprender las decisiones metodológicas posteriores. Además, una visualización ordenada favorece la elaboración de figuras académicas para el libro o para un artículo. La claridad gráfica no reemplaza la validez estadística, pero sí mejora la transparencia del análisis. En investigaciones **SEM**, la representación del modelo debe permitir que otros investigadores comprendan la relación entre variables observables, constructos y rutas estructurales. Por ello, organizar indicadores no es solo una acción estética; es parte de la comunicación científica del modelo. Un diagrama claro facilita la revisión, la enseñanza y la reproducción del análisis (Gefen et al., 2000).

El **guardado y gestión de versiones del modelo**. Una vez creados los constructos, asignados los indicadores y trazadas las rutas estructurales, el investigador debe guardar el modelo dentro del proyecto correspondiente. En análisis aplicados, es frecuente realizar ajustes, revisar configuraciones o generar versiones alternativas. Por ello, conviene trabajar con nombres claros y conservar la trazabilidad de las modificaciones. En el caso **CAITIZEN**, puede existir una versión inicial del modelo completo, una versión después de revisar indicadores y otra versión final para reporte. Sin embargo, en el Capítulo 1 solo debe explicarse la importancia de guardar correctamente y evitar confusiones entre versiones. **SmartPLS4** permite trabajar con modelos dentro del proyecto, pero la responsabilidad de nombrarlos de manera clara corresponde al investigador. Una mala gestión de versiones puede provocar que se interpreten resultados de un modelo distinto al que se pretendía reportar. Desde la

lógica **PLS-SEM**, la reproducibilidad exige saber qué modelo fue estimado, qué indicadores fueron incluidos y qué rutas estructurales se evaluaron. Por ello, la gestión de versiones debe formar parte de las buenas prácticas de construcción del modelo. La ruta aplicada de **PLS-SEM** requiere orden no solo conceptual y estadístico, sino también documental y operativo (Hair et al., 2019).

La **verificación final del modelo antes de ejecutar el algoritmo**. Antes de calcular resultados, el investigador debe revisar que todos los constructos estén presentes, que todos los indicadores hayan sido asignados correctamente, que no existan indicadores sueltos, que las flechas estructurales estén orientadas en la dirección correcta y que el tipo de medición sea congruente con la definición conceptual del modelo. En **CAITIZEN**, esta verificación implica confirmar que **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** estén conectados hacia **CAITIZEN**, que cada constructo tenga sus diez indicadores correspondientes y que la configuración sea reflectiva. Esta revisión previa evita errores que podrían detectarse tarde, después de generar reportes o interpretar resultados. También ayuda a distinguir entre errores de construcción y problemas estadísticos reales. Por ejemplo, una carga externa inesperada podría deberse a un indicador asignado al constructo equivocado, no necesariamente a debilidad conceptual del ítem. Por ello, antes de ejecutar el algoritmo **PLS-SEM**, debe realizarse una inspección visual y metodológica del modelo. El control operativo no termina al dibujar el diagrama; incluye verificar que el modelo construido corresponda exactamente al modelo teórico previsto. En el enfoque PLS, la especificación correcta del modelo es condición previa para cualquier estimación e interpretación posterior (Chin, 1998).

Los **controles básicos para construir el modelo** en **SmartPLS4** permiten convertir una propuesta conceptual en un modelo **PLS-SEM** visual y estimable. **Create model** inicia el espacio de construcción; **Latent variable** permite crear constructos; el **panel de indicadores** permite asignar variables observables; **Select** ayuda a organizar el modelo; **Connect** permite establecer relaciones estructurales; **Construct properties** permite revisar la especificación de cada constructo; la alineación de indicadores mejora la claridad visual; el guardado permite conservar versiones; y la verificación final asegura que el modelo esté listo para el algoritmo. En el caso **CAITIZEN**, estos controles permitirán representar la relación entre cinco predictores —**CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP**— y la variable dependiente **CAITIZEN**. No obstante, su uso debe comprenderse dentro de una lógica metodológica. **SmartPLS4** facilita la construcción visual del modelo, pero no decide por el investigador qué constructos usar, qué indicadores asignar o qué rutas justificar. Por ello, el aprendizaje de estos controles debe orientarse a desarrollar criterio analítico. Cada acción en el software debe responder a una decisión conceptual previa y preparar una evaluación estadística posterior. De esta manera, el lector podrá usar **SmartPLS4** no como una herramienta mecánica, sino como una plataforma para expresar, estimar y evaluar modelos estructurales aplicados (SmartPLS GmbH, s. f.-a).

Controles básicos para estimar, revisar y reportar resultados

Una vez construido el modelo en **SmartPLS4**, el investigador debe pasar de la representación gráfica a la **estimación, revisión y reporte** de resultados. Esta transición es fundamental porque el modelo dibujado en el lienzo todavía no demuestra calidad metodológica ni apoyo empírico. En el caso del modelo **CAITIZEN**, la construcción gráfica permite visualizar los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, así como las rutas estructurales propuestas hacia la ciudadanía sostenible asistida por inteligencia artificial. Sin embargo, para saber si el modelo funciona adecuadamente, es necesario ejecutar procedimientos de cálculo y revisar criterios de calidad. En **SmartPLS4**, estos procedimientos se concentran en controles como **Calculate**, **PLS-SEM Algorithm**, **Bootstrapping**, **Open full report**, **Quality criteria**, **Collinearity statistics (VIF)** y las secciones de resultados del modelo de medición y del modelo estructural. Estos controles no deben entenderse como botones aislados, sino como etapas de una ruta analítica. El algoritmo proporciona una primera estimación; el reporte organiza los resultados; los criterios de calidad permiten evaluar la medición; la revisión estructural permite interpretar relaciones; y el **Bootstrapping** permite contrastar hipótesis. Por ello, este apartado introduce los controles básicos que el lector utilizará en los capítulos posteriores. La finalidad no es agotar todavía la evaluación estadística, sino comprender la función de cada control dentro del proceso **PLS-SEM**. **SmartPLS4** facilita el cálculo, pero el investigador debe decidir qué revisar, en qué orden y con qué criterios metodológicos interpretar los resultados (SmartPLS GmbH, s. f.-a) y la siguiente lista de controles estratégicos, a tener en cuenta:

Calculate, que funciona como punto de entrada a los procedimientos de estimación disponibles en **SmartPLS4**. Desde esta opción, el investigador puede ejecutar el **PLS-SEM Algorithm**, aplicar **Bootstrapping** y, según la configuración del proyecto, acceder a otros procedimientos analíticos. En el modelo **CAITIZEN**, **Calculate** será utilizado inicialmente para ejecutar el algoritmo **PLS-SEM** y obtener los resultados preliminares del modelo. Es importante comprender que calcular no significa validar automáticamente. El cálculo solo genera estimaciones; la validación exige interpretación posterior. Al seleccionar **Calculate**, el investigador debe saber qué procedimiento requiere y qué pregunta metodológica desea responder. Si necesita obtener cargas externas, coeficientes de trayectoria, confiabilidad, **AVE**, **VIF** y **R²** preliminar, debe ejecutar el algoritmo **PLS-SEM**. Si necesita evaluar significancia estadística, **t-values**, **p-values** e intervalos de confianza, debe ejecutar **Bootstrapping**. Esta distinción evita confundir estimación inicial con contraste de hipótesis. En muchos análisis aplicados, un error frecuente consiste en interpretar directamente los coeficientes del algoritmo como si fueran pruebas de significancia. En realidad, el algoritmo ofrece la estimación del modelo, mientras que el **Bootstrapping** permite evaluar la estabilidad estadística de los resultados. Por ello, el control **Calculate** debe enseñarse como una puerta de acceso a distintos procedimientos, no como una operación única. La práctica adecuada de **PLS-SEM** exige reconocer

cuándo estimar, cuándo contrastar, cuándo diagnosticar y cuándo reportar (Hair et al., 2021).

El **PLS-SEM Algorithm** es el procedimiento que permite obtener la estimación inicial del modelo de rutas. En **SmartPLS4**, este algoritmo utiliza el modelo especificado y los datos de los indicadores para estimar puntuaciones de variables latentes, cargas externas, coeficientes estructurales y distintos criterios de evaluación. En el caso **CAITIZEN**, el algoritmo permitirá observar cómo se comportan los indicadores de **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITIZEN**, así como las rutas estructurales que conectan a los predictores con la variable dependiente. Esta etapa es indispensable porque proporciona una primera imagen del desempeño del modelo. Sin embargo, los resultados iniciales deben interpretarse con cautela. Una carga externa alta sugiere que un indicador representa bien al constructo, pero no basta por sí sola para afirmar validez total. Un coeficiente de trayectoria positivo indica dirección de relación, pero no significa necesariamente que la hipótesis esté apoyada estadísticamente. Un valor R^2 permite conocer la varianza explicada del constructo dependiente, pero debe interpretarse junto con la teoría, el contexto y los tamaños de efecto. Por ello, el algoritmo **PLS-SEM** debe verse como el primer diagnóstico empírico del modelo. En el libro, su ejecución se desarrollará con detalle en el Capítulo 4. En este apartado introductorio basta con reconocer que el algoritmo transforma la construcción gráfica del modelo en resultados analíticos que deberán ser evaluados ordenadamente (SmartPLS GmbH, s. f.-c).

Después de ejecutar el algoritmo, el investigador debe revisar primero el **modelo de medición**. Esta revisión es prioritaria porque las relaciones estructurales solo pueden interpretarse con rigor si los constructos están medidos adecuadamente. En un modelo reflectivo como **CAITIZEN**, la evaluación del modelo de medición incluye criterios como **cargas externas**, **confiabilidad interna**, **rho_A**, **confiabilidad compuesta** y **validez convergente** mediante **AVE**. Las cargas externas permiten identificar qué tan bien cada indicador representa a su constructo. La confiabilidad interna permite revisar la consistencia del conjunto de indicadores. La validez convergente indica si los indicadores comparten suficiente varianza para representar una misma variable latente. En el caso de **CAIL**, por ejemplo, se esperaría que los indicadores vinculados con alfabetización crítica en **IA** carguen adecuadamente sobre el constructo. Lo mismo aplica para **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITIZEN**. Esta etapa evita interpretar hipótesis estructurales sobre constructos débiles o mal medidos. Por tanto, los controles de revisión de resultados deben usarse siguiendo una secuencia lógica: primero se revisa la medición, después la estructura y finalmente la significancia. Un investigador que salta directamente a los **p-values** sin revisar cargas, confiabilidad y **AVE** corre el riesgo de reportar relaciones estadísticamente llamativas pero metodológicamente frágiles. En **PLS-SEM**, reportar bien implica demostrar primero que los constructos poseen calidad de medición suficiente (Hair et al., 2019).

Dentro de los criterios de medición, la **validez convergente** ocupa un lugar central. Esta se revisa comúnmente mediante el **Average Variance Extracted** o **AVE**, que indica la proporción de varianza que un constructo explica en sus indicadores. En

términos generales, un **AVE** adecuado sugiere que los indicadores comparten suficiente varianza común y, por tanto, representan de manera razonable al constructo latente. En el modelo **CAITIZEN**, esta revisión permitirá observar si cada bloque de indicadores converge sobre su constructo correspondiente. Por ejemplo, los indicadores de **HAIC** deben converger en torno a la colaboración creativo-funcional entre personas e inteligencia artificial; los indicadores de **MTPP** deben converger en torno a prácticas metacognitivas y transparentes en el uso de prompts. Si el **AVE** de un constructo fuera bajo, el investigador tendría que revisar indicadores, cargas externas y coherencia conceptual. No obstante, la decisión de eliminar indicadores no debe ser automática. Un indicador puede tener valor conceptual aunque su carga no sea ideal; eliminarlo sin análisis puede empobrecer el contenido del constructo. Por ello, la revisión de **AVE** debe acompañarse de criterio teórico. En un libro aplicado, es necesario enseñar que los umbrales son guías, no sustitutos del juicio metodológico. La validez convergente permite diagnosticar si el constructo se sostiene empíricamente, pero su interpretación debe estar articulada con la definición conceptual, la escala utilizada y el propósito del modelo (Fornell & Larcker, 1981).

Además de la validez convergente, el investigador debe revisar la **validez discriminante**, que permite determinar si los constructos son empíricamente distintos entre sí. En el modelo **CAITIZEN**, esta revisión resulta especialmente importante porque algunos constructos pueden estar conceptualmente relacionados. Por ejemplo, **EAR** y **AFDJ** comparten una preocupación ética, pero no deben medir exactamente lo mismo. **CAIL** y **MTPP** pueden relacionarse con capacidades críticas ante la **IA**, pero uno se orienta a alfabetización crítica y otro a transparencia metacognitiva en prácticas de prompting. La validez discriminante ayuda a verificar que cada constructo conserve identidad propia dentro del modelo. **SmartPLS4** permite revisar criterios como **Fornell-Larcker**, **Cross Loadings** y **HTMT**. Entre ellos, **HTMT** se ha convertido en un criterio especialmente relevante para detectar problemas de discriminación entre constructos en modelos basados en varianza. Si dos constructos presentan valores excesivamente altos de **HTMT**, el investigador debe revisar si existe solapamiento conceptual, indicadores ambiguos o una especificación incorrecta del modelo. Esta revisión no debe interpretarse como simple requisito estadístico, sino como una prueba de claridad conceptual. Si los constructos no se distinguen empíricamente, las relaciones estructurales entre ellos pueden perder sentido. Por ello, antes de afirmar que **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** explican **CAITIZEN**, debe comprobarse que cada uno representa una dimensión diferenciada del fenómeno de ciudadanía asistida por **IA** (Henseler et al., 2015).

Otro control básico en la revisión de resultados es **Collinearity statistics (VIF)**, que permite evaluar posibles problemas de colinealidad. En **PLS-SEM**, la colinealidad puede afectar la interpretación de relaciones estructurales, especialmente cuando varios predictores explican un mismo constructo dependiente. En el modelo **CAITIZEN**, los predictores **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** se dirigen hacia la variable dependiente **CAITIZEN**. Por ello, es necesario revisar si estos constructos presentan niveles excesivos de colinealidad entre sí. Si dos o más predictores son demasiado similares, los coeficientes de trayectoria pueden volverse inestables o difíciles de

interpretar. **SmartPLS4** permite distinguir entre **VIF** del modelo externo y **VIF** del modelo interno. En un modelo reflectivo como **CAITIZEN**, el **Inner model VIF** resulta especialmente relevante para evaluar colinealidad entre constructos predictores. No obstante, el **VIF** debe interpretarse con prudencia. Un valor elevado puede advertir solapamiento conceptual o redundancia entre predictores; un valor aceptable permite continuar con la interpretación estructural. Esta revisión es importante porque **PLS-SEM** no exige que los predictores sean completamente independientes, pero sí requiere ausencia de colinealidad crítica. En el libro, el análisis **VIF** será tratado con detalle al revisar el algoritmo y el modelo estructural. Desde este apartado, debe quedar claro que **VIF** es un control de diagnóstico, no una prueba de hipótesis (Kock, 2015).

Una vez revisada la medición y la colinealidad, el investigador puede avanzar hacia los resultados del **modelo estructural**. En esta etapa se analizan los **coeficientes de trayectoria**, el R^2 , el f^2 y otros indicadores de capacidad explicativa o predictiva. Los coeficientes de trayectoria muestran la dirección y magnitud de las relaciones entre constructos. En **CAITIZEN**, permitirán observar qué tan fuerte es la relación de **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** con la variable dependiente. El R^2 indica la proporción de varianza explicada en **CAITIZEN** por sus predictores. El f^2 permite valorar el tamaño del efecto de cada predictor dentro del modelo. Estos resultados deben interpretarse en conjunto, no de manera aislada. Una ruta puede ser significativa, pero tener un efecto pequeño; otra puede tener un efecto moderado, pero requerir discusión teórica. Del mismo modo, un R^2 alto puede sugerir capacidad explicativa, pero debe evaluarse considerando la naturaleza del fenómeno y el contexto de investigación. En ciencias de la administración, la interpretación debe articular estadística, teoría y relevancia práctica. El modelo **CAITIZEN** no busca únicamente saber si existen relaciones numéricas, sino identificar qué dimensiones contribuyen de manera relevante a explicar la ciudadanía sostenible asistida por **IA**. Por ello, los controles estructurales deben usarse para pasar de los resultados al significado académico del modelo (Cohen, 1992).

El control **Bootstrapping** cumple una función diferente a la del algoritmo **PLS-SEM**. Mientras el algoritmo proporciona la estimación inicial del modelo, **Bootstrapping** permite evaluar la **significancia estadística** de los resultados mediante un procedimiento de remuestreo. En **SmartPLS4**, este procedimiento genera **t-values**, **p-values** e **intervalos de confianza**, que permiten contrastar hipótesis. En el modelo **CAITIZEN**, **Bootstrapping** será utilizado para evaluar si las rutas **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN** cuentan con apoyo estadístico. Es fundamental que el lector comprenda que una relación estimada por el algoritmo no equivale automáticamente a una hipótesis apoyada. Para afirmar que una hipótesis tiene respaldo empírico, se requiere revisar la significancia, la dirección del efecto, la magnitud y el intervalo de confianza. Además, la decisión no debe basarse únicamente en el valor **p**. Un resultado estadísticamente significativo debe interpretarse junto con el tamaño del efecto y su coherencia conceptual. **Bootstrapping** es especialmente útil en **PLS-SEM** porque no depende de supuestos estrictos de normalidad y permite evaluar la estabilidad de los coeficientes

estimados. Por ello, este control será una herramienta central para el capítulo dedicado a la contrastación de hipótesis. En este apartado, se presenta solo como parte de la ruta básica de estimación y revisión de resultados (SmartPLS GmbH, s. f.-b).

El control **Open full report** permite acceder a una visión organizada de los resultados generados por **SmartPLS4**. Este reporte completo es indispensable porque concentra información del modelo de medición, del modelo estructural y, cuando corresponde, del **Bootstrapping**. En el caso **CAITIZEN**, el reporte permitirá revisar cargas externas, confiabilidad, **AVE**, **HTMT**, **VIF**, coeficientes de trayectoria, **R^2** , **f^2** , **t -values**, **p -values** e intervalos de confianza. Sin embargo, el reporte no debe leerse como una colección de tablas independientes. Su interpretación exige una ruta ordenada. Primero se revisan los indicadores y constructos; después, la validez discriminante; posteriormente, la colinealidad y los resultados estructurales; finalmente, la significancia de las hipótesis. Esta secuencia ayuda a evitar interpretaciones erróneas. Por ejemplo, no se debe comenzar reportando que una hipótesis fue apoyada si aún no se ha demostrado que el constructo dependiente y los predictores fueron medidos con calidad suficiente. Tampoco debe omitirse el diagnóstico de colinealidad si varios constructos predicen a **CAITIZEN**. En un libro aplicado, el reporte completo debe enseñarse como instrumento de lectura crítica. El investigador debe saber localizar resultados, seleccionar los criterios relevantes y convertir tablas en interpretación académica. La calidad del reporte final dependerá de que el usuario no copie mecánicamente salidas del software, sino que las organice conforme a los estándares de reporte **PLS-SEM** (Hair et al., 2019).

El apartado **Quality criteria** agrupa resultados clave para evaluar la calidad del modelo. En **SmartPLS4**, esta sección permite consultar distintos indicadores asociados con medición, validez, colinealidad y estructura. Para el modelo **CAITIZEN**, **Quality criteria** será una de las secciones más importantes porque permitirá revisar si los constructos cumplen con los requisitos mínimos antes de interpretar hipótesis. En modelos reflectivos, el lector deberá prestar atención a cargas externas, confiabilidad, **AVE** y validez discriminante. Después, deberá revisar colinealidad y resultados estructurales. Esta sección ayuda a ordenar el análisis porque concentra criterios que, de otra manera, podrían revisarse de forma dispersa. No obstante, los criterios de calidad no deben interpretarse como una lista mecánica de “*cumple*” o “*no cumple*”. Cada indicador requiere lectura metodológica. Una confiabilidad muy baja puede revelar problemas de consistencia; una confiabilidad excesivamente alta puede sugerir redundancia de indicadores; un **AVE** insuficiente puede indicar falta de convergencia; un **HTMT** elevado puede advertir falta de distinción entre constructos; un **VIF** alto puede sugerir colinealidad. Por tanto, **Quality criteria** es una sección de diagnóstico. En el caso **CAITIZEN**, estos criterios permitirán determinar si el modelo puede avanzar desde la evaluación de medición hacia la interpretación estructural. La función del investigador es transformar estos diagnósticos en decisiones: conservar, revisar, justificar, ajustar o reportar limitaciones del modelo (Hair et al., 2021).

Otro aspecto importante de la revisión de resultados es distinguir entre **explicación**, **predicción** y **significancia**. En **PLS-SEM**, el modelo puede explicar una proporción

Juan Mejía Trejo

de varianza en el constructo dependiente, puede mostrar capacidad predictiva y puede presentar rutas significativas. Estos tres elementos están relacionados, pero no son equivalentes. El R^2 ayuda a comprender la capacidad explicativa del modelo sobre **CAITIZEN**; los análisis predictivos permiten valorar si el modelo tiene utilidad para anticipar valores fuera de la muestra o en escenarios aplicados; y el **Bootstrapping** permite determinar si las relaciones estimadas son estadísticamente significativas. Un error frecuente consiste en reducir toda la interpretación al valor p . Sin embargo, un modelo **PLS-SEM** debe valorar también magnitud, dirección, relevancia sustantiva y capacidad explicativa. En el modelo **CAITIZEN**, puede ocurrir que un predictor sea significativo, pero tenga un efecto pequeño; también puede ocurrir que un predictor no sea significativo, pero tenga valor teórico para discutir el fenómeno. Por ello, el reporte debe ser equilibrado y no exagerar resultados. La interpretación debe responder preguntas como: qué constructos explican mejor **CAITIZEN**, qué efectos son relevantes, qué hipótesis reciben apoyo, qué hallazgos requieren cautela y qué limitaciones metodológicas deben reconocerse. Esta distinción es fundamental para que el análisis **PLS-SEM** no se convierta en una lectura automática de umbrales (Shmueli, 2010).

Los controles de estimación y revisión también deben conducir al investigador hacia un **reporte académico defendible**. En un libro aplicado, no basta con mostrar capturas de **SmartPLS4**; es necesario enseñar cómo convertir resultados en narrativa metodológica. Un reporte adecuado debe explicar qué algoritmo se usó, cómo se evaluó el modelo de medición, qué criterios se aplicaron para confiabilidad y validez, cómo se revisó la colinealidad, qué resultados estructurales se obtuvieron y cómo se contrastaron las hipótesis. En el modelo **CAITIZEN**, esto implicará reportar tablas de cargas externas, confiabilidad, **AVE**, **HTMT**, **VIF**, coeficientes de trayectoria, R^2 , f^2 y **Bootstrapping**. Cada tabla debe acompañarse de interpretación. Por ejemplo, no basta decir que el **AVE** supera 0.50; debe explicarse que ello respalda la validez convergente del constructo. No basta indicar que una ruta tiene $p < 0.05$; debe comentarse su dirección, magnitud y relevancia para la ciudadanía asistida por **IA**. Tampoco basta reportar que no hay colinealidad crítica; debe explicarse por qué eso permite interpretar los predictores con mayor confianza. En este sentido, los controles de **SmartPLS4** son medios para producir evidencia, pero el reporte académico exige criterio, síntesis y argumentación. El investigador debe presentar los resultados de manera transparente, ordenada y proporcional a la evidencia obtenida (Hair et al., 2019).

En síntesis, los **controles básicos para estimar, revisar y reportar resultados** en **SmartPLS4** permiten avanzar desde el modelo gráfico hacia la evidencia empírica. **Calculate** abre los procedimientos de estimación; **PLS-SEM Algorithm** genera la primera estimación; **Quality criteria** organiza diagnósticos de medición y estructura; **Collinearity statistics (VIF)** permite revisar posibles problemas de colinealidad; **Bootstrapping** evalúa significancia estadística; y **Open full report** concentra los resultados necesarios para la interpretación y el reporte. En el modelo **CAITIZEN**, estos controles permitirán evaluar si los indicadores representan adecuadamente a sus constructos, si los constructos son confiables y válidos, si los predictores presentan

colinealidad aceptable, si las rutas estructurales tienen magnitud relevante y si las hipótesis cuentan con apoyo estadístico. El lector debe comprender que la estimación no es el final del análisis, sino el inicio de una revisión metodológica ordenada. **SmartPLS4** ofrece los resultados, pero no decide por sí solo qué significan. La interpretación depende del investigador, de la teoría, de los criterios de calidad y del propósito del estudio. Por ello, este apartado prepara al lector para los capítulos posteriores, donde cada control será aplicado con mayor detalle al caso **CAITIZEN**. La meta no es aprender a oprimir botones, sino entender cómo cada control contribuye a construir evidencia metodológicamente válida y académicamente reportable (SmartPLS GmbH, s. f.-a). Ver **Tabla 1.1**

Tabla 1.1. Cien controles, vistas, funciones y salidas principales de SmartPLS4 para PLS-SEM

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
1	Adjusted R-square / R^2 ajustado	Coeficiente de determinación ajustado; corrige el R^2 según número de predictores.	Open Full Report → Quality Criteria → R-square.	SmartPLS GmbH (s. f.-e)
2	Algorithms and Techniques	Sección general de técnicas disponibles: PLS-SEM , Bootstrapping , PLSpredict, IPMA, NCA, MGA, etc.	Documentation → Algorithms and Techniques.	SmartPLS GmbH (s. f.-a)
3	AVE	Average Variance Extracted ; varianza media extraída para evaluar validez convergente.	Open Full Report → Quality Criteria → Construct Reliability and Validity → AVE.	SmartPLS GmbH (s. f.-e)
4	Blindfolding	Procedimiento de omisión sistemática usado tradicionalmente para revisar relevancia predictiva Q^2 .	Calculate → Blindfolding, si está disponible según modelo.	SmartPLS GmbH (s. f.-a)
5	Bootstrapping	Remuestreo no paramétrico para evaluar significancia estadística.	Calculate → Bootstrapping .	SmartPLS GmbH (s. f.-b)
6	Bootstrapping Confidence Intervals	Intervalos de confianza generados por Bootstrapping .	Bootstrapping Report → Confidence Intervals.	SmartPLS GmbH (s. f.-b)
7	Bootstrapping p Values	p-values para evaluar apoyo estadístico de relaciones o criterios.	Bootstrapping Report → p Values.	SmartPLS GmbH (s. f.-b)
8	Bootstrapping t Values	t-values generados por remuestreo para evaluar significancia.	Bootstrapping Report → t Values.	SmartPLS GmbH (s. f.-b)
9	Calculate	Menú principal para ejecutar algoritmos y procedimientos de análisis.	Modeling View → Calculate.	SmartPLS GmbH (s. f.-a)
10	Case Processing / Cases	Revisión del número de casos incluidos en el análisis.	Data View o Report View → Data / Descriptives.	SmartPLS GmbH (s. f.-d)

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
11	CB-SEM	<i>Covariance-Based Structural Equation Modeling</i> ; modelado SEM basado en covarianzas.	Create Model o Calculate → CB-SEM, si aplica.	SmartPLS GmbH (s. f.-a)
12	CB-SEM Bootstrapping	Bootstrapping aplicado a modelos CB-SEM.	Calculate → CB-SEM Bootstrapping .	SmartPLS GmbH (s. f.-a)
13	CCA	<i>Confirmatory Composite Analysis</i> ; análisis confirmatorio de compuestos.	Calculate → Confirmatory Composite Analysis.	SmartPLS GmbH (s. f.-a)
14	Collinearity Statistics (VIF)	Estadísticos de colinealidad; VIF significa <i>Variance Inflation Factor</i> .	Open Full Report → Quality Criteria → Collinearity Statistics (VIF).	SmartPLS GmbH (s. f.-e)
15	Confidence Intervals Bias-Corrected	Intervalos de confianza corregidos por sesgo, si se seleccionan en Bootstrapping .	Bootstrapping Settings → Confidence Interval Method.	SmartPLS GmbH (s. f.-b)
16	Confidence Intervals Percentile	Intervalos de confianza por percentiles.	Bootstrapping Settings → Confidence Interval Method.	SmartPLS GmbH (s. f.-b)
17	Connect	Dibuja relaciones estructurales entre constructos.	Modeling View → Connect → origen → destino.	SmartPLS GmbH (s. f.-e)
18	Consistent Bootstrapping	Bootstrapping para Consistent PLS-SEM , también llamado PLSc.	Calculate → Consistent Bootstrapping .	SmartPLS GmbH (s. f.-a)
19	Consistent PLS-SEM / PLSc	Estimación PLS consistente para ciertos modelos de factores comunes.	Calculate → Consistent PLS-SEM .	SmartPLS GmbH (s. f.-a)
20	Construct Properties	Revisión o edición de nombre, tipo de medición y configuración del constructo.	Clic derecho sobre el constructo → Properties.	SmartPLS GmbH (s. f.-e)
21	Construct Reliability and Validity	Sección de confiabilidad y validez convergente: Alpha, rho_A , rho_c y AVE.	Open Full Report → Quality Criteria → Construct Reliability and Validity.	SmartPLS GmbH (s. f.-e)
22	Copy Results	Copia tablas o resultados para usarlos en reportes externos.	Report View → seleccionar tabla → Copy.	SmartPLS GmbH (s. f.-e)
23	Create Model / New Model	Inicia un nuevo modelo dentro del proyecto activo.	Project View → Create Model / PLS-SEM .	SmartPLS GmbH (s. f.-e)
24	Create Project / New Project	Crea un nuevo proyecto dentro del Workspace .	Project View → New Project.	SmartPLS GmbH (s. f.-d)
25	Cronbach Alpha	Coefficiente tradicional de consistencia interna de los indicadores.	Quality Criteria → Construct Reliability and Validity → Cronbach Alpha.	SmartPLS GmbH (s. f.-e)

Juan Mejía Trejo

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
26	Cross Loadings	Cargas cruzadas para revisar si cada indicador carga más en su constructo.	Quality Criteria → Discriminant Validity → Cross Loadings.	SmartPLS GmbH (s. f.-e)
27	CVPAT	<i>Cross-Validated Predictive Ability Test</i> ; prueba de capacidad predictiva validada cruzadamente.	PLSpredict Results → CVPAT, si se ejecuta PLSpredict.	SmartPLS GmbH (s. f.-k)
28	Data Groups	Grupos de datos para análisis por submuestras o comparaciones.	Data View / Project View → Data Groups, según versión.	SmartPLS GmbH (s. f.-d)
29	Data View	Vista para revisar variables, casos y características del archivo importado.	Doble clic sobre archivo de datos.	SmartPLS GmbH (s. f.-d)
30	Descriptive Statistics	Estadísticos descriptivos de indicadores: media, desviación estándar, etc.	Data View → Descriptives.	SmartPLS GmbH (s. f.-d)
31	Discriminant Validity	Validez discriminante; revisa si los constructos son empíricamente distintos.	Quality Criteria → Discriminant Validity.	SmartPLS GmbH (s. f.-e)
32	Drag and Drop Indicators	Asignación de indicadores a constructos mediante arrastre.	Modeling View → arrastrar indicadores al constructo.	SmartPLS GmbH (s. f.-e)
33	Duplicate Model	Duplica un modelo para conservar versiones alternativas.	Project View → clic derecho sobre modelo → Duplicate.	SmartPLS GmbH (s. f.-d)
34	Edit Data Set	Revisión o edición limitada del conjunto de datos dentro del proyecto.	Data View → herramientas de datos.	SmartPLS GmbH (s. f.-d)
35	Excess Kurtosis	Curtosis excedente de las variables; ayuda a revisar forma de distribución.	Data View → Descriptives → Excess Kurtosis.	SmartPLS GmbH (s. f.-p)
36	Export Data	Exporta datos o información de variables del proyecto.	Data View → Export.	SmartPLS GmbH (s. f.-d)
37	Export Report	Exporta reportes a formatos disponibles, por ejemplo Excel o HTML.	Report View → Export.	SmartPLS GmbH (s. f.-e)
38	f-square / f^2	Tamaño del efecto; mide contribución de un predictor sobre un constructo endógeno.	Quality Criteria → f-square.	SmartPLS GmbH (s. f.-e)
39	Final Results	Sección general de resultados finales del algoritmo ejecutado.	Open Full Report → Final Results.	SmartPLS GmbH (s. f.-e)
40	Fornell-Larcker Criterion	Criterio de validez discriminante basado en AVE y correlaciones.	Quality Criteria → Discriminant Validity → Fornell-Larcker Criterion.	SmartPLS GmbH (s. f.-e)
41	Generate Data Groups	Crea grupos para análisis comparativo o multigrupo.	Data View → Data Groups / Generate Groups.	SmartPLS GmbH (s. f.-d)

Juan Mejía Trejo

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
42	Goodness of Fit / Model Fit	Criterios de ajuste del modelo, especialmente relevantes en CB-SEM y análisis complementarios.	Report View → Model Fit.	SmartPLS GmbH (s. f.-o)
43	Graphical Results	Resultados visuales sobre el modelo, como coeficientes y cargas en el diagrama.	Después del cálculo → Results on Model.	SmartPLS GmbH (s. f.-e)
44	Heatmap / Result Colors	Colores o resaltados de resultados según umbrales.	Report View → resultados coloreados.	SmartPLS GmbH (s. f.-f)
45	HTMT	Heterotrait-Monotrait Ratio; razón heterorrasgo-monorrasgo para validez discriminante.	Quality Criteria → Discriminant Validity → HTMT.	SmartPLS GmbH (s. f.-a)
46	HTMT Confidence Intervals	Intervalos de confianza para HTMT mediante Bootstrapping .	Bootstrapping Report → HTMT Confidence Intervals.	SmartPLS GmbH (s. f.-b)
47	Import Data File	Importa archivo de datos al proyecto.	Project View → Import Data File.	SmartPLS GmbH (s. f.-d)
48	Import Project	Importa un proyecto previamente guardado o compartido.	Workspace / Project View → Import Project.	SmartPLS GmbH (s. f.-d)
49	Indicator Correlations	Correlaciones entre indicadores, útiles para revisión preliminar.	Data View o Report View → Correlations, si está disponible.	SmartPLS GmbH (s. f.-d)
50	Indicators / Variables List	Lista de variables observables disponibles para modelar.	Modeling View → panel izquierdo de indicadores.	SmartPLS GmbH (s. f.-e)
51	Inner Model VIF	VIF del modelo interno; revisa colinealidad entre constructos predictores.	Quality Criteria → Collinearity Statistics → Inner Model.	SmartPLS GmbH (s. f.-e)
52	IPMA	<i>Importance-Performance Map Analysis</i> ; análisis importancia-desempeño.	Calculate → IPMA.	SmartPLS GmbH (s. f.-h)
53	IPMA Construct Performance	Desempeño de constructos en IPMA .	IPMA Report → Construct Performance.	SmartPLS GmbH (s. f.-h)
54	IPMA Indicator Performance	Desempeño de indicadores en IPMA .	IPMA Report → Indicator Performance.	SmartPLS GmbH (s. f.-h)
55	IPMA Target Construct	Selección del constructo objetivo para análisis IPMA .	Calculate → IPMA → Target Construct.	SmartPLS GmbH (s. f.-h)
56	Latent Variable / Construct	Crea variables latentes o constructos en el modelo.	Modeling View → crear constructo o arrastrar indicadores.	SmartPLS GmbH (s. f.-e)
57	Linear Regression	Modelo de regresión lineal disponible en SmartPLS .	Create Model / Calculate → Regression.	SmartPLS GmbH (s. f.-n)
58	MGA	<i>Multigroup Analysis</i> ; análisis multigrupo para comparar parámetros entre grupos.	Calculate → Multigroup Analysis.	SmartPLS GmbH (s. f.-j)

Juan Mejía Trejo

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
59	MICOM	<i>Measurement Invariance of Composites</i> ; prueba de invariancia de medición de compuestos.	Calculate → MICOM / Permutation, según flujo de MGA.	SmartPLS GmbH (s. f.-a)
60	Missing Values	Revisión de valores perdidos en el conjunto de datos.	Data View → Missing Values / Descriptives.	SmartPLS GmbH (s. f.-d)
61	Model Fit	Evalúa qué tan bien el modelo reproduce la matriz observada, según técnica usada.	Report View → Model Fit.	SmartPLS GmbH (s. f.-o)
62	Model Name	Campo para nombrar el modelo.	New Model → Model Name → Save.	SmartPLS GmbH (s. f.-e)
63	Model Type: PLS-SEM	Selección del tipo de modelo PLS-SEM .	New Model → Model Type → PLS-SEM .	SmartPLS GmbH (s. f.-e)
64	Modeling Canvas / Área gráfica	Lienzo donde se colocan constructos, indicadores y relaciones.	Modeling View.	SmartPLS GmbH (s. f.-e)
65	Modeling View / Model Editor	Vista gráfica de construcción y edición del modelo.	Se abre después de crear el modelo.	SmartPLS GmbH (s. f.-e)
66	Moderation	Configuración y análisis de efectos moderadores.	Modeling View → Create Moderation / Calculate → Moderation outputs.	SmartPLS GmbH (s. f.-l)
67	NCA	<i>Necessary Condition Analysis</i> ; análisis de condiciones necesarias.	Calculate → NCA.	SmartPLS GmbH (s. f.-i)
68	NCA Bottleneck Table	Tabla de cuellos de botella en NCA .	NCA Report → Bottleneck Table.	SmartPLS GmbH (s. f.-i)
69	NCA Ceiling Line	Línea de techo usada en análisis NCA .	NCA Report → Ceiling Line.	SmartPLS GmbH (s. f.-i)
70	NCA Effect Size	Tamaño del efecto en análisis de condición necesaria.	NCA Report → Effect Size.	SmartPLS GmbH (s. f.-i)
71	Outer Loadings	Cargas externas de indicadores reflectivos.	Quality Criteria → Outer Loadings.	SmartPLS GmbH (s. f.-e)
72	Outer Model VIF	VIF del modelo externo; revisa colinealidad entre indicadores.	Quality Criteria → Collinearity Statistics → Outer Model.	SmartPLS GmbH (s. f.-e)
73	Outer Weights	Pesos externos, especialmente relevantes para medición formativa.	Open Full Report → Outer Weights.	SmartPLS GmbH (s. f.-e)
74	Path Analysis	Análisis de rutas basado en regresión.	Create Model / Calculate → Path Analysis.	SmartPLS GmbH (s. f.-m)

Juan Mejía Tréjo

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
75	Path Coefficients	Coefficientes de trayectoria entre constructos.	Final Results → Path Coefficients.	SmartPLS GmbH (s. f.-c)
76	Permutation Test	Prueba de permutación para comparaciones de grupos o invariancia.	Calculate → Permutation.	SmartPLS GmbH (s. f.-a)
77	PLS-MGA	Análisis multigrupo PLS basado en Bootstrapping .	Calculate → Multigroup Analysis → PLS-MGA.	SmartPLS GmbH (s. f.-j)
78	PLS-SEM Algorithm	Algoritmo principal de estimación PLS-SEM .	Calculate → PLS-SEM Algorithm.	SmartPLS GmbH (s. f.-c)
79	PLSpredict	Evaluación predictiva mediante validación cruzada k-fold.	Calculate → PLSpredict.	SmartPLS GmbH (s. f.-g)
80	PLSpredict k-folds	Número de particiones para validación cruzada.	PLSpredict Settings → Number of Folds.	SmartPLS GmbH (s. f.-g)
81	PLSpredict MAE	<i>Mean Absolute Error</i> ; error absoluto medio en predicción.	PLSpredict Report → Prediction Summary.	SmartPLS GmbH (s. f.-g)
82	PLSpredict MAPE	<i>Mean Absolute Percentage Error</i> ; error porcentual absoluto medio.	PLSpredict Report → Prediction Summary.	SmartPLS GmbH (s. f.-g)
83	PLSpredict Q² predict	Indicador de capacidad predictiva fuera de muestra.	PLSpredict Report → Q ² predict.	SmartPLS GmbH (s. f.-g)
84	PLSpredict RMSE	<i>Root Mean Square Error</i> ; raíz del error cuadrático medio.	PLSpredict Report → Prediction Summary.	SmartPLS GmbH (s. f.-g)
85	PROCESS Models	Modelos tipo PROCESS para mediación, moderación y procesos condicionales.	Create Model / Calculate → PROCESS.	SmartPLS GmbH (s. f.-m)
86	Project View	Vista para administrar proyectos y archivos.	Workspace → seleccionar proyecto.	SmartPLS GmbH (s. f.-d)
87	Quality Criteria	Sección de criterios de calidad del modelo.	Open Full Report → Quality Criteria.	SmartPLS GmbH (s. f.-e)
88	Regression Bootstrapping	Bootstrapping para modelos de regresión.	Calculate → Regression Bootstrapping .	SmartPLS GmbH (s. f.-n)
89	Report Navigation	Panel de navegación entre resultados del reporte.	Report View → panel lateral.	SmartPLS GmbH (s. f.-e)
90	Report View / Open Full Report	Vista del reporte completo de resultados.	Después del cálculo → Open Full Report.	SmartPLS GmbH (s. f.-e)
91	rho_A	Coefficiente de confiabilidad del constructo usado en PLS-SEM .	Quality Criteria → Construct Reliability and Validity → rho_A.	SmartPLS GmbH (s. f.-e)

Juan Mejía Tréjo

CAPÍTULO 1. Bases del análisis estructural aplicado

Núm.	Control SmartPLS4	Descripción operativa	Cómo se accesa	Referencia SmartPLS
92	<i>rho_c</i> / Composite Reliability	Confiabilidad compuesta; evalúa consistencia interna considerando cargas externas.	Quality Criteria → Construct Reliability and Validity → <i>rho_c</i> / Composite Reliability.	SmartPLS GmbH (s. f.-e)
93	R-square / R^2	Coefficiente de determinación; varianza explicada del constructo endógeno.	Quality Criteria → R-square.	SmartPLS GmbH (s. f.-e)
94	Sample Projects	Proyectos de muestra para aprendizaje y práctica.	Documentation → Resources → Sample Projects.	SmartPLS GmbH (s. f.-d)
95	SAVE Model	Guarda el modelo dentro del proyecto.	Model Editor → Save.	SmartPLS GmbH (s. f.-e)
96	Select	Selecciona, mueve y organiza elementos del modelo.	Modeling View → Select.	SmartPLS GmbH (s. f.-e)
97	Simple Slope Plot	Gráfico de pendientes simples para moderación.	Moderation Results → Simple Slope Plot.	SmartPLS GmbH (s. f.-l)
98	Skewness	Asimetría de las variables del conjunto de datos.	Data View → Descriptives → Skewness.	SmartPLS GmbH (s. f.-p)
99	Specific Indirect Effects	Efectos indirectos específicos en mediación.	Final Results → Specific Indirect Effects.	SmartPLS GmbH (s. f.-a)
100	Thresholds / Result Colors	Colores o umbrales visuales usados para resaltar resultados.	Report View → Result Colors / Thresholds.	SmartPLS GmbH (s. f.-f)

Fuente: Elaboración propia con base en SmartPLS GmbH (s. f.-a, s. f.-d) y Hair et al. (2021).

CAPÍTULO 2. Bases operativas del modelado PLS-SEM en SmartPLS4



El Capítulo 2. Bases operativas del modelado PLS-SEM en SMARTPLS4 , constituye una **etapa metodológica decisiva** dentro del proceso de modelación estructural mediante **PLS-SEM**, porque transforma los datos recolectados en una **matriz analítica válida para su estimación**. En este sentido, la preparación del archivo no debe entenderse como una operación meramente informática, sino como un **procedimiento de aseguramiento de calidad de los datos**. La correcta organización de filas, columnas, indicadores, escalas y códigos permite que SmartPLS reconozca adecuadamente cada variable observada y la vincule con su constructo correspondiente. Esta fase inicial es congruente con la lógica de **PLS-SEM**, en la que la validez del modelo depende de la consistencia entre **teoría, medición y datos empíricos** (Hair et al., 2022).

La descarga del archivo desde **Google Forms o Google Sheets** en formato CSV representa el **tránsito entre la aplicación del cuestionario y el análisis estadístico propiamente dicho**. Cada fila debe corresponder a un **caso de análisis** y cada columna a un **ítem, indicador o variable observada**. Por ello, el investigador debe verificar que la base conserve una estructura tabular limpia, sin campos ajenos al análisis, sin encabezados extensos y sin valores que puedan alterar la lectura del software. Esta depuración es indispensable porque la medición de constructos en

Juan Mejía Trejo

ciencias sociales exige coherencia entre el **diseño del instrumento, la codificación de las respuestas y la posterior evaluación estadística de los indicadores** (Churchill, 1979).

Asimismo, el capítulo enfatiza la necesidad de sustituir las preguntas largas por **códigos breves, claros y sistemáticos**, como **CAIL1, EAR1, AFDJ1, HAIC1, MTPP1 o CAITIZEN1**. Esta decisión facilita la **trazabilidad entre el cuestionario, la base de datos, el modelo de medición y el reporte de resultados**. En estudios basados en constructos latentes, los indicadores no son columnas aisladas, sino **representaciones observables de conceptos teóricos**. Por tanto, una nomenclatura ordenada ayuda a evitar errores de asignación, mejora la lectura del modelo y fortalece la correspondencia entre los ítems y las dimensiones conceptuales evaluadas (Jarvis et al., 2003).

La revisión inicial del archivo importado en **SmartPLS4** permite diagnosticar si la base cumple **condiciones mínimas para continuar con el análisis**. Entre los criterios principales se encuentran el número de indicadores importados, el número de casos válidos, la ausencia de valores perdidos, el tipo de dato reconocido por el software, el rango observado de la escala y la variabilidad de las respuestas. Esta revisión preliminar es relevante porque **PLS-SEM**, aunque es flexible frente a supuestos de normalidad estricta, requiere **datos correctamente codificados y consistentes** para estimar relaciones entre constructos latentes de manera confiable (Chin, 1998).

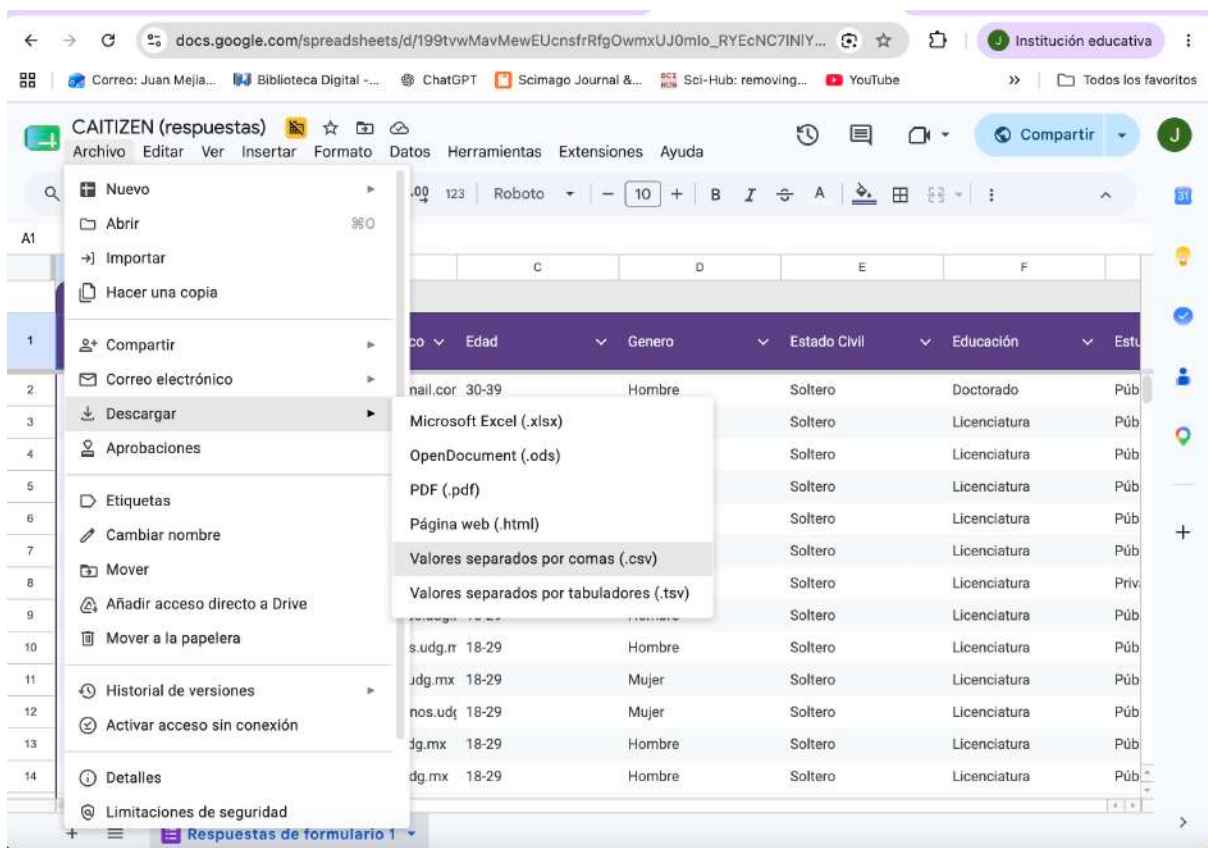
El diagnóstico de medias, medianas, valores mínimos y máximos, desviación estándar, asimetría, curtosis y prueba de normalidad permite valorar la **calidad operativa de la base antes de construir el modelo**. En particular, la ausencia de datos perdidos, la permanencia de las respuestas dentro del rango de la escala Likert y la existencia de variabilidad suficiente son condiciones que reducen riesgos de estimación e interpretación. Aunque **PLS-SEM** puede trabajar con distribuciones no normales, el uso de **Bootstrapping** resulta fundamental para evaluar la significancia de las relaciones del modelo cuando los indicadores no cumplen normalidad estricta (SmartPLS GmbH, s. f.-b).

Finalmente, la creación del **Workspace** y del proyecto en **SmartPLS4** organiza el **entorno técnico donde se almacenan la base, el modelo y los resultados**. En el caso **CAITIZEN**, la importación de **60 indicadores y 511 casos** permite avanzar hacia la construcción del modelo de medición y del modelo estructural. Esta preparación aporta al libro una ruta aplicada, replicable y verificable para que el lector comprenda que la calidad del análisis **PLS-SEM** comienza antes de ejecutar el algoritmo: inicia con una **base bien depurada, correctamente codificada y metodológicamente diagnosticada**. De este modo, el capítulo fortalece la **transparencia, la trazabilidad y la reproducibilidad del proceso analítico** (Hair et al., 2019).

Preparación del archivo para SmartPLS4

Una vez aplicado el cuestionario en **Google Forms**, las respuestas se almacenan automáticamente en una hoja de cálculo de **Google Sheets**. Esta hoja constituye la base inicial de trabajo para preparar los datos que posteriormente serán analizados en **SmartPLS** mediante **PLS-SEM**. Para utilizar esta información en **SmartPLS**, es necesario descargar el archivo en formato **CSV**. Este formato es recomendable porque conserva la estructura tabular de la información y facilita su importación al software. En **Google Sheets**, el procedimiento consiste en seleccionar la opción (Ver **Figura 2.1.**): **Archivo** → **Descargar** → **Valores separados por comas (.csv)**

Figura 2.1. Proceso de descarga del archivo de datos en formato CSV



Fuente: Elaboración propia con **SmartPLS4**.

Al descargar el archivo en formato **CSV**, cada fila representa un caso de análisis, es decir, una persona que respondió el cuestionario. A su vez, cada columna representa una pregunta, ítem o indicador del instrumento. Esta organización es indispensable para que SmartPLS pueda reconocer correctamente los datos y vincularlos con los constructos del modelo de medición.

Antes de importar el archivo en **SmartPLS**, la base **debe revisarse y depurarse**. Es recomendable sustituir los encabezados largos de las preguntas por códigos breves

Juan Mejía Trejo

y claros, por ejemplo: **CAIL1**, **CAIL2**, **EAR1**, **AFDJ1**, **HAIC1**, **MTPP1** o **CAITIZEN1**. Estos códigos permiten identificar con precisión cada indicador dentro del modelo y evitan errores ocasionados por textos extensos, acentos, signos de interrogación o caracteres especiales. También es necesario verificar que las respuestas de **escala Likert** estén codificadas numéricamente. Por ejemplo, cuando se utiliza una escala de cinco puntos, puede emplearse la siguiente codificación: 1 = Totalmente en desacuerdo, 2 = En desacuerdo, 3 = Ni de acuerdo ni en desacuerdo, 4 = De acuerdo y 5 = Totalmente de acuerdo. Esta codificación facilita el análisis estadístico y permite que **SmartPLS** procese adecuadamente los indicadores del modelo.

Asimismo, deben eliminarse o separarse las columnas que no formen parte del análisis **PLS-SEM**, tales como marca temporal, correo electrónico, nombre del participante u otros datos personales, salvo que se requieran para análisis descriptivos o de segmentación. También debe revisarse que no existan celdas vacías, respuestas duplicadas, valores fuera de rango o inconsistencias en la captura. Una vez depurada la base, el archivo **CSV** queda listo para importarse en **SmartPLS**. Esta etapa es fundamental porque la calidad del análisis depende, en gran medida, de la correcta preparación de los datos. Una base mal estructurada puede generar errores de importación, dificultades en la asignación de indicadores o problemas en la estimación del modelo. En síntesis, la descarga en formato **CSV** representa el paso de transición entre la recolección de datos en Google Forms y el análisis estadístico en **SmartPLS**. Por ello, no debe considerarse únicamente como una descarga técnica, sino como una fase metodológica de preparación de la matriz de datos para garantizar la validez operativa del análisis **PLS-SEM**.

Revisión inicial del archivo de datos importado en SmartPLS4

Una vez importado el archivo en **SmartPLS4**, el programa muestra una tabla descriptiva de los indicadores. Esta tabla permite verificar si la base fue importada correctamente y si los datos cumplen condiciones mínimas para continuar con la construcción del modelo **PLS-SEM**. Ver **Figura 2.2**.

Figura 2.2. Revisión inicial del archivo de datos importado en SmartPLS4

Name	No.	Type	Missings	Mean	Median	Scale min	Scale max	Observed min	Observed max	Standard deviation	Excess
CAIL1	1	MET	0	4.992	5.000	1.000	7.000	1.000	7.000	1.479	
CAIL2	2	MET	0	4.272	4.000	1.000	7.000	1.000	7.000	1.566	
CAIL3	3	MET	0	4.928	5.000	1.000	7.000	1.000	7.000	1.765	
CAIL4	4	MET	0	5.143	5.000	1.000	7.000	1.000	7.000	1.447	
CAIL5	5	MET	0	5.166	5.000	1.000	7.000	1.000	7.000	1.415	
CAIL6	6	MET	0	5.227	5.000	1.000	7.000	1.000	7.000	1.497	
CAIL7	7	MET	0	5.838	6.000	1.000	7.000	1.000	7.000	1.476	

Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

En el panel izquierdo se observan tres datos generales:

Indicators: indica el número total de indicadores importados. En este caso son 60, lo cual corresponde a los ítems cuantitativos del modelo.

Samples: indica el número de casos o respuestas válidas. En este caso son 511 participantes.

Missing values: indica el número total de valores perdidos detectados. En este caso es 0, por lo que la base no presenta datos faltantes.

Name

La columna **Name** muestra el nombre de cada indicador importado. Por ejemplo: **CAIL1**, **CAIL2**, **CAIL3**, etc. Los nombres deben ser breves, claros y consistentes con el constructo al que pertenecen. Lo recomendable es usar códigos sin acentos, sin espacios, sin signos especiales y sin preguntas largas.

Valor esperado: nombres cortos, por ejemplo: **CAIL1**, **EAR1**, **AFDJ1**, **HAIC1**, **MTPP1**, **CAITIZEN1**.

Si no cumple: si aparecen preguntas largas, acentos, comas, paréntesis o textos extensos, conviene regresar al CSV, renombrar las columnas y volver a importar la base.

No.

La columna **No.** indica el número consecutivo asignado por **SmartPLS4** a cada indicador. Sirve únicamente para identificar el orden de las variables en la base.

Valor esperado: numeración continua, por ejemplo, 1, 2, 3, 4, etc.

Si no cumple: si faltan números o el orden no coincide con la estructura esperada, se debe revisar si el archivo CSV contiene columnas eliminadas, duplicadas o mal importadas.

Type

La columna **Type** indica el tipo de dato reconocido por **SmartPLS4**. En la imagen aparece **MET**, que significa métrico. Para análisis **PLS-SEM** con escalas tipo Likert, es adecuado que los indicadores sean tratados como métricos, ya que **SmartPLS4** requiere que los datos sean numéricos.

Valor esperado: **MET**.

Si no cumple: si algún indicador no aparece como métrico, probablemente contiene texto, símbolos, espacios vacíos o respuestas no numéricas. En ese caso, debe corregirse la base en Excel o Google Sheets y volver a importar el archivo.

Missings

La columna **Missings** muestra cuántos valores perdidos tiene cada indicador. En la imagen, todos los indicadores muestran **0**, lo cual significa que no hay respuestas faltantes en esos ítems.

Valor esperado: 0. **Límite aceptable:** idealmente 0. Si existen datos perdidos, deben ser **menos del 10%** y estar justificados.

Si no cumple: si un indicador presenta valores perdidos, se debe revisar la base. Si los datos perdidos fueron codificados como **999**, ese valor debe declararse como marcador de dato perdido al importar. Si hay muchos datos faltantes en un indicador, puede considerarse su eliminación o imputación, según el criterio metodológico del estudio.

Mean

La columna **Mean** muestra la media o promedio de las respuestas de cada indicador. En una escala de 1 a 7, la media debe estar dentro de ese rango. Por ejemplo, una media de 4.992 indica que las respuestas promedio se ubican cerca del punto 5 de la escala.

Valor esperado: entre 1 y 7.

Interpretación: medias cercanas a 1 indican baja valoración; medias cercanas a 7 indican alta valoración; medias alrededor del centro indican posiciones intermedias.

Si no cumple: si aparece una media menor a 1 o mayor a 7, hay valores fuera de rango en la base, como 0, 8, 99 o 999. En ese caso, se debe corregir el archivo antes de continuar.

Median

La columna **Median** muestra la mediana, es decir, el valor central de las respuestas.

En escalas **Likert**, la mediana ayuda a identificar la tendencia central sin verse tan afectada por valores extremos.

Valor esperado: entre 1 y 7.

Interpretación: una mediana de 5 indica que el centro de las respuestas está en 5; una mediana de 6 indica mayor concentración hacia respuestas altas.

Si no cumple: si la mediana aparece fuera del rango de la escala, existe un error en la base. Deben revisarse valores atípicos o códigos mal tratados como respuestas válidas.

Scale min

La columna **Scale min** muestra el valor mínimo teórico de la escala definida para los indicadores. En este caso, como la escala va de 1 a 7, el valor mínimo esperado es 1.000.

Valor esperado: 1.000.

Juan Mejía Trejo

Si no cumple: si aparece 0 u otro valor, debe revisarse si **SmartPLS4** interpretó incorrectamente la escala o si en la base existen valores fuera del rango permitido.

Scale max

La columna **Scale max** muestra el valor máximo teórico de la escala.

En este caso, el valor máximo esperado es 7.000.

Valor esperado: 7.000.

Si no cumple: si aparece un valor mayor, como 8, 99 o 999, se debe revisar la base. Si 999 representa dato perdido, debe declararse como marcador de dato perdido y no dejarse como respuesta válida.

Observed min

La columna **Observed min** muestra el valor mínimo realmente observado en las respuestas. En una escala de 1 a 7, el valor mínimo observado debe ser igual o mayor que 1.000.

Valor esperado: entre 1.000 y 7.000.

Si no cumple: si aparece 0, -1 o algún valor inferior a 1, existe un error de captura o codificación. Debe corregirse antes de construir el modelo.

Observed max

La columna **Observed max** muestra el valor máximo realmente observado en las respuestas. En una escala de 1 a 7, el valor máximo observado debe ser igual o menor que 7.000.

Valor esperado: entre 1.000 y 7.000.

Si no cumple: si aparece 8, 99, 999 u otro valor mayor a 7, el dato no es válido como respuesta Likert. Debe corregirse, eliminarse o declararse como dato perdido si corresponde.

Standard deviation

La columna **Standard deviation** muestra la desviación estándar, es decir, qué tanto varían las respuestas respecto a la media. Una desviación estándar adecuada indica que existe variabilidad en las respuestas. En la imagen, **los valores cercanos a 1.4, 1.5 o 1.7** indican una dispersión razonable para una escala de 1 a 7.

Valor esperado: mayor que 0.

Rango práctico aceptable: en escalas de 1 a 7, valores aproximados entre **1.0 y 2.0** suelen indicar **variabilidad suficiente**.

Si no cumple: si la desviación estándar es cercana a 0, significa que casi todos respondieron lo mismo. En ese caso, el indicador puede tener poca capacidad para discriminar entre participantes y debe revisarse. **Si es excesivamente alta, debe verificarse si existen valores fuera de rango. Ver Figura 2.3.**

Juan Mejía Trejo

Figura 2.3. Continuación de la revisión descriptiva del archivo de datos importado en SmartPLS4

SmartPLS Edit Themes

 Back

 Setup

 Derive indicators

 Add group

 Generate groups

 Clear groups

 Export to Excel / CSV

CAITZENPLS

Indicators

Copy to Excel/Word

Indicators 60

Samples 511

Missing values 0

● Indicators

○ Correlations

○ Data groups

○ Raw data

Name	an	Scale min	Scale max	Observed min	Observed max	Standard deviation	Excess kurtosis	Skewness	Cramér-von Mises p value
CAIL1	00	1.000	7.000	1.000	7.000	1.479	0.184	-0.671	0.000
CAIL2	00	1.000	7.000	1.000	7.000	1.566	-0.590	-0.162	0.000
CAIL3	00	1.000	7.000	1.000	7.000	1.765	-0.695	-0.483	0.000
CAIL4	00	1.000	7.000	1.000	7.000	1.447	-0.207	-0.519	0.000
CAIL5	00	1.000	7.000	1.000	7.000	1.415	-0.051	-0.625	0.000
CAIL6	00	1.000	7.000	1.000	7.000	1.497	0.155	-0.792	0.000
CAIL7	00	1.000	7.000	1.000	7.000	1.476	1.157	-1.300	0.000

Fuente: Elaboración propia con SmartPLS4.

Excess kurtosis

La columna **Excess kurtosis** muestra la curtosis, es decir, el grado de concentración de las respuestas. Valores cercanos a 0 indican una distribución más equilibrada. Valores muy altos indican concentración excesiva en ciertos puntos. Valores muy negativos indican una distribución más plana.

Valor esperado: cercano a 0.

Rango práctico aceptable: entre -2 y +2.

Si no cumple: si la curtosis supera esos límites, no necesariamente se elimina el indicador, pero debe revisarse. En PLS-SEM la normalidad estricta no es obligatoria, pero una curtosis extrema puede afectar la estabilidad de los resultados. En ese caso, se recomienda revisar datos atípicos y aplicar **Bootstrapping**.

Skewness

La columna **Skewness** muestra la asimetría de las respuestas. Un valor cercano a 0 indica una distribución equilibrada. Valores negativos indican concentración hacia respuestas altas. Valores positivos indican concentración hacia respuestas bajas.

Valor esperado: cercano a 0.

Rango práctico aceptable: entre -2 y +2.

Si no cumple: si la asimetría es muy alta, se debe revisar si el ítem está sesgado o si los participantes respondieron mayoritariamente en un extremo de la escala. En PLS-SEM no se exige normalidad perfecta, pero una asimetría extrema debe reportarse y considerarse en la interpretación.

Cramér-von Mises p value

La columna **Cramér-von Mises p value** evalúa la normalidad de la distribución del indicador. Cuando el valor p es menor a 0.05, se rechaza la hipótesis de

Juan Mejía Trejo

normalidad. En la imagen se observan valores 0.000, lo cual indica que varios indicadores no siguen una distribución normal.

Valor esperado para normalidad: p mayor que 0.05.

Si no cumple: si el valor p es menor a 0.05, significa que el indicador no presenta distribución normal. En **PLS-SEM** esto no invalida el análisis, porque el método no exige normalidad multivariada estricta. La decisión recomendada es continuar con el análisis, pero utilizar **Bootstrapping** para evaluar la significancia estadística de las relaciones del modelo.

Decisiones metodológicas si algún criterio no se cumple

Si existen valores perdidos, se debe revisar su origen. Cuando se usa **999** como código de dato perdido, debe declararse durante la importación. Si los **valores perdidos** son pocos, puede considerarse imputación. Si son muchos en un mismo indicador, puede evaluarse su eliminación. Si aparecen valores fuera del rango 1–7, se debe corregir la base antes de continuar. Ningún valor menor a 1 o mayor a 7 debe permanecer como respuesta válida.

Si un indicador tiene desviación estándar muy baja, debe revisarse porque puede no discriminar adecuadamente entre participantes. Si todos responden casi lo mismo, el indicador puede aportar poca información al modelo. Si existen problemas de asimetría o curtosis, no se elimina automáticamente el indicador. En **PLS-SEM** estos problemas se atienden principalmente mediante **Bootstrapping** y revisión de la calidad del modelo de medición. Si la prueba de **Cramér-von Mises** indica no normalidad, se puede continuar con **PLS-SEM**, ya que este enfoque es adecuado para datos que no cumplen normalidad estricta. Sin embargo, debe utilizarse **Bootstrapping** y reportar los resultados con cuidado.

En síntesis, esta tabla permite verificar **cuatro aspectos básicos** antes de construir el modelo: que la base tenga los indicadores correctos, que no existan perdidos, que todos los datos estén dentro del rango de la escala y que los indicadores presenten variabilidad suficiente para ser analizados.

Diagnóstico de cumplimiento de la base CAITIZEN en SmartPLS4

La tabla inicial de **SmartPLS4** permite realizar un diagnóstico preliminar de la calidad del archivo antes de construir el modelo **PLS-SEM**. En el caso de la base **CAITIZEN**, se observa que fueron importados **60 indicadores y 511 casos**, sin **valores perdidos** detectados. Esto indica que la base fue leída correctamente por el software y que puede utilizarse para la construcción del modelo. Ver **Tabla 2.1**.

Tabla 2.1. Diagnóstico general del archivo CAITIZEN en SmartPLS4

Criterio evaluado	Resultado observado	Criterio esperado	Diagnóstico
Indicadores importados	60	Coincidir con los ítems del modelo	Cumple
Casos válidos	511	Muestra suficiente para PLS-SEM	Cumple
Valores perdidos	0	Idealmente 0	Cumple
Rango teórico de la escala	1 a 7	1 a 7	Cumple
Valores mínimos observados	1	No menor a 1	Cumple
Valores máximos observados	7	No mayor a 7	Cumple
Tipo de dato	MET	Métrico / numérico	Cumple
Desviación estándar	Aprox. 1.33 a 1.77	Mayor que 0; deseable con variabilidad suficiente	Cumple
Asimetría	Aprox. -1.30 a -0.16	Preferentemente entre -2 y +2	Cumple
Curtosis excesiva	Aprox. -0.69 a 1.46	Preferentemente entre -2 y +2	Cumple

Fuente: Elaboración propia

Interpretación del diagnóstico

La base presenta **60** indicadores cuantitativos, correspondientes a los constructos **CAIL, EAR, AFDJ, HAIC, MTPP y CAITIZEN**. El número de indicadores coincide con la estructura prevista del instrumento, ya que cada constructo cuenta con **10 ítems**. El número de casos válidos es de **511**, lo cual resulta adecuado para continuar con el análisis **PLS-SEM**. Además, **SmartPLS4** reporta **0** valores perdidos, por lo que no es necesario aplicar imputación ni eliminar casos por datos faltantes. Los valores observados se encuentran dentro del rango esperado de la escala **Likert de 1 a 7**. El valor mínimo observado es 1 y el valor máximo observado es 7, lo que confirma que no existen valores fuera de rango, como **0, 8, 99 o 999**. Por tanto, no se detectan errores de codificación en la escala. La columna Type muestra que los indicadores fueron reconocidos como **MET**, es decir, como variables métricas. Esto es adecuado para **SmartPLS4**, ya que los ítems tipo Likert se procesan como variables numéricas dentro del análisis **PLS-SEM**.

La desviación estándar se encuentra aproximadamente entre **1.33 y 1.77**, lo cual indica que los **indicadores presentan variabilidad suficiente**. Esto es importante porque un indicador con desviación estándar cercana a 0 tendría poca capacidad para diferenciar entre participantes. La asimetría se mantiene aproximadamente entre **-1.30 y -0.16**, dentro del rango práctico aceptable de **-2 a +2**. Los valores negativos indican que las respuestas tienden hacia valores altos de la escala, lo cual es común en cuestionarios actitudinales. Sin embargo, no se observa asimetría extrema.

La **curtosis excesiva** se encuentra aproximadamente entre **-0.69 y 1.46**, también dentro del rango práctico aceptable de **-2 a +2**. Esto indica que no existen

Juan Mejía Trejo

concentraciones extremas de respuestas que comprometan la revisión inicial de los datos. Ver **Tabla 2.2**.

Tabla 2.2. Diagnóstico por constructo del archivo CAITIZEN

Constructo	Indicadores	Valores perdidos	Rango observado	Variabilidad	Asimetría / curtosis	Diagnóstico
CAIL	10	0	1 a 7	Adecuada	Dentro de rango	Cumple
EAR	10	0	1 a 7	Adecuada	Dentro de rango	Cumple
AFDJ	10	0	1 a 7	Adecuada	Dentro de rango	Cumple
HAIC	10	0	1 a 7	Adecuada	Dentro de rango	Cumple
MTPP	10	0	1 a 7	Adecuada	Dentro de rango	Cumple
CAITIZEN	10	0	1 a 7	Adecuada	Dentro de rango	Cumple

Fuente: Elaboración propia

Decisión metodológica

Con base en la revisión inicial de **SmartPLS4**, la base **CAITIZEN** cumple con los criterios mínimos para continuar con el análisis **PLS-SEM**. No se observan datos perdidos, valores fuera de rango, errores de codificación ni indicadores sin variabilidad suficiente. Por tanto, la decisión metodológica recomendada es **continuar con la construcción del modelo de medición y del modelo estructural en SmartPLS4**. No es necesario eliminar indicadores en esta etapa. La evaluación definitiva de los ítems deberá realizarse posteriormente mediante los criterios del modelo de medición, tales como cargas externas, confiabilidad interna, validez convergente y validez discriminante.

Conclusión del diagnóstico

El archivo **CAITIZEN** se encuentra correctamente importado y cumple con las condiciones preliminares para el análisis **PLS-SEM**. Los 60 indicadores presentan datos completos, valores dentro del rango esperado de la escala de 1 a 7 y variabilidad suficiente. En consecuencia, se puede avanzar a la construcción del modelo en **SmartPLS4**.

Creación del área de trabajo en SmartPLS

Una vez descargado el archivo en formato **CSV** desde **Google Sheets**, se debe abrir **SmartPLS 4**, versión utilizada en esta guía para desarrollar el análisis **PLS-SEM**. Esta versión permite crear un área de trabajo, importar bases de datos, construir modelos de medición y estructurales, así como ejecutar procedimientos de estimación

Juan Mejía Tréjo

como el algoritmo PLS y **Bootstrapping**. Antes de importar el archivo CSV, es importante que la base haya sido revisada previamente, verificando que los nombres de los ítems sean breves, que las respuestas estén codificadas numéricamente y que no existan columnas innecesarias, celdas vacías o datos duplicados.

El primer paso dentro de **SmartPLS4** consiste en crear o seleccionar el área de trabajo, denominada en el programa como **Workspace**. El **Workspace** es la carpeta local donde se almacenarán todos los archivos relacionados con el análisis, incluyendo el proyecto, el archivo, el modelo **PLS-SEM** y los resultados generados. Por ello, no debe considerarse únicamente como una carpeta técnica, sino como el espacio principal de organización del estudio. Si al ingresar a **SmartPLS4** aparece el mensaje **“No Workspace directory selected”**, significa que todavía no se ha definido un área de trabajo. Para continuar, se debe seleccionar la opción **New Workspace** y elegir o crear una carpeta en la computadora. Para mantener el orden del análisis, se recomienda nombrar esta carpeta como **CAITIZEN**. De esta manera, todos los archivos del modelo quedarán concentrados en un mismo lugar.

Una vez creada el área de trabajo, el siguiente paso es generar un nuevo proyecto. Para ello, se selecciona la opción **New project**. El nombre del proyecto será simplemente **CAITIZEN**, ya que identifica de manera clara, directa y suficiente el modelo que será analizado. Este proyecto funcionará como el entorno principal donde se incorporará el archivo, se construirán los constructos, se asignarán los indicadores y se estimará el modelo estructural. Después de crear el proyecto **CAITIZEN**, se procede a importar el archivo **CSV** previamente descargado desde Google Sheets. **SmartPLS4** reconocerá la base como una matriz de datos: cada fila representa un participante o caso de análisis, mientras que cada columna corresponde a un ítem, indicador o variable observada del cuestionario.

Antes de construir el modelo, es indispensable verificar que la importación se haya realizado correctamente. Debe revisarse que todos los indicadores aparezcan con nombres claros, que las respuestas estén en formato numérico y que no existan errores de lectura. Esta revisión es fundamental porque una base mal estructurada puede ocasionar problemas en la asignación de indicadores, en la estimación del modelo o en la interpretación de los resultados. Después de validar el archivo, se procede a crear los constructos del modelo **CAITIZEN** y a vincular cada uno con sus indicadores correspondientes. Por ejemplo, los ítems de alfabetización crítica en inteligencia artificial se asignan al constructo **CAIL**; los de responsabilidad ética a **EAR**; los de justicia algorítmica a **AFDJ**; los de colaboración humano-IA a **HAIC**; los de metacognición y participación a **MTPP**; y los ítems de ciudadanía asistida por IA al constructo **CAITIZEN**.

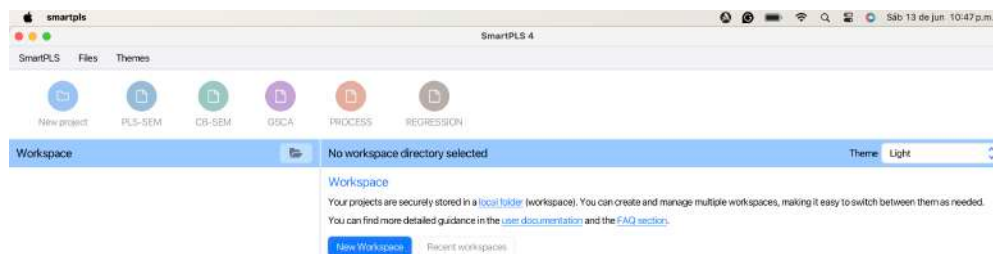
Finalmente, se establecen las relaciones estructurales entre los constructos conforme al modelo teórico propuesto. Con ello, el área de trabajo y el proyecto **CAITIZEN** quedan organizados y preparados para ejecutar el análisis **PLS-SEM** en **SmartPLS 4.1.1.8**, incluyendo el algoritmo PLS, la evaluación del modelo de medición, la evaluación del modelo estructural y el procedimiento de **Bootstrapping**.

Procedimiento paso a paso

Para realizar este procedimiento en **SmartPLS4**, se recomienda seguir la siguiente secuencia:

- Abrir **SmartPLS4** en la computadora.
- Revisar si en la pantalla inicial aparece el mensaje “**No Workspace directory selected**”.
- Si aparece ese mensaje, seleccionar la opción **New Workspace**.
- Elegir una ubicación en la computadora donde se guardarán los archivos del análisis.
- Crear una carpeta nueva con el nombre **CAITIZEN**.
- Seleccionar esa carpeta como área de trabajo o **Workspace**. Ver **Figura 2.4**.

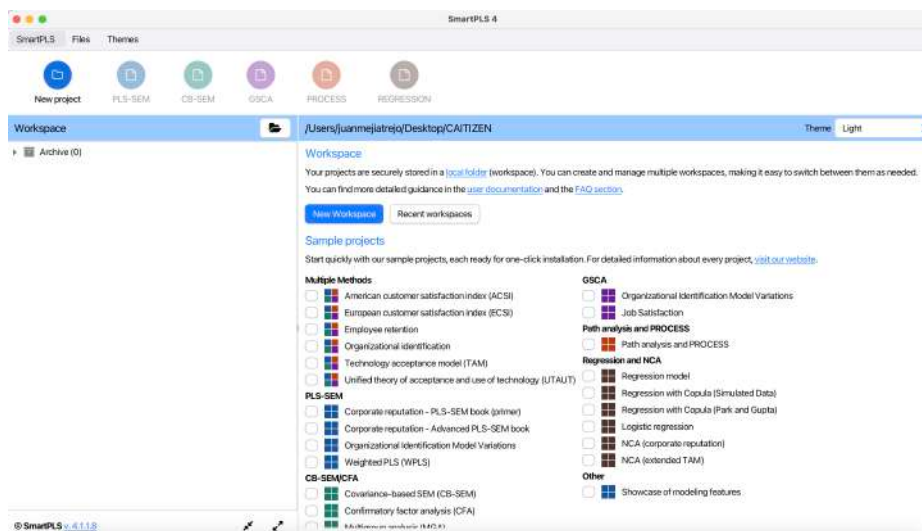
Figura 2.4. Creación del Workspace CAITIZEN en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

- Verificar que **CAITIZEN** aparezca como **Workspace**, donde **SmartPLS4** reconoce el área de trabajo. Ver **Figura 2.5**.

Figura 2.5. Reconocimiento del Workspace CAITIZEN en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Trejo

- Una vez definida el área de trabajo, seleccionar la opción **New project**.
- Escribir como nombre del proyecto: **CAITIZEN**. Ver **Figura 2.6**.

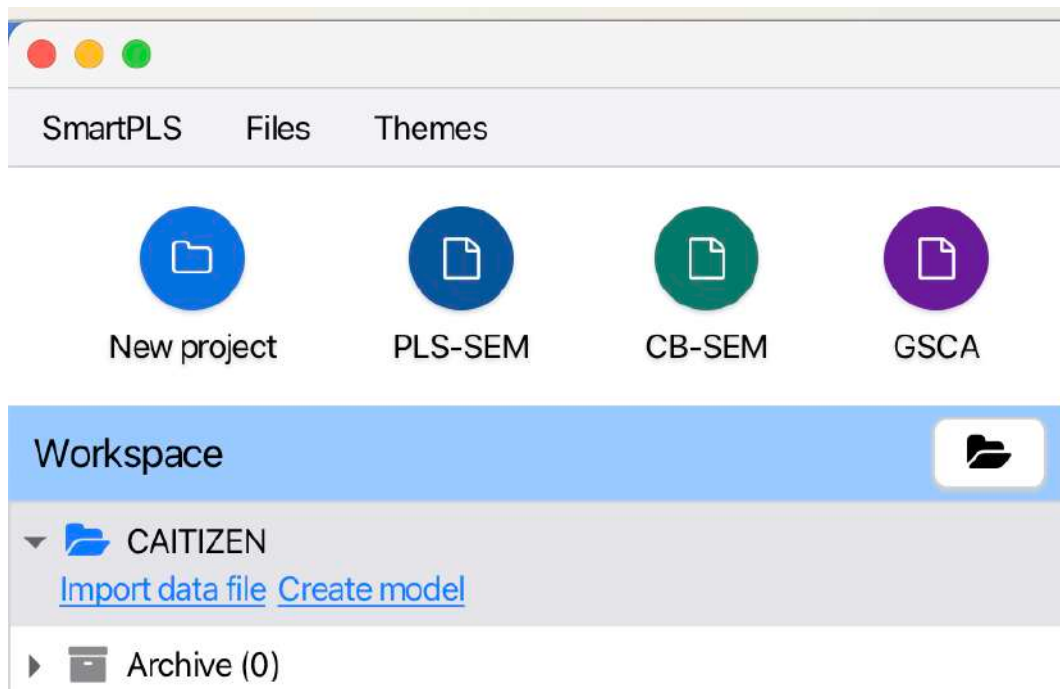
Figura 2.6. Asignación del nombre del proyecto CAITIZEN en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

- Confirmar la creación del proyecto.
- Ingresar al proyecto **CAITIZEN** dentro del **Workspace**.
- Seleccionar la opción para importar el archivo de datos. Ver **Figura 2.7**.

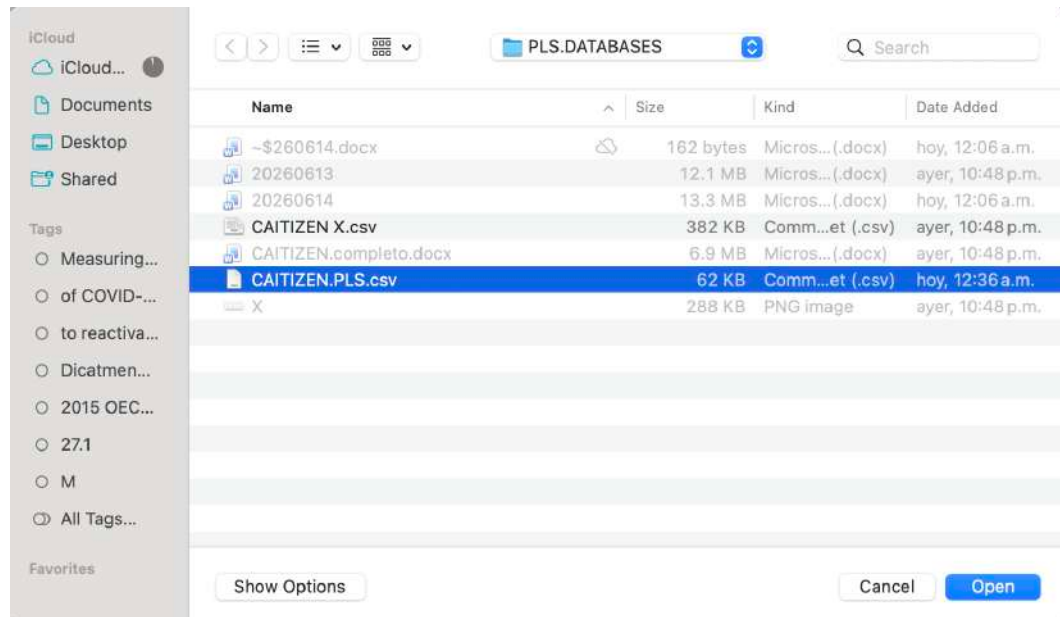
Figura 2.7. Verificación del proyecto CAITIZEN dentro del Workspace



Fuente: Elaboración propia con **SmartPLS4**.

- Buscar en la computadora el archivo **CAITIZEN.PLS.CSV** previamente descargado desde **Google Sheets**. Ver **Figura 2.8**.

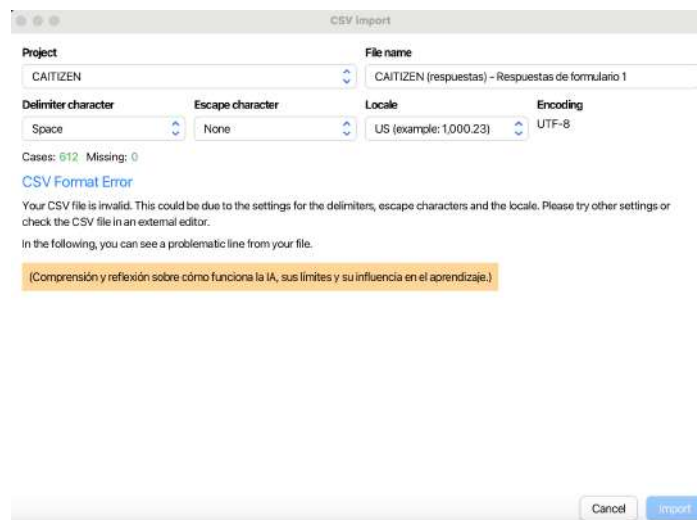
Figura 2.8. Búsqueda del archivo CAITIZEN.PLS.CSV para importación



Fuente: Elaboración propia desde navegador Google Chrome

- Seleccionar el archivo **CSV** correspondiente al cuestionario.
- Importar el archivo dentro del proyecto **CAITIZEN**. Si hay problemas se mostrará con este aviso. Ver **Figura 2.9**.

Figura 2.10. Configuración de importación del archivo CSV en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

- Si todo marcha bien, se mostrará.
 - a. Verificar que el archivo se lea correctamente.
 - b. Configurar **Delimiter character: Comma**.
 - c. Configurar **Escape character: None**.
 - d. Mantener **Locale: US**.
 - e. Confirmar **Encoding: UTF-8**.
 - f. Si hay datos perdidos codificados como **999**, seleccionar **Define a missing value marker**.
 - g. Escribir **999**.
- Revisar que **SmartPLS4** reconozca correctamente las columnas del archivo de datos.
- Verificar que cada columna corresponda a un ítem o indicador del cuestionario.
- Confirmar que los nombres de los indicadores sean breves y claros, por ejemplo: **CAIL1, CAIL2, EAR1, AFDJ1, HAIC1, MTPP1** o **CAITIZEN1**.
- Revisar que las respuestas estén registradas en formato numérico.
- **Verificar que no existan columnas innecesarias**, celdas vacías, valores fuera de rango o errores de lectura.
- Guardar los cambios realizados dentro del proyecto.
- Crear los constructos del modelo **CAITIZEN**.
- Asignar a cada constructo sus indicadores correspondientes.
- Vincular los indicadores de alfabetización crítica en **IA** al constructo **CAIL**.
- Vincular los indicadores de responsabilidad ética al constructo **EAR**.
- Vincular los indicadores de justicia algorítmica al constructo **AFDJ**.
- Vincular los indicadores de colaboración humano-IA al constructo **HAIC**.
- Vincular los indicadores de metacognición y participación al constructo **MTPP**.
- Vincular los indicadores de ciudadanía asistida por **IA** al constructo **CAITIZEN**.
- Establecer las relaciones estructurales entre los constructos conforme al modelo teórico.
- Guardar el modelo dentro del proyecto **CAITIZEN**.
- Continuar con la ejecución del análisis **PLS-SEM** en **SmartPLS4**. Oprimir **Importar** para incorporar el archivo CSV al proyecto **CAITIZEN**. Ver **Figura 2.10**.

Figura 2.10. Importación del archivo CSV en SmartPLS4

Variable	Source	No.	Type	Missing	Mean	Model	Scale min	Scale max	Observed min	Observed max	Standard deviation	Bias
CAIL1	1	CAIL1	0	4.892	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAIL2	2	CAIL2	0	4.272	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
EAR1	3	EAR1	0	4.938	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
AFDJ1	4	AFDJ1	0	5.362	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
HAIC1	5	HAIC1	0	5.368	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
MTPP1	6	MTPP1	0	5.327	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN1	7	CAITIZEN1	0	5.376	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN2	8	CAITIZEN2	0	5.304	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN3	9	CAITIZEN3	0	5.162	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN4	10	CAITIZEN4	0	5.346	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN5	11	CAITIZEN5	0	4.753	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN6	12	CAITIZEN6	0	4.708	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN7	13	CAITIZEN7	0	5.296	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN8	14	CAITIZEN8	0	5.337	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN9	15	CAITIZEN9	0	5.018	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN10	16	CAITIZEN10	0	5.431	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN11	17	CAITIZEN11	0	5.041	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN12	18	CAITIZEN12	0	5.465	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000
CAITIZEN13	19	CAITIZEN13	0	5.363	0.000	1.000	1.000	1.000	1.000	1.000	1.000	0.000

Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Trejo

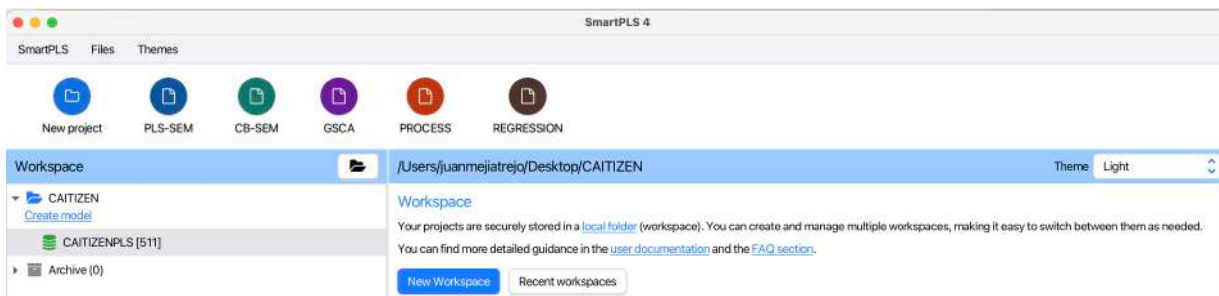
Pantalla posterior a la creación del proyecto e importación del CSV

Después de crear el área de trabajo, generar el proyecto **CAITIZEN** e importar el archivo **CSV**, **SmartPLS4** muestra en el panel izquierdo la estructura del proyecto. En esta pantalla se observa que el **Workspace CAITIZEN** fue creado correctamente y que el archivo de datos **CAITIZENPLS [511]** ya fue importado. El número entre corchetes indica la cantidad de casos disponibles para el análisis.

En esta etapa, el archivo de datos ya está incorporado al proyecto; por tanto, el siguiente paso consiste en crear el modelo **PLS-SEM**. Para ello, se debe seleccionar la opción **Create model**, ubicada debajo del proyecto **CAITIZEN**. Esta opción permite abrir el espacio gráfico donde se construirán los constructos, se asignarán los indicadores y se establecerán las relaciones estructurales del modelo.

A partir de esta pantalla se pueden realizar tres acciones principales: revisar que el archivo de datos esté correctamente cargado, crear el modelo de análisis y organizar los archivos del proyecto. Si la base aparece con el nombre **CAITIZENPLS [511]**, significa que **SmartPLS4** reconoce los datos importados y que se puede continuar con la construcción del modelo. Así, con base en la pantalla sobre el **Proyecto CAITIZEN con archivo de datos importado en SmartPLS4**. La pantalla muestra el **Workspace CAITIZEN**, el proyecto creado y la base **CAITIZEN** con **511 casos**. A partir de este punto, se debe seleccionar **Create model** para construir el modelo **PLS-SEM**. Ver **Figura 2.11**.

Figura 2.11. Confirmación del proyecto CAITIZEN con archivo de datos importado



Fuente: Elaboración propia con **SmartPLS4**.

En consecuencia, se puede avanzar a la construcción del modelo en **SmartPLS4**.

Juan Mejía Tréjo

Síntesis del procedimiento

La preparación e importación de la base de datos en SmartPLS4 constituye la fase inicial para asegurar que los indicadores, casos y valores cumplan las condiciones mínimas requeridas antes de construir el modelo PLS-SEM. El procedimiento se desarrolló de forma secuencial y quedó documentado mediante las figuras y tablas del capítulo.

- **Descarga de la base:** exportar desde Google Sheets el archivo en formato CSV, como se muestra en la **Figura 2.1**.
- **Revisión inicial:** comprobar en SmartPLS4 el número de indicadores, casos válidos, valores perdidos y tipo de dato, de acuerdo con la **Figura 2.2**.
- **Diagnóstico descriptivo:** revisar medias, medianas, mínimos, máximos, desviación estándar, asimetría, curtosis y normalidad, como se presenta en la **Figura 2.3**.
- **Validación general:** confirmar que la base contiene 60 indicadores, 511 casos, cero valores perdidos, valores entre 1 y 7 y suficiente variabilidad, según la **Tabla 2.1**.
- **Validación por constructo:** verificar que CAIL, EAR, AFDJ, HAIC, MTPP y CAITIZEN cuenten con diez indicadores y cumplan los criterios establecidos, conforme a la **Tabla 2.2**.
- **Creación del Workspace:** establecer el área de trabajo CAITIZEN en SmartPLS4, como se observa en las **Figuras 2.4 y 2.5**.
- **Creación del proyecto:** asignar el nombre CAITIZEN al proyecto y verificar su ubicación dentro del Workspace, según las **Figuras 2.6 y 2.7**.
- **Selección del archivo:** localizar y seleccionar el archivo CAITIZEN.PLS.CSV, como se muestra en la **Figura 2.8**.
- **Configuración e importación:** definir delimitador, codificación y parámetros de lectura, e importar el archivo, conforme a las **Figuras 2.9 y 2.10**.
- **Confirmación final:** comprobar que la base CAITIZENPLS [511] aparezca correctamente dentro del proyecto, como se presenta en la **Figura 2.11**.
- **Decisión metodológica:** continuar con la creación del modelo, la asignación de indicadores y la estimación **PLS-SEM**, debido a que la base cumple las condiciones preliminares requeridas.

CAPÍTULO 3. Construcción del modelo en SmartPLS4



El **Capítulo 3. Construcción del modelo en SmartPLS4** introduce al lector en una fase central del análisis estructural aplicado: la transformación de una base de datos correctamente importada en un **modelo gráfico, teóricamente organizado y metodológicamente estimable**. Después de preparar el archivo y verificar la calidad inicial de los datos, el investigador debe construir el modelo **PLS-SEM** dentro del área de trabajo de **SmartPLS4**. Esta etapa resulta fundamental porque permite articular visualmente los **constructos, sus indicadores y las relaciones hipotéticas** que orientan el estudio. En este sentido, el capítulo aporta una ruta aplicada para pasar del dato organizado al modelo estructural formal (Hair et al., 2022).

La creación del modelo **CAITIZEN** inicia con la **selección del proyecto, la definición del tipo de modelo y la asignación de un nombre operativo** dentro del software. Aunque estas acciones parecen técnicas, tienen una función metodológica relevante: asegurar que el análisis se realice en el proyecto correcto, con la base correspondiente y bajo el enfoque **PLS-SEM**. La elección de **PLS-SEM** responde a su utilidad para analizar modelos orientados a la **explicación y predicción de constructos latentes**, especialmente cuando se trabaja con relaciones entre variables teóricas complejas. Por ello, la configuración inicial del modelo constituye un punto de control entre la preparación de datos y la estimación estadística (Chin, 1998).

Un componente esencial del capítulo es la **creación de constructos latentes reflectivos**. En el modelo **CAITIZEN**, los constructos **CAIL, EAR, AFDJ, HAIC, MTPP**

Juan Mejía Trejo

y **CAITIZEN** deben representarse gráficamente y vincularse con sus respectivos indicadores. Esta tarea no solo organiza visualmente el modelo, sino que define la **lógica de medición** que se aplicará posteriormente. En modelos reflectivos, los indicadores se entienden como **manifestaciones observables del constructo latente**; por tanto, deben mantener correspondencia conceptual con la variable que representan. Esta precisión evita problemas de especificación y fortalece la congruencia entre **teoría, medición y análisis empírico** (Jarvis et al., 2003).

La asignación de indicadores a cada constructo también cumple una función de **trazabilidad metodológica**. Cada ítem importado desde la base debe incorporarse al constructo correspondiente, verificando que no existan **indicadores faltantes, duplicados o mal asignados**. En el caso **CAITIZEN**, cada constructo integra diez indicadores, lo cual exige revisar cuidadosamente la estructura completa antes de estimar el modelo. Esta revisión previa es relevante porque los errores de asignación pueden alterar las **cargas externas, la confiabilidad, la validez convergente y la validez discriminante**. En consecuencia, construir correctamente el modelo de medición es una condición indispensable para evaluar después la calidad estadística del instrumento (Fornell & Larcker, 1981).

El capítulo también destaca la importancia de establecer las **relaciones estructurales** entre los constructos conforme al modelo teórico propuesto. En **CAITIZEN**, CAIL, EAR, AFDJ, HAIC y MTPP se configuran como **predictores directos** de la ciudadanía sostenible asistida por inteligencia artificial. La dirección de las flechas no es un asunto visual secundario, sino una **decisión teórica** que representa hipótesis causales o predictivas. Por ello, cada relación debe dibujarse cuidadosamente desde la variable independiente hacia la variable dependiente. Esta organización permite que el modelo estructural exprese de manera clara las rutas analíticas que serán evaluadas mediante el algoritmo **PLS-SEM** (Hair et al., 2019).

Finalmente, la organización visual del modelo, la revisión de propiedades de los constructos y la confirmación del tipo de medición preparan el análisis para su **estimación posterior**. Antes de ejecutar el algoritmo, el investigador debe comprobar que todos los constructos estén correctamente nombrados, que los indicadores hayan sido asignados y que las relaciones estructurales coincidan con el planteamiento conceptual. De este modo, el capítulo fortalece la **reproducibilidad del procedimiento** y evita que el análisis estadístico se realice sobre una estructura incompleta o equivocada. Su aporte principal consiste en mostrar que la calidad del **PLS-SEM** comienza con un **modelo bien construido, ordenado y teóricamente coherente** (SmartPLS GmbH, s. f.-e).

Inicio de la construcción del modelo CAITIZEN

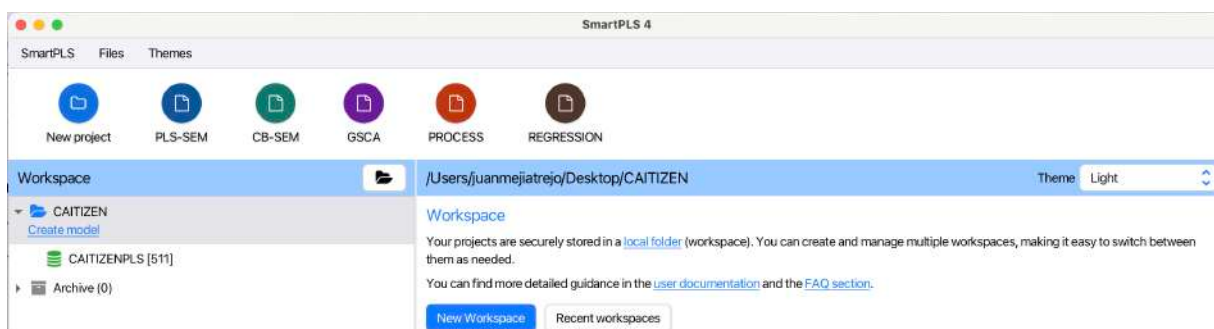
Una vez creado el **Workspace**, generado el proyecto **CAITIZEN** e importada correctamente la base de datos **CAITIZENPLS [511]**, el siguiente paso consiste en construir el modelo **PLS-SEM** dentro de **SmartPLS4**. Esta etapa permite representar gráficamente los constructos latentes, asignar los indicadores observables y establecer

Juan Mejía Tréjo

las relaciones estructurales que serán evaluadas posteriormente mediante el algoritmo **PLS-SEM** y el procedimiento de **Bootstrapping**. En esta fase todavía no se interpretan resultados estadísticos; el objetivo es organizar correctamente el modelo gráfico para que la estimación posterior corresponda al planteamiento conceptual de **CAITIZEN**.

La **Figura 3.1** muestra que el proyecto **CAITIZEN** ya se encuentra creado dentro del **Workspace** y que el archivo de datos **CAITIZENPLS [511]** fue importado correctamente. El número entre corchetes indica que la base contiene 511 casos disponibles para el análisis. A partir de esta pantalla, se debe seleccionar la opción **Create model**, ubicada debajo del proyecto. Esta acción abre el proceso para construir el modelo **PLS-SEM**. Por tanto, la figura documenta el punto de transición entre la importación de datos y la construcción formal del modelo gráfico en **SmartPLS4**.

Figura 3.1 Proyecto CAITIZEN con archivo de datos importado y opción Create model en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

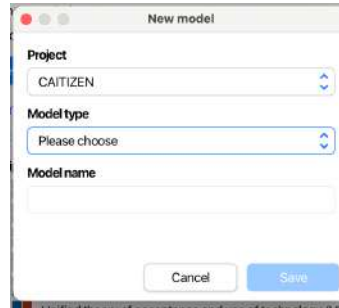
Para iniciar la construcción del modelo, se selecciona **Create model** dentro del proyecto **CAITIZEN**. **SmartPLS4** muestra entonces la ventana **New model**, donde se definen tres elementos básicos: el proyecto, el tipo de modelo y el nombre del modelo. En el campo **Project**, debe aparecer seleccionado **CAITIZEN**, lo cual confirma que el modelo será creado dentro del proyecto correcto y asociado con la base de datos previamente importada. Esta verificación es importante porque evita crear modelos fuera del proyecto correspondiente o usar accidentalmente una base de datos equivocada.

Creación del modelo PLS-SEM en SmartPLS4

La **Figura 3.2** presenta la ventana inicial **New model**, donde todavía no se ha seleccionado el tipo de modelo ni se ha escrito el nombre del modelo. Esta imagen es útil porque muestra el estado previo a la configuración del análisis. En este punto, el usuario debe confirmar que el campo **Project** indique **CAITIZEN**. Después debe desplegar el campo **Model type** para elegir **PLS-SEM**. Finalmente, debe escribir el nombre del modelo. Esta captura permite explicar al lector que la construcción del modelo inicia con una configuración explícita antes de acceder al área gráfica de trabajo.

Juan Mejía Tréjo

Figura 3.2. Ventana *New model* para crear el modelo CAITIZEN en SmartPLS4



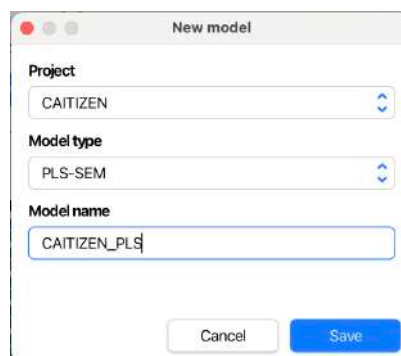
Fuente: Elaboración propia con **SmartPLS4**.

En la misma ventana ***New model***, el campo ***Model type*** debe configurarse como **PLS-SEM**, ya que el objetivo es construir y evaluar un modelo de medición y un modelo estructural mediante mínimos cuadrados parciales. En el campo ***Model name***, se recomienda escribir **CAITIZEN** para mantener consistencia conceptual. Si en la captura aparece **CAITIZEN_PLS**, puede conservarse como nombre operativo del archivo de **SmartPLS4**, pero editorialmente conviene explicar que el modelo conceptual se denomina **CAITIZEN**, mientras que **PLS-SEM** corresponde a la técnica estadística utilizada.

Apertura del área gráfica de trabajo

La **Figura 3.3** muestra la configuración completa de la ventana ***New model***, con el proyecto **CAITIZEN**, el tipo de modelo **PLS-SEM** y el nombre operativo **CAITIZEN_PLS**. Para fines metodológicos, esta captura documenta que el análisis se realizará mediante **PLS-SEM**. Para fines editoriales, debe aclararse que **CAITIZEN** es el nombre conceptual del modelo y **CAITIZEN_PLS** puede ser solo una etiqueta operativa dentro del software. Una vez revisados estos campos, se presiona **SAVE** para abrir el área gráfica donde se construirán los constructos y relaciones estructurales.

Figura 3.3. Selección del modelo PLS-SEM y asignación del nombre operativo del modelo en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

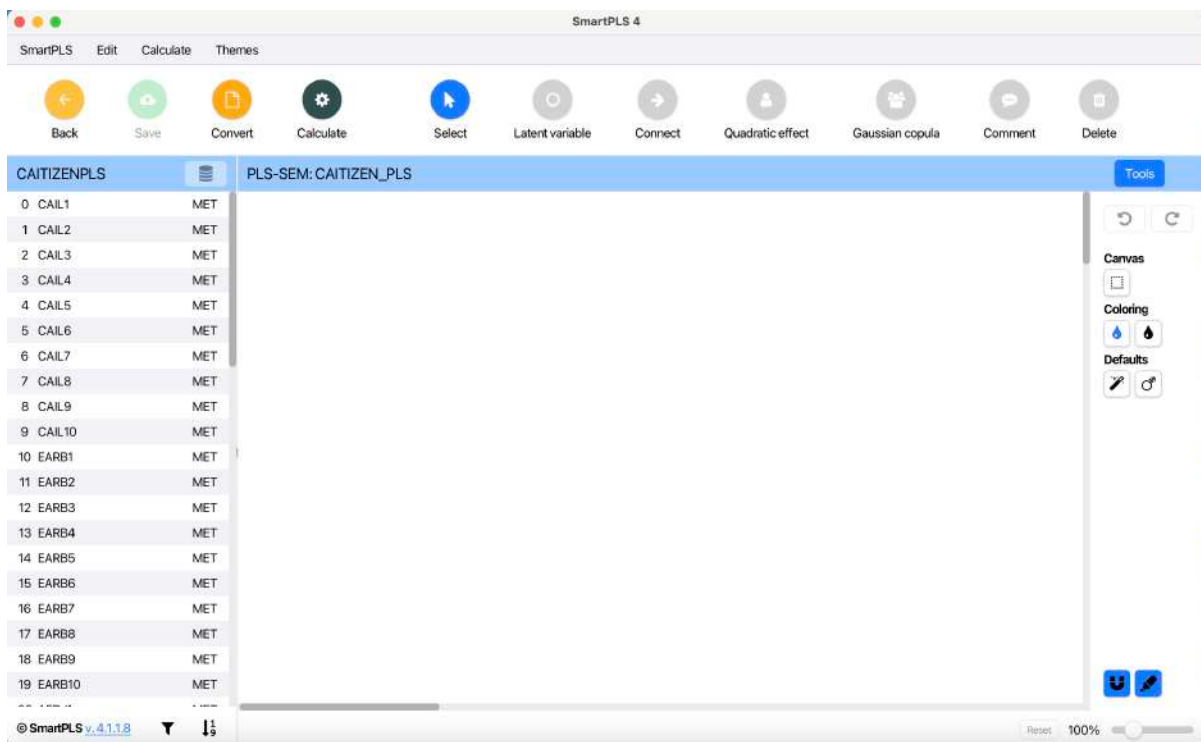
Juan Mejía Tréjo

Después de presionar **SAVE**, **SmartPLS4** abre el área gráfica de trabajo. En esta pantalla se observa el archivo de datos importado en el panel izquierdo y, al centro, el lienzo donde se crearán los constructos. La lista de indicadores aparece con sus nombres abreviados, por ejemplo **CAIL1**, **CAIL2**, **EARB1**, **AFDJ1**, **HAIC1**, **MTPP1** y **CAITIZEN1**. Esta pantalla confirma que el modelo ya está creado, pero aún no contiene variables latentes ni relaciones. Por ello, representa el punto de partida para construir visualmente el modelo **CAITIZEN**.

Creación de constructos latentes

La **Figura 3.4** muestra el área gráfica vacía y el listado de indicadores disponibles. Esta etapa es importante porque permite comprobar que **SmartPLS4** reconoce correctamente los campos importados desde el archivo **CSV**. A partir de esta pantalla se deben crear los constructos latentes. Para ello, se selecciona la herramienta **Latent variable** y se coloca cada constructo en el lienzo. Posteriormente, se asignan los indicadores correspondientes. La correcta identificación de los ítems en el panel izquierdo evita errores de asignación y facilita que cada variable latente quede vinculada con sus diez indicadores.

Figura 3.4. Área gráfica inicial para construir el modelo CAITIZEN en SmartPLS4



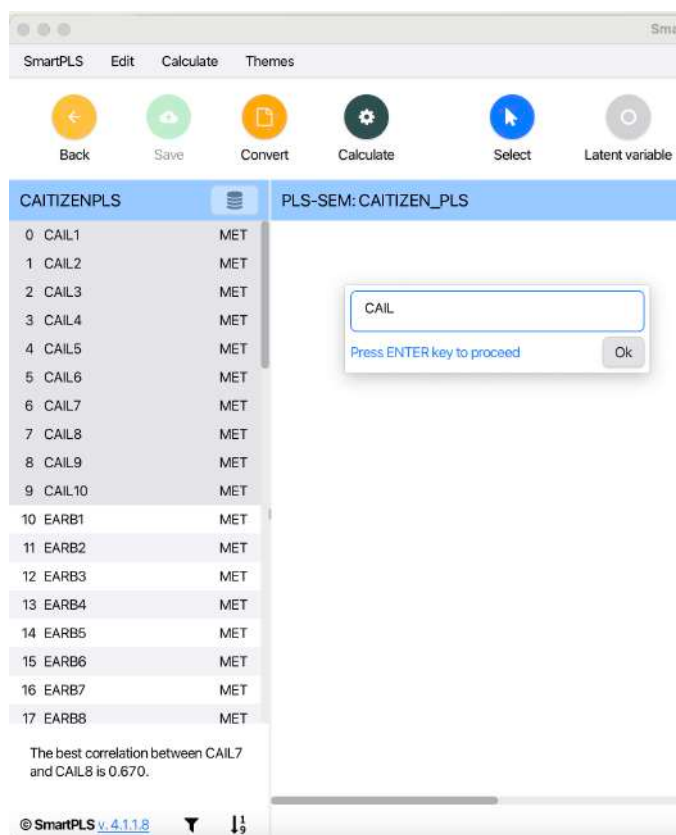
Fuente: Elaboración propia con **SmartPLS4**.

El primer constructo puede crearse seleccionando **Latent variable** y escribiendo el nombre correspondiente, por ejemplo **CAIL**. Este constructo representa **Critical Artificial Intelligence Literacy** y debe recibir los indicadores **CAIL1** a **CAIL10**. La creación del primer constructo permite explicar el procedimiento que después se repetirá con **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. En esta fase todavía no se establecen relaciones estructurales; primero se crean los constructos y se asignan sus indicadores para completar el modelo de medición reflectivo.

Asignación y alineación de indicadores

La **Figura 3.5** muestra el momento en que se nombra el constructo **CAIL** dentro del área gráfica. Esta captura es útil porque ilustra cómo **SmartPLS4** solicita el nombre de la variable latente antes de incorporarla al modelo. El mismo procedimiento se sigue para los demás constructos del modelo **CAITIZEN**. Debe cuidarse que los nombres sean breves, consistentes y sin errores tipográficos. Si un constructo queda mal nombrado, puede corregirse posteriormente desde sus propiedades. Una denominación clara facilita la interpretación del modelo y evita confusión durante la revisión de resultados.

Figura 3.5. Creación del constructo CAIL en el área gráfica de SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

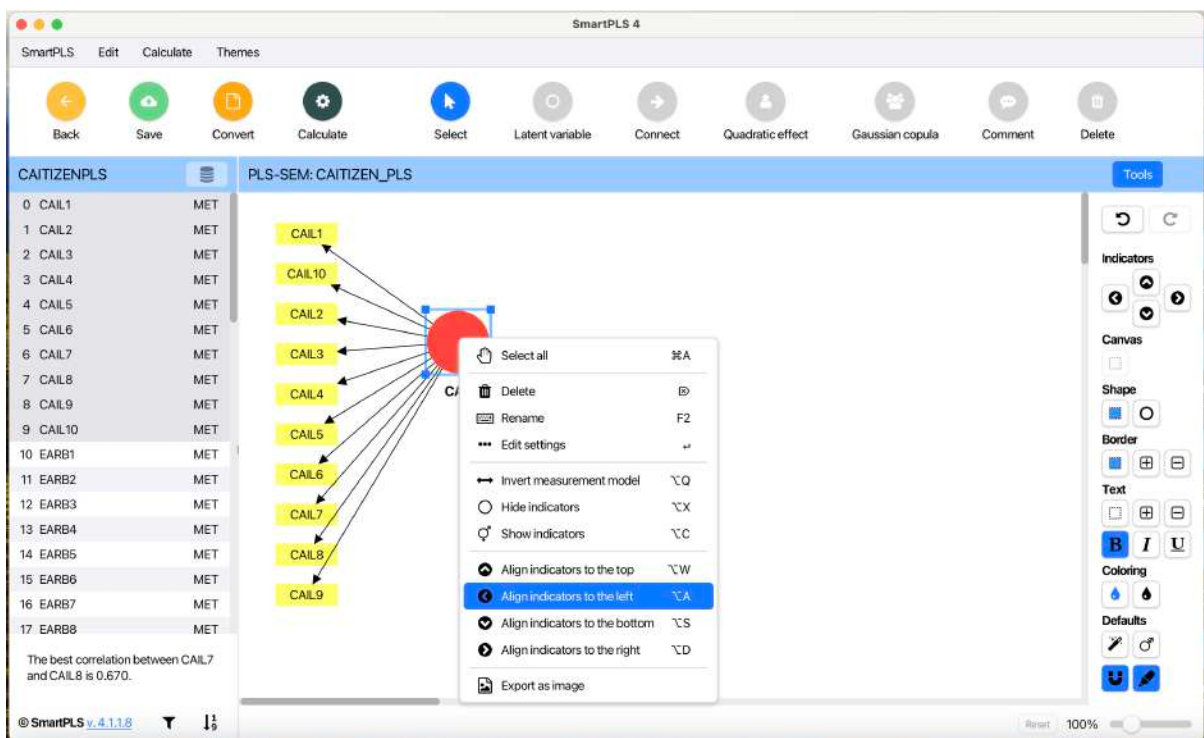
Juan Mejía Tréjo

Una vez creado el constructo, se le asignan sus indicadores y se organiza su presentación visual. **SmartPLS4** permite alinear los indicadores hacia arriba, abajo, izquierda o derecha. Esta organización no modifica los resultados estadísticos, pero mejora la legibilidad del modelo. En el caso de **CAIL**, los indicadores pueden alinearse a la izquierda para aprovechar mejor el espacio del lienzo. Esta operación también facilita revisar que estén incorporados los diez indicadores correctos. Un modelo visualmente ordenado permite detectar con mayor rapidez errores de asignación, duplicación de ítems o indicadores faltantes. Ver **Figura 3.6**.

Organización del modelo de medición reflectivo

La **Figura 3.6** muestra la alineación de los indicadores del constructo **CAIL** mediante el menú contextual de **SmartPLS4**. Esta herramienta permite organizar el modelo de medición de forma más clara. En esta pantalla se observa que los indicadores **CAIL1** a **CAIL10** se encuentran vinculados al constructo y pueden alinearse hacia la izquierda. Esta función debe utilizarse con todos los constructos cuando sea necesario. Aunque se trata de una acción visual, tiene importancia práctica porque facilita la revisión del modelo antes de ejecutar el algoritmo **PLS-SEM** y reduce el riesgo de interpretar un modelo desordenado.

Figura 3.6. Alineación de indicadores del constructo CAIL en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Después de crear y organizar los constructos, se debe construir el modelo completo. El modelo **CAITIZEN** está integrado por seis constructos reflectivos: **CAIL**, **EAR**,

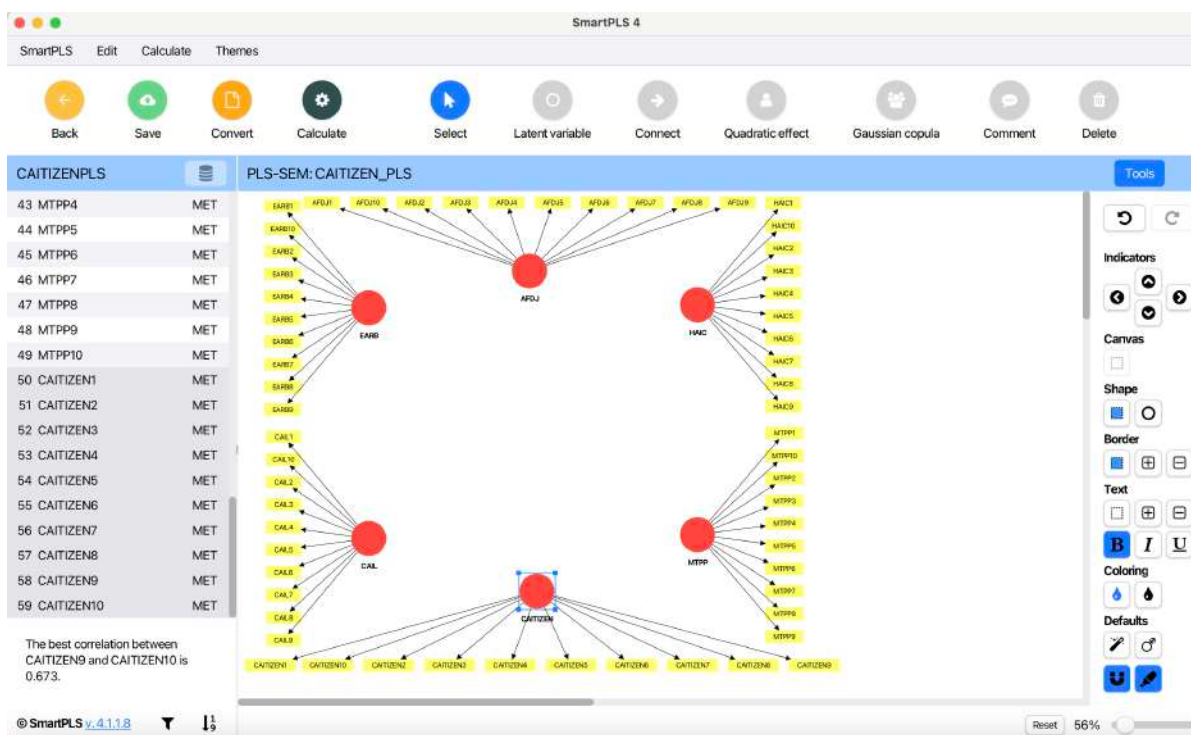
Juan Mejía Trejo

AFDJ, HAIC, MTPP y CAITIZEN. Los cinco primeros funcionan como constructos predictores, mientras que **CAITIZEN** representa la variable dependiente. Cada constructo debe tener asignados diez indicadores. En esta etapa, todos los indicadores deben aparecer conectados a su variable latente correspondiente. Si algún indicador permanece en el panel izquierdo sin asignar, debe incorporarse antes de continuar con las relaciones estructurales.

Revisión de propiedades de los constructos

La **Figura 3.7** muestra el modelo con los seis constructos ya creados y con sus indicadores asignados. Aunque todavía puede requerir ajustes visuales, esta captura permite verificar la estructura general del modelo de medición. Se observa que cada constructo tiene sus indicadores asociados mediante flechas de medición. También se aprecia la distribución espacial de las variables latentes en el lienzo. Esta organización es necesaria antes de establecer relaciones estructurales, porque primero debe completarse el modelo de medición. El modelo gráfico debe ser claro, ordenado y coherente con la estructura conceptual de **CAITIZEN**.

Figura 3.7. Organización inicial de constructos e indicadores del modelo CAITIZEN



Fuente: Elaboración propia con **SmartPLS4**.

SmartPLS4 permite revisar y editar las propiedades de cada constructo. Al hacer doble clic sobre una variable latente, se abre una ventana donde se puede modificar el nombre visible, revisar el tipo de medición y configurar la alineación de indicadores.

Juan Mejía Trejo

Esta función es especialmente útil cuando un constructo aparece con un nombre operativo diferente al esperado. Por ejemplo, si aparece **EARB**, puede ajustarse la leyenda visible a **EAR** en el campo **Caption in report**. También debe verificarse que el modelo de medición permanezca configurado como **reflective**. Ver **Figura 3.8**

Establecimiento de relaciones estructurales

La **Figura 3.8** muestra la ventana de propiedades del constructo, donde puede corregirse la leyenda del reporte y confirmar que el modelo de medición sea reflectivo. Esta captura es relevante porque permite explicar que no basta con dibujar constructos; también es necesario revisar su configuración. En este caso, si el constructo aparece como **EARB**, debe ajustarse a **EAR** para mantener consistencia con el modelo conceptual. Además, el campo **Measurement model** debe permanecer como **reflective**, ya que los indicadores representan manifestaciones observables de la variable latente y no componentes formativos.

Figura 3.8. Configuración de propiedades del constructo EAR en SmartPLS4

The screenshot shows the 'Latent variable' settings window in SmartPLS4. The construct name is 'EARB'. Under the 'Settings' tab, the 'Caption in report' is 'EAR'. The 'Measurement model' is 'reflective'. The 'Weighting mode' is 'Automatic', 'Reliability' is '1', 'Indicator alignment' is 'left', 'Indicator sort order' is 'Alphabetic order', 'Margin' is '100', and 'Folded' is 'no'. At the bottom are buttons for 'Previous', 'Next', 'Cancel', and 'Apply'.

Fuente: Elaboración propia con **SmartPLS4**.

Una vez creados y revisados los constructos, se deben establecer las relaciones estructurales. Conforme al modelo teórico, **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** predicen directamente a **CAITIZEN**. Para dibujar estas relaciones, se utiliza la herramienta

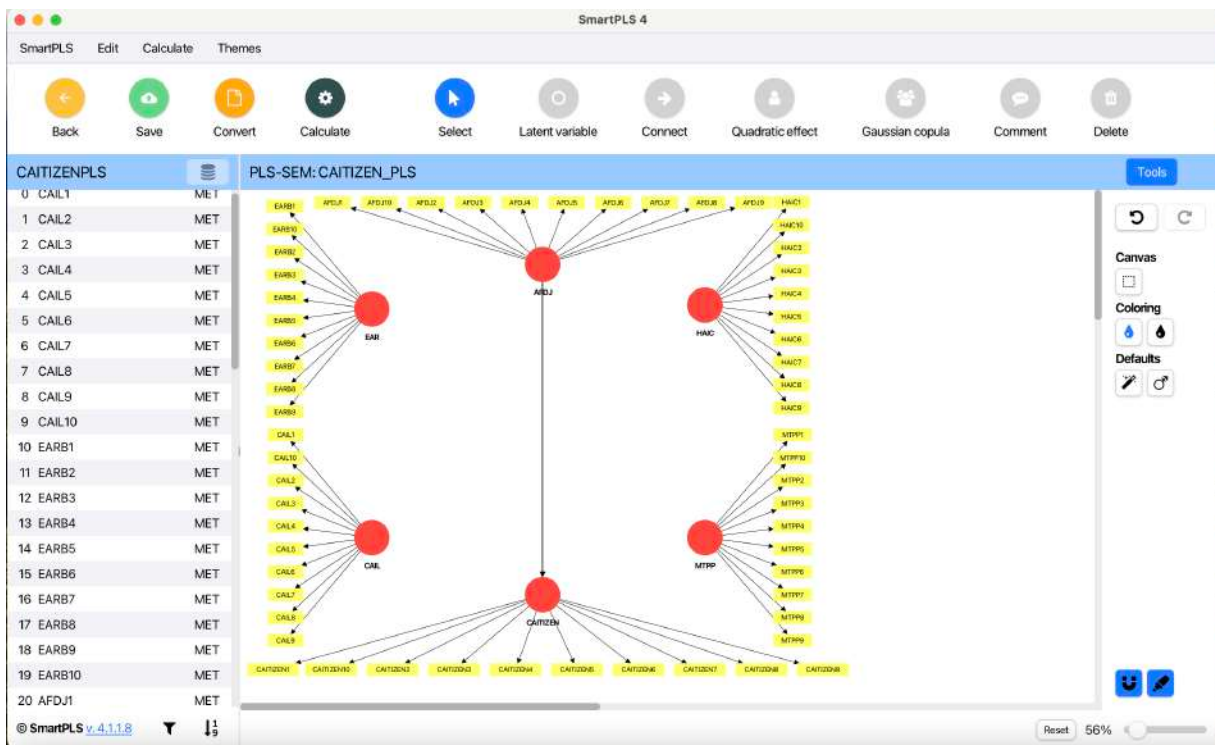
Juan Mejía Tréjo

Connect ubicada en la barra superior. Se selecciona primero el constructo de origen y después el constructo destino. La dirección de la flecha es fundamental, porque define qué variable actúa como predictor y cuál como dependiente. En este modelo, todas las flechas estructurales deben dirigirse hacia **CAITIZEN**. Ver **Figura 3.9**

Verificación final del modelo CAITIZEN

La **Figura 3.9** presenta el modelo **CAITIZEN** con las relaciones estructurales ya incorporadas. Se observa que los constructos predictores están conectados hacia **CAITIZEN**, lo cual representa el planteamiento hipotético del modelo. En esta etapa todavía no se han calculado coeficientes, cargas externas ni niveles de significancia. La imagen solo documenta la estructura gráfica previa a la estimación. Es importante revisar que las flechas estén correctamente orientadas y que no existan relaciones accidentales entre constructos predictores. Cualquier error en esta fase afectaría los resultados posteriores del algoritmo **PLS-SEM**.

Figura 3.9. Modelo CAITIZEN con relaciones estructurales preliminares en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Finalmente, se organiza el modelo para dejarlo listo antes de ejecutar el análisis. La variable dependiente **CAITIZEN** puede colocarse al centro inferior o en una posición que facilite la convergencia visual de las flechas estructurales. Los constructos predictores pueden distribuirse alrededor: **EAR** y **CAIL** a la izquierda, **AFDJ** arriba,

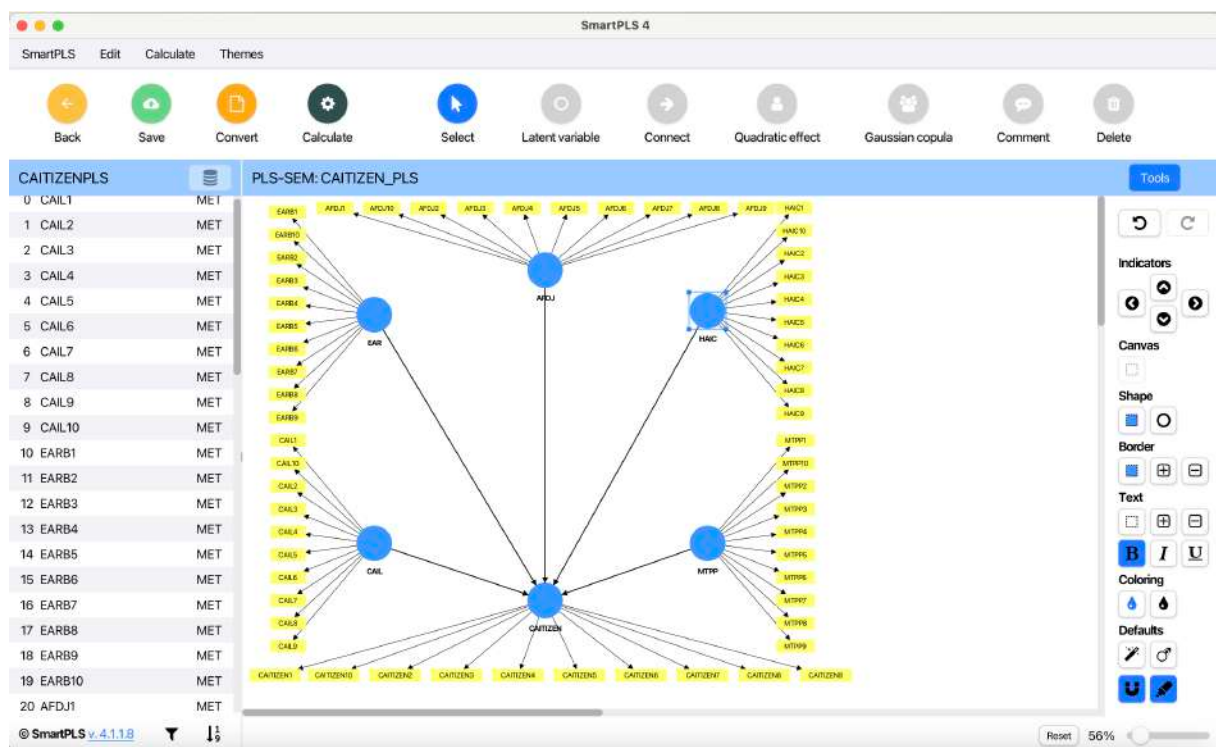
Juan Mejía Tréjo

HAIC y **MTPP** a la derecha. Esta organización no altera los resultados estadísticos, pero mejora la claridad del diagrama. El modelo debe guardarse antes de ejecutar el algoritmo **PLS-SEM** para asegurar que la estructura final quede registrada.

Cierre metodológico de la construcción del modelo

La **Figura 3.10** muestra el modelo **CAITIZEN** organizado y preparado para el análisis. En esta versión se observan los constructos, sus indicadores y las relaciones estructurales dirigidas hacia **CAITIZEN**. Si los constructos aparecen en color azul, esto indica que están seleccionados o activos dentro del área de trabajo; no representa un error. Antes de continuar, se debe comprobar que ningún indicador haya quedado sin asignar, que todos los constructos estén configurados como reflectivos y que las relaciones estructurales coincidan con el modelo teórico. Con ello, el modelo queda listo para la estimación **PLS-SEM**.

Figura 3.10. Modelo CAITIZEN organizado con relaciones estructurales finales en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

En síntesis, la construcción del modelo **CAITIZEN** en **SmartPLS4** integra tres tareas principales: crear el modelo **PLS-SEM**, asignar indicadores a constructos reflectivos y establecer relaciones estructurales entre predictores y variable dependiente. En esta

Juan Mejía Trejo

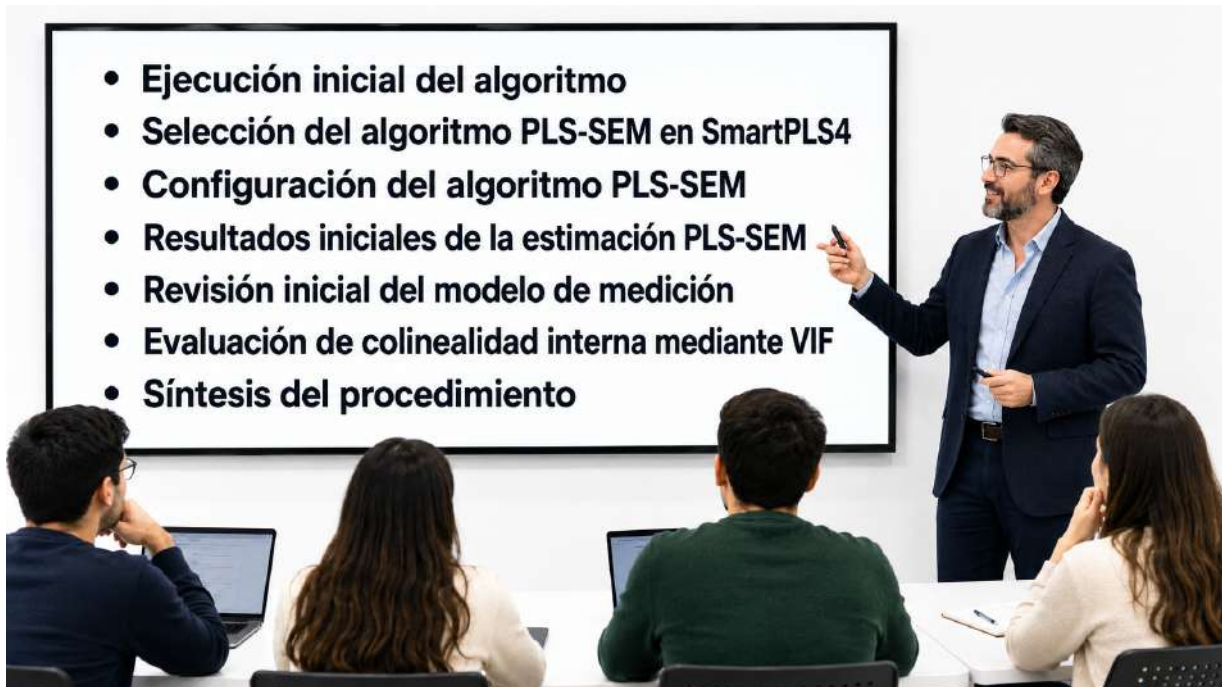
etapa no se interpretan resultados estadísticos, porque aún no se ha ejecutado el algoritmo **PLS-SEM**. La finalidad es dejar el modelo correctamente construido, ordenado y guardado. Una vez completada esta fase, se puede avanzar a la ejecución del algoritmo, la revisión del modelo de medición, la evaluación del modelo estructural y el procedimiento de **Bootstrapping**.

Síntesis del procedimiento

- Abrir **SmartPLS4**.
- Verificar que el **Workspace** activo sea **CAITIZEN**.
- Ingresar al proyecto **CAITIZEN**.
- Confirmar que el archivo de datos **CAITIZENPLS [511]** aparezca correctamente importado.
- Seleccionar la opción **Create model**. Ver **Figura 3.1**.
- En la ventana **New model**, verificar que el campo Project muestre **CAITIZEN** como se aprecia en la **Figura 3.2**.
- En Model type, seleccionar **PLS-SEM**.
- En Model name, escribir **CAITIZEN** o, si se conserva la captura, aclarar que **CAITIZEN_PLS** es un nombre operativo. Ver **Figura 3.3**.
- Presionar **SAVE** para abrir el área gráfica. Ver **Figura 3.4**.
- Crear el constructo **CAIL**. Ver **Figura 3.5**.
- Asignar los indicadores **CAIL1 a CAIL10** al constructo **CAIL**.
- Alinear los indicadores para mejorar la visualización. Ver **Figura 3.6**.
- Crear los constructos **EAR, AFDJ, HAIC, MTPP y CAITIZEN**.
- Asignar diez indicadores a cada constructo. Ver **Figura 3.7**.
- Revisar las propiedades de cada constructo y confirmar que el modelo de medición sea reflectivo. Ver **Figura 3.8**.
- Cambiar de **Select a Connect** en la barra superior.
- Dibujar las flechas **CAIL → CAITIZEN, EAR → CAITIZEN, AFDJ → CAITIZEN, HAIC → CAITIZEN y MTPP → CAITIZEN**. Ver **Figura 3.9**.
- Organizar visualmente el modelo completo.
- Confirmar que ningún indicador haya quedado sin asignar.
- Guardar el modelo dentro del proyecto **CAITIZEN**. Ver **Figura 3.10**.

Con estos pasos, el modelo **CAITIZEN** queda construido en **SmartPLS4** y preparado para ejecutar el algoritmo **PLS-SEM**, evaluar el modelo de medición, revisar el modelo estructural y aplicar **Bootstrapping**.

CAPÍTULO 4. PLS-SEM Algorithm y colinealidad



El **Capítulo 4. PLS-SEM Algorithm y colinealidad** en **SmartPLS4** representa el primer momento de estimación estadística del modelo **CAITIZEN**, una vez que la base de datos fue importada, los constructos fueron creados y las relaciones estructurales quedaron definidas. Su importancia radica en que permite pasar del modelo gráfico al **primer diagnóstico cuantitativo formal**, sin asumir todavía que las hipótesis han sido comprobadas. En esta etapa, el **algoritmo PLS-SEM** genera resultados preliminares sobre indicadores, constructos y relaciones, lo que permite verificar si el modelo puede procesarse adecuadamente y si cuenta con condiciones iniciales para avanzar hacia evaluaciones más profundas del modelo de medición y del modelo estructural (Hair et al., 2022).

La ejecución del algoritmo **PLS-SEM** cumple una función metodológica de **transición entre la construcción visual del modelo y su evaluación estadística formal**. En el caso **CAITIZEN**, esta primera estimación permite observar cómo los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTTP** se relacionan inicialmente con la variable dependiente **CAITIZEN**. Sin embargo, los coeficientes obtenidos en esta fase no deben interpretarse todavía como evidencia inferencial definitiva, ya que la significancia estadística se evaluará posteriormente mediante Bootstrapping. Por ello, el capítulo enfatiza que el algoritmo inicial sirve para diagnosticar, revisar y preparar el modelo antes de contrastar hipótesis (Chin, 1998).

Juan Mejía Tréjo

Antes de ejecutar el algoritmo, resulta indispensable confirmar que el modelo esté correctamente configurado. Cada constructo debe tener asignados sus diez indicadores, el tipo de medición debe mantenerse como **reflectivo** y las flechas estructurales deben dirigirse desde los constructos predictores hacia **CAITIZEN**. Esta verificación previa evita problemas de estimación, errores en la asignación de indicadores y relaciones estructurales incorrectas. En modelos reflectivos, los indicadores representan manifestaciones observables del constructo latente; por tanto, una asignación equivocada puede afectar la confiabilidad, la validez convergente y la validez discriminante del modelo (Jarvis et al., 2003).

La configuración del algoritmo en **SmartPLS4** también requiere atención metodológica. El capítulo recomienda conservar los parámetros estándar: **Weighting scheme = Path, Type of results = Standardized e Initial weight = Default**. Esta decisión facilita una estimación inicial clara, comparable y compatible con reportes convencionales de **PLS-SEM**. La estandarización de resultados permite interpretar de forma más directa las cargas externas, los coeficientes de trayectoria y los valores preliminares de **R²**. En consecuencia, mantener la configuración predeterminada resulta adecuado cuando no existe una razón teórica o estadística específica para modificar los parámetros del algoritmo (SmartPLS GmbH, s. f.-c).

Uno de los aportes centrales del capítulo es la revisión de la **colinealidad interna** mediante el **Variance Inflation Factor**, conocido como **VIF**. En el modelo **CAITIZEN**, el interés principal se concentra en el **Inner model VIF**, porque permite evaluar si los constructos predictores presentan solapamiento excesivo al explicar la variable dependiente. Aunque el **Outer model VIF** puede revisarse como diagnóstico complementario, su relevancia es mayor en modelos formativos. En modelos reflectivos, como **CAITIZEN**, la atención metodológica debe centrarse en la colinealidad entre constructos predictores, ya que esta puede distorsionar la interpretación individual de los coeficientes estructurales (Kock & Lynn, 2012).

Los resultados **VIF** reportados para **CAITIZEN** muestran valores aceptables: **AFDJ** → **CAITIZEN = 3.394**, **CAIL** → **CAITIZEN = 2.776**, **EAR** → **CAITIZEN = 3.205**, **HAIC** → **CAITIZEN = 2.666** y **MTPP** → **CAITIZEN = 2.604**. Todos se encuentran por debajo del umbral general de 5.0, por lo que no existe **colinealidad crítica** en el modelo estructural. No obstante, **AFDJ** supera ligeramente el criterio conservador de 3.3, lo que exige reportar el resultado con cautela, sin eliminar el constructo, debido a su relevancia conceptual dentro del modelo **CAITIZEN**. Esta decisión confirma que la depuración del modelo debe combinar criterio estadístico y justificación teórica (Hair et al., 2019).

En síntesis, el capítulo aporta una ruta operativa y metodológica para ejecutar el algoritmo **PLS-SEM**, obtener los resultados iniciales y revisar la colinealidad estructural antes de avanzar hacia la evaluación completa del modelo de medición. Su contribución principal consiste en mostrar que la estimación inicial no es un cierre analítico, sino una **fase de control de calidad estadística**. Con ello, el modelo **CAITIZEN** queda preparado para examinar cargas externas, confiabilidad interna, validez convergente, validez discriminante y significancia mediante Bootstrapping.

Juan Mejía Trejo

Esta secuencia fortalece la transparencia, la trazabilidad y la reproducibilidad del análisis **PLS-SEM** aplicado (Ringle et al., 2012).

Ejecución inicial del algoritmo

Una vez construido el modelo **CAITIZEN** en **SmartPLS4**, asignados los indicadores a cada constructo y establecidas las relaciones estructurales, el siguiente paso consiste en ejecutar el algoritmo **PLS-SEM**. Esta etapa permite obtener la primera estimación estadística del modelo y revisar si los indicadores y constructos cumplen con criterios mínimos de calidad. En esta fase se generan los resultados iniciales del modelo de medición y del modelo estructural, los cuales sirven como base para evaluar cargas externas, confiabilidad interna, validez convergente, colinealidad y capacidad explicativa. Por tanto, este procedimiento constituye el primer diagnóstico estadístico formal del modelo **CAITIZEN**. En este capítulo, la ejecución del algoritmo **PLS-SEM** se entiende como una etapa de **diagnóstico inicial** y no como la comprobación definitiva de las hipótesis. Su función es generar la primera estimación del modelo **CAITIZEN**, verificar que el modelo gráfico construido en el capítulo anterior pueda ser procesado correctamente y revisar criterios preliminares de calidad, especialmente la colinealidad entre constructos predictores. La interpretación completa de las cargas externas, la confiabilidad interna, la validez convergente y la validez discriminante se desarrollará en el capítulo siguiente. De esta manera, el Capítulo 4 cumple una función de transición metodológica: conecta la construcción visual del modelo con su evaluación estadística formal.

Función del algoritmo PLS-SEM en el modelo CAITIZEN

El algoritmo **PLS-SEM** estima simultáneamente las relaciones entre indicadores y constructos, así como las relaciones entre constructos latentes. En el caso del modelo **CAITIZEN**, permite analizar cómo **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** se relacionan con la variable dependiente **CAITIZEN**. Esta primera ejecución no debe interpretarse todavía como prueba definitiva de hipótesis, porque la significancia estadística de las relaciones estructurales se evaluará posteriormente mediante **Bootstrapping**. En consecuencia, el algoritmo **PLS-SEM** se utiliza primero para obtener estimaciones iniciales, revisar la calidad del modelo y preparar la evaluación inferencial posterior.

Verificación previa del modelo construido

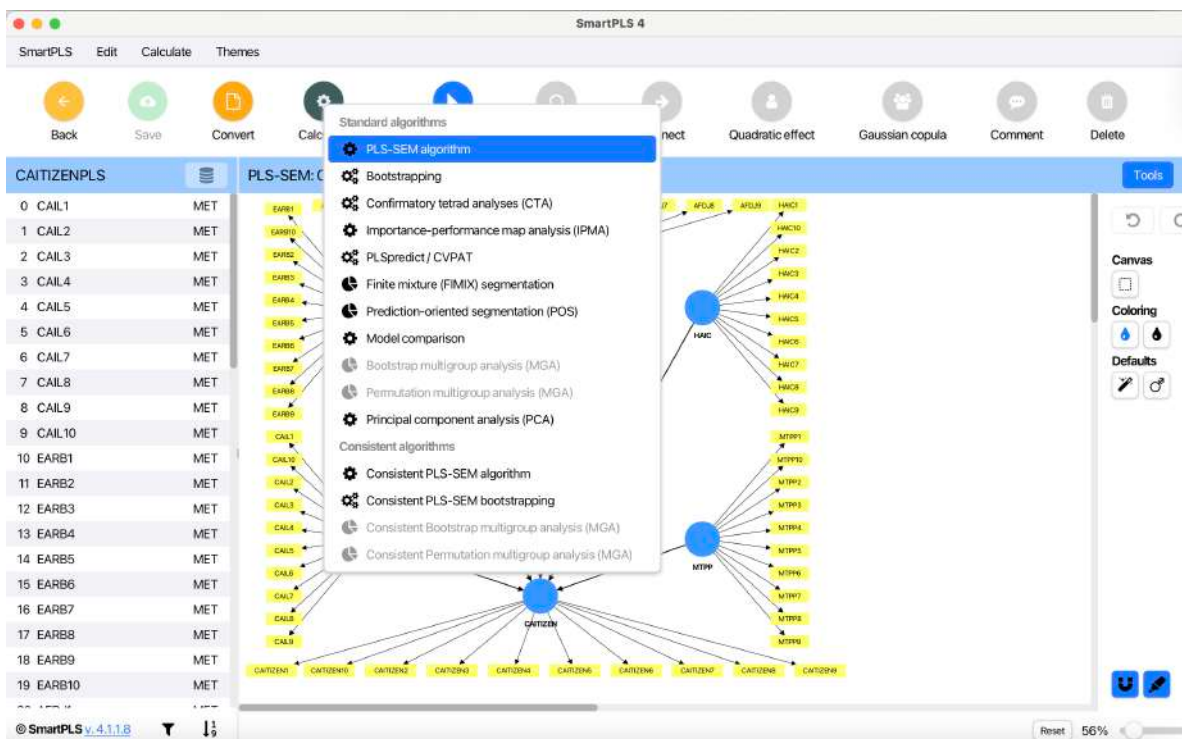
Antes de ejecutar el algoritmo, es necesario confirmar que el modelo esté correctamente construido. Cada constructo debe tener asignados sus diez indicadores, los constructos deben estar configurados como reflectivos y las flechas estructurales deben dirigirse desde **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** hacia **CAITIZEN**. Esta revisión previa evita errores de estimación, indicadores sin asignar o relaciones estructurales incorrectas. Si el modelo gráfico no representa adecuadamente el planteamiento teórico, los resultados generados por **SmartPLS4** pueden ser inconsistentes, incompletos o difíciles de interpretar en las etapas posteriores del análisis.

Juan Mejía Trejo

Selección del algoritmo PLS-SEM en SmartPLS4

Para iniciar la estimación, se selecciona la opción **Calculate** en la barra superior de **SmartPLS4** y posteriormente **PLS-SEM Algorithm**. Esta función ejecuta la primera estimación del modelo mediante mínimos cuadrados parciales. **SmartPLS4** utiliza el archivo de datos **CAITIZENPLS [511]** y calcula los valores iniciales del modelo. En esta etapa se obtienen cargas externas, coeficientes de trayectoria, valores de R^2 , criterios de confiabilidad, validez convergente y estadísticas de colinealidad, incluyendo el **Variance Inflation Factor (VIF)**. Por ello, esta ejecución es indispensable antes de realizar **Bootstrapping** o contrastar hipótesis. Ver **Figura 4.1**.

Figura 4.1. Selección de PLS-SEM Algorithm en SMARTPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Configuración del algoritmo PLS-SEM

Después de seleccionar **PLS-SEM Algorithm**, **SmartPLS4** muestra la ventana de configuración del algoritmo. Para este análisis se recomienda mantener la configuración predeterminada: **Weighting scheme = Path**, **Type of results = Standardized** e **Initial weight = Default**. Esta configuración es adecuada para obtener la estimación inicial del modelo **CAITIZEN** y los criterios de calidad asociados. Si no existe una justificación metodológica específica para modificar estos parámetros, conviene conservar la configuración estándar, ya que permite una estimación clara,

Juan Mejía Trejo

comparable y compatible con reportes **PLS-SEM** convencionales aplicados a modelos reflectivos. Ver **Figura 4.2**.

Figura 4.2. Configuración inicial del algoritmo PLS-SEM

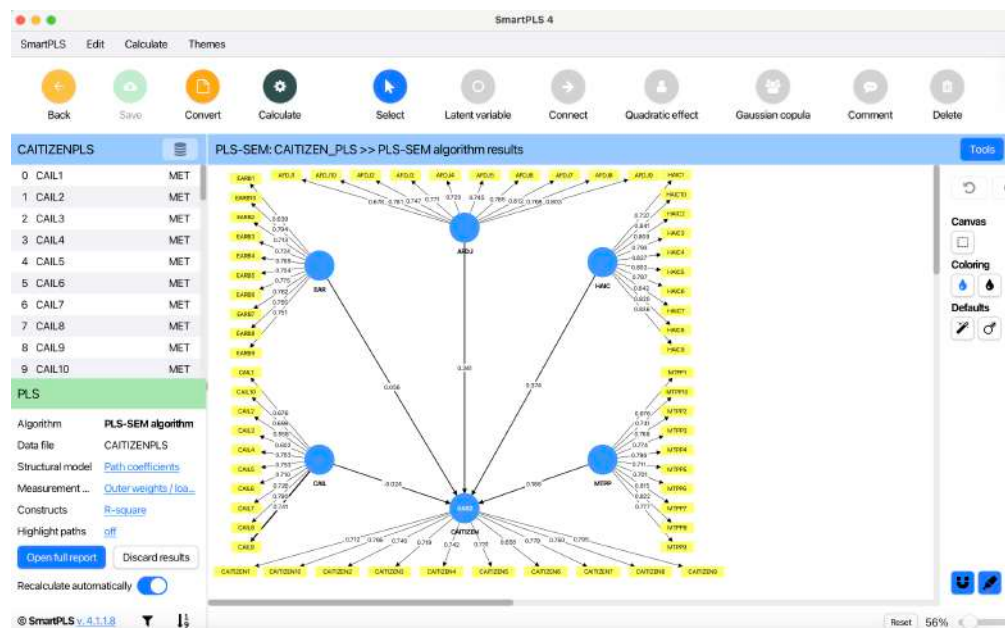


Fuente: Elaboración propia con **SmartPLS4**.

Resultados iniciales de la estimación PLS-SEM

Una vez revisada la configuración, se debe presionar **Start calculation**. **SmartPLS4** procesa el modelo y muestra los resultados iniciales sobre el diagrama. En esta pantalla se observan valores estimados en los indicadores y en las relaciones estructurales. También puede visualizarse el valor de R^2 del constructo dependiente **CAITIZEN**, lo cual permite observar de manera preliminar la capacidad explicativa del modelo. Estos resultados confirman que el modelo fue estimado correctamente y que ya es posible revisar los criterios de calidad en el reporte completo generado por el programa. Ver **Figura 4.3**.

Figura 4.3. Resultados iniciales PLS-SEM Algorithm para el modelo CAITIZEN



Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

Revisión inicial del modelo de medición

Las **cargas externas** indican qué tan bien cada indicador representa al constructo al que fue asignado. Como criterio general, se esperan valores iguales o superiores a **0.70**. La confiabilidad interna se evalúa mediante **Cronbach Alpha**, **rho_A** y **Composite Reliability**, cuyos valores deben ser superiores a **0.70**. La validez convergente se revisa mediante el **Average Variance Extracted (AVE)**, cuyo valor debe ser superior a **0.50**. Estos criterios permiten determinar si el modelo de medición presenta calidad suficiente para avanzar hacia la validez discriminante y la evaluación estructural.

Si algún indicador presenta una carga externa baja, **no debe eliminarse automáticamente**. Primero debe revisarse si su eliminación mejora la confiabilidad, el **AVE** o la validez discriminante del constructo. Además, debe considerarse si el ítem tiene importancia conceptual dentro del modelo **CAITIZEN**. La depuración de indicadores debe combinar criterio estadístico y justificación teórica. El objetivo no es eliminar ítems únicamente para mejorar cifras, sino conservar una medición válida, coherente y representativa de cada dimensión del modelo: alfabetización crítica, responsabilidad ética, justicia de datos, creatividad humano-IA y metacognición.

Evaluación de colinealidad interna mediante VIF

Para evaluar la colinealidad en el modelo **CAITIZEN** se utiliza el **Variance Inflation Factor (VIF)**. En **SmartPLS4**, este criterio se consulta después de ejecutar el **PLS-SEM Algorithm**, mediante la opción **Open full report**. Dentro del reporte se selecciona **Quality criteria** y posteriormente **Collinearity statistics (VIF)**. **SmartPLS4** muestra dos opciones relevantes: **Outer model** e **Inner model**. El **Outer model VIF** permite revisar la colinealidad entre indicadores de un mismo constructo, mientras que el **Inner model VIF** permite revisar la colinealidad entre constructos predictores dentro del modelo estructural.

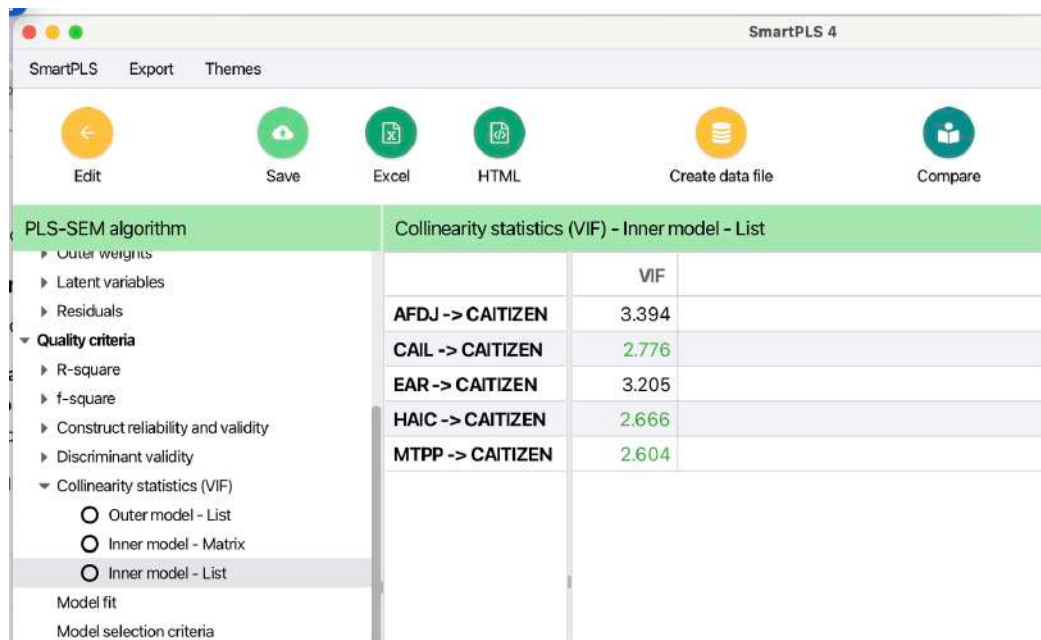
Diferencia entre Outer model VIF e Inner model VIF

En modelos con **constructos formativos**, el **Outer model VIF** adquiere especial importancia porque los indicadores actúan como componentes que forman el constructo. En cambio, cuando los constructos son reflectivos, como ocurre en el modelo **CAITIZEN**, los indicadores son manifestaciones observables de una variable latente. Por ello, el **Outer model VIF** puede revisarse como diagnóstico complementario para detectar posibles solapamientos excesivos entre indicadores, pero no constituye el criterio principal para evaluar la colinealidad estructural del modelo. En este caso, la atención metodológica debe centrarse en los valores VIF internos.

Revisión del Inner model VIF en CAITIZEN

En el modelo **CAITIZEN**, el interés principal es revisar si los constructos predictores **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** presentan colinealidad problemática al explicar la variable dependiente **CAITIZEN**. Para ello, debe seleccionarse la opción **Inner model – List**, dentro de **Collinearity statistics (VIF)**. Esta sección muestra los valores **VIF** correspondientes a las rutas estructurales dirigidas hacia **CAITIZEN**. Si los valores se mantienen por debajo del umbral general de **5.0**, se considera que no existe colinealidad crítica entre los predictores del modelo estructural. Ver **Figura 4.4**.

Figura 4.4. Consulta de Collinearity Statistics (VIF) en Inner model-list



Collinearity statistics (VIF) - Inner model - List	
	VIF
AFDJ -> CAITIZEN	3.394
CAIL -> CAITIZEN	2.776
EAR -> CAITIZEN	3.205
HAIC -> CAITIZEN	2.666
MTPP -> CAITIZEN	2.604

Fuente: Elaboración propia con **SmartPLS4**.

Criterios de interpretación del VIF

En términos metodológicos, valores **VIF** inferiores a **5.0** indican ausencia de colinealidad crítica. Un criterio más conservador recomienda valores inferiores a **3.3**. Por tanto, cuando los valores **VIF** se encuentran por debajo de estos límites, se considera que los constructos predictores pueden interpretarse sin problemas graves de redundancia estadística. En ciencias sociales no se exige independencia perfecta entre constructos, porque las variables suelen estar conceptualmente relacionadas.

Resultados VIF del modelo CAITIZEN

Sin embargo, sí se requiere que no exista solapamiento excesivo que distorsione los coeficientes estructurales o dificulte la interpretación individual de cada predictor. En el modelo **CAITIZEN**, los valores **VIF** internos fueron los siguientes: **AFDJ** →

Juan Mejía Trejo

CAITIZEN = 3.394, CAIL → CAITIZEN = 2.776, EAR → CAITIZEN = 3.205, HAIC → CAITIZEN = 2.666 y MTPP → CAITIZEN = 2.604. Estos resultados muestran que todos los valores se encuentran por debajo del umbral general de **5.0**, por lo que no existe colinealidad crítica en el modelo estructural. En consecuencia, las relaciones entre los predictores y la variable dependiente pueden seguir evaluándose dentro del modelo **CAITIZEN**.

Decisión metodológica sobre colinealidad

Sin embargo, el valor de **AFDJ → CAITIZEN = 3.394** supera ligeramente el criterio conservador de **3.3**. **Esto no invalida el modelo ni obliga a eliminar el constructo, pero indica que AFDJ mantiene cierta cercanía conceptual con otros predictores**, especialmente aquellos relacionados con ética, responsabilidad, transparencia y justicia. Por ello, este resultado debe reportarse con cautela. La decisión metodológica recomendada es conservar el constructo, ya que su valor VIF sigue siendo aceptable bajo el umbral general de **5.0** y mantiene relevancia teórica dentro del modelo.

Síntesis del procedimiento

- Ejecutar el modelo **CAITIZEN** mediante la ruta **Calculate → PLS-SEM Algorithm**. Esta acción inicia la primera estimación estadística del modelo en **SmartPLS4**. Ver **Figura 4.1**.
- Configurar el algoritmo con los parámetros recomendados para la estimación inicial: **Weighting scheme = Path, Type of results = Standardized e Initial weight = Default**. Ver **Figura 4.2**.
- Presionar **Start calculation** para obtener los resultados iniciales del modelo. En esta etapa se visualizan valores estimados en indicadores, relaciones estructurales y el valor preliminar de R^2 del constructo dependiente **CAITIZEN**. Ver **Figura 4.3**.
- Abrir el reporte completo mediante la opción **Open full report**.
- Ingresar a **Quality criteria**.
- Seleccionar **Collinearity statistics (VIF)**.
- Distinguir entre **Outer model VIF** e **Inner model VIF**.
- Reconocer que el **Outer model VIF** permite revisar colinealidad entre indicadores de un mismo constructo.
- Considerar que, en el modelo **CAITIZEN**, los constructos son reflectivos; por ello, el **Outer model VIF** funciona como diagnóstico complementario.
- Seleccionar **Inner model – List** para revisar la colinealidad entre constructos predictores del modelo estructural. Ver **Figura 4.4**.
- Identificar los valores **VIF** de las rutas dirigidas hacia **CAITIZEN**: **AFDJ → CAITIZEN, CAIL → CAITIZEN, EAR → CAITIZEN, HAIC → CAITIZEN y MTPP → CAITIZEN**.
- Comparar los valores obtenidos con el umbral general de **5.0** y con el criterio conservador de **3.3**.
- Confirmar que no existe colinealidad crítica cuando los valores VIF se mantienen por debajo de **5.0**.

- Reportar con cautela los valores cercanos o ligeramente superiores al criterio conservador de **3.3**.
- Conservar los constructos predictores cuando su valor VIF sea aceptable y exista justificación teórica para mantenerlos en el modelo.

Con este procedimiento, las **Figuras 4.1 a 4.4** documentan la secuencia operativa que va desde la ejecución del algoritmo **PLS-SEM** hasta la revisión del **Inner model VIF**. En el modelo **CAITIZEN**, los resultados obtenidos permiten concluir que no existe colínealidad crítica entre los constructos predictores y que el análisis puede continuar con la evaluación del modelo de medición reflectivo, particularmente cargas externas, confiabilidad interna y validez convergente.

CAPÍTULO 5. Evaluación del modelo de medición reflectivo



El **Capítulo 5. Evaluación del modelo de medición reflectivo**, constituye una **fase metodológica indispensable** antes de interpretar cualquier relación estructural del modelo. Su propósito central es comprobar si los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** están adecuadamente medidos por sus indicadores. En **PLS-SEM**, no es suficiente construir gráficamente un modelo ni ejecutar el algoritmo; primero debe verificarse que cada variable latente tenga una **base empírica confiable y válida**. Por ello, este capítulo aporta una revisión sistemática de **cargas externas**, **confiabilidad interna**, **validez convergente** y **validez discriminante**, como condición previa para evaluar hipótesis y efectos estructurales (Hair et al., 2021).

La naturaleza **reflectiva** del modelo **CAITIZEN** implica que los indicadores son **manifestaciones observables de un constructo latente**, no componentes independientes que lo forman. Así, los ítems **CAIL1–CAIL10** reflejan **alfabetización crítica en inteligencia artificial**; **EAR1–EAR10** reflejan **responsabilidad ética**; **AFDJ1–AFDJ10** reflejan **justicia de datos**; **HAIC1–HAIC10** reflejan **colaboración humano–IA**; **MTPP1–MTPP10** reflejan **transparencia metacognitiva**; y **CAITIZEN1–CAITIZEN10** reflejan **ciudadanía sostenible asistida por IA**. Esta especificación exige revisar si los indicadores comparten suficiente varianza con su constructo y si operan de manera coherente dentro de cada dimensión conceptual (Jarvis et al., 2003).

La primera evaluación corresponde a las **cargas externas**, porque estas indican qué tan bien cada indicador representa al constructo asignado. Como criterio general,

Juan Mejía Trejo

se esperan cargas iguales o superiores a **0.70**; las cargas entre **0.40 y 0.70** requieren revisión; y las menores a **0.40** suelen considerarse problemáticas. Sin embargo, la decisión de conservar o eliminar indicadores no debe ser automática. En el modelo **CAITIZEN**, un ítem con carga moderada puede conservarse si posee **relevancia conceptual** para explicar dimensiones como alfabetización crítica, ética, justicia algorítmica o ciudadanía asistida por **IA**. La depuración debe equilibrar **evidencia estadística y fundamento teórico** (Hair et al., 2019).

Después de revisar las cargas externas, el capítulo aborda la **confiabilidad interna** mediante **Cronbach Alpha, rho_A y confiabilidad compuesta**. Estos criterios permiten determinar si los indicadores de un mismo constructo funcionan de manera consistente como manifestaciones de una variable latente común. En el modelo **CAITIZEN**, los constructos superan el umbral recomendado de **0.70**, lo cual respalda su **consistencia interna**. No obstante, también se advierte que valores cercanos o superiores a **0.95** pueden sugerir redundancia entre indicadores. Por tanto, la confiabilidad debe interpretarse junto con las cargas externas y la validez convergente, no como criterio aislado (Gefen et al., 2000).

La **validez convergente** se evalúa mediante la **Average Variance Extracted o AVE**, cuyo umbral recomendado es **0.50**. Este criterio permite saber si el constructo explica al menos la mitad de la varianza de sus indicadores. En **CAITIZEN**, cinco constructos cumplen este criterio: **AFDJ, CAITIZEN, EAR, HAIC y MTPP**. El caso que requiere atención es **CAIL**, con **AVE = 0.497**, apenas por debajo del valor esperado. Esta diferencia mínima no invalida el constructo, pero sí justifica revisar sus indicadores antes de decidir cualquier eliminación. La decisión debe considerar si el ajuste mejora la medición sin debilitar la **cobertura conceptual de la alfabetización crítica en IA** (Fornell & Larcker, 1981).

La **validez discriminante** permite verificar si los constructos son **empíricamente distintos entre sí**. En este capítulo se revisan tres criterios complementarios: **Fornell-Larcker, Cross Loadings y HTMT**. El criterio Fornell-Larcker muestra algunos puntos de atención en relaciones como **AFDJ–EAR, CAIL–AFDJ y CAITIZEN–HAIC**, lo cual resulta comprensible porque las dimensiones del modelo pertenecen a un campo conceptual común. Sin embargo, las **Cross Loadings** muestran que los indicadores cargan principalmente en sus constructos asignados, y los valores **HTMT** se mantienen por debajo de **0.85**. En conjunto, estos resultados respaldan una **validez discriminante aceptable** (Henseler et al., 2015).

Finalmente, el capítulo enfatiza la importancia de conservar la **trazabilidad metodológica** mediante **copias del modelo original y exportación de reportes a Excel**. Cuando se evalúa eliminar un indicador, se recomienda duplicar el modelo, ejecutar nuevamente el algoritmo y comparar resultados antes y después de la depuración. Esta práctica permite **documentar decisiones, conservar evidencia y evitar modificaciones irreversibles** sobre la versión inicial. En síntesis, el capítulo demuestra que la evaluación del modelo de medición reflectivo no es una verificación mecánica de umbrales, sino una lectura integrada de **confiabilidad, validez, teoría y**

transparencia metodológica. Con ello, **CAITIZEN** queda preparado para avanzar hacia la evaluación del modelo estructural (Franke & Sarstedt, 2019).

Naturaleza reflectiva del modelo CAITIZEN

En el **Capítulo 1** se explicó que el análisis **PLS-SEM** distingue entre **modelo de medición** y **modelo estructural**. A partir de esa base, este capítulo se concentra exclusivamente en la evaluación del **modelo de medición reflectivo** del modelo **CAITIZEN**. Esta precisión es importante porque, antes de interpretar rutas estructurales, coeficientes, significancia estadística o hipótesis, se debe comprobar que cada constructo esté adecuadamente medido por sus indicadores. En otras palabras, el objetivo de esta etapa no es todavía determinar si **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** explican a **CAITIZEN**, sino verificar si los indicadores asignados a cada constructo reflejan de manera consistente y válida la variable latente correspondiente. Por ello, la evaluación del modelo de medición constituye una condición previa para cualquier interpretación posterior del modelo estructural. Si los constructos no están bien medidos, las relaciones entre ellos pierden solidez metodológica, aunque los coeficientes parezcan estadísticamente relevantes (Hair et al., 2021).

El modelo **CAITIZEN** se trabaja como un modelo integrado por **constructos reflectivos**. En este tipo de especificación, la variable latente explica las respuestas observadas en sus indicadores. Por tanto, los indicadores no se entienden como piezas independientes que forman el constructo, sino como manifestaciones empíricas de una misma dimensión subyacente. Por ejemplo, los ítems **CAIL1–CAIL10** reflejan la **alfabetización crítica en inteligencia artificial**; no la forman como componentes aislados. De manera equivalente, **EAR1–EAR10** reflejan la **conciencia ética y responsabilidad**; **AFDJ1–AFDJ10** reflejan la **conciencia de equidad y justicia de datos**; **HAIC1–HAIC10** reflejan la **colaboración creativa humano–IA**; **MTPP1–MTPP10** reflejan la **transparencia metacognitiva en prácticas de prompting**; y **CAITIZEN1–CAITIZEN10** reflejan la **ciudadanía sostenible asistida por IA**. Esta lógica obliga a revisar si cada indicador comparte suficiente varianza con su constructo y si el conjunto de indicadores funciona de manera coherente (Jarvis et al., 2003).

La especificación reflectiva tiene consecuencias directas sobre los criterios de evaluación. En primer lugar, se revisan las **cargas externas**, porque indican qué tan bien cada indicador representa al constructo al que fue asignado. Una carga externa alta sugiere que el ítem es una buena manifestación empírica de la variable latente. En segundo lugar, se revisa la **confiabilidad interna**, porque los indicadores de un mismo constructo reflectivo deben mostrar consistencia entre sí. Para ello se utilizan criterios como **Cronbach Alpha**, **rho_A** y **Composite Reliability**. En tercer lugar, se evalúa la **validez convergente** mediante la **AVE**, porque el constructo debe explicar una proporción suficiente de la varianza de sus indicadores. Finalmente, se revisa la **validez discriminante**, con el propósito de verificar que cada constructo sea empíricamente distinto de los demás. Estos criterios no son independientes; en conjunto permiten determinar si el modelo de medición reflectivo tiene calidad suficiente para avanzar hacia el análisis estructural (Hair et al., 2019).

Juan Mejía Trejo

En el caso de **CAIL**, la evaluación reflectiva busca verificar si los diez ítems asignados representan adecuadamente la alfabetización crítica en inteligencia artificial. Si los indicadores presentan cargas externas suficientes, puede afirmarse que expresan una misma orientación latente relacionada con comprensión crítica, reconocimiento de sesgos, identificación de límites, verificación de información y uso responsable de la IA. En cambio, si algún indicador muestra una carga externa baja, debe revisarse si realmente representa el contenido del constructo o si está capturando una dimensión distinta. La misma lógica aplica a **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. En todos los casos, el análisis no debe reducirse a eliminar indicadores para mejorar cifras. Una decisión metodológica adecuada debe equilibrar evidencia estadística y pertinencia conceptual. Un ítem puede tener una carga menor a la esperada y, aun así, conservar valor teórico dentro del constructo; por ello, la depuración debe realizarse con cautela (Gefen et al., 2000).

La naturaleza reflectiva también permite distinguir la evaluación de indicadores de la evaluación de relaciones estructurales. En esta etapa todavía no interesa saber cuál predictor tiene mayor efecto sobre **CAITIZEN**, ni si una hipótesis es estadísticamente significativa. Lo que se revisa es si cada bloque de indicadores funciona adecuadamente como medida de su constructo. Por ejemplo, antes de interpretar la ruta **AFDJ** → **CAITIZEN**, debe verificarse que **AFDJ** esté correctamente medido por sus indicadores. Antes de analizar la ruta **HAIC** → **CAITIZEN**, debe comprobarse que **HAIC1–HAIC10** reflejen de manera coherente la colaboración creativa humano–IA. Esta secuencia protege la interpretación del modelo, porque evita aceptar o rechazar hipótesis sobre constructos que aún no han demostrado calidad de medición. En consecuencia, la evaluación reflectiva funciona como una etapa de validación previa, necesaria para que el modelo estructural pueda interpretarse con rigor (Chin, 1998).

En síntesis, la naturaleza reflectiva del modelo **CAITIZEN** significa que sus indicadores son manifestaciones observables de constructos latentes. Por ello, la evaluación del modelo de medición debe centrarse en cuatro preguntas metodológicas: si los indicadores cargan adecuadamente en su constructo; si los indicadores de cada constructo muestran consistencia interna; si el constructo explica suficiente varianza de sus indicadores; y si cada constructo es distinto de los demás. Estas preguntas orientan las secciones siguientes del capítulo: primero se revisarán las **cargas externas**; después, la **confiabilidad interna**; posteriormente, la **validez convergente mediante AVE**; y finalmente, la **validez discriminante**. Solo después de cumplir esta revisión será metodológicamente válido avanzar hacia la evaluación del modelo estructural y la contrastación de hipótesis. De esta manera, el Capítulo 5 no repite la explicación general del **Capítulo 1**, sino que aplica esa base conceptual a la evaluación formal del modelo de medición reflectivo (Hair et al., 2021).

Evaluación de cargas externas

La primera revisión del modelo de medición reflectivo corresponde a las **cargas externas** de los indicadores. En un modelo reflectivo, cada ítem debe representar adecuadamente al constructo latente al que fue asignado. Por ello, se revisan las

Juan Mejía Trejo

cargas de **CAIL1–CAIL10**, **EAR1–EAR10**, **AFDJ1–AFDJ10**, **HAIC1–HAIC10**, **MTPP1–MTPP10** y **CAITIZEN1–CAITIZEN10**. Como criterio general, una carga externa igual o superior a **0.70** se considera adecuada, porque indica que el indicador comparte suficiente varianza con su constructo. Ver **Tabla 5.1**.

Tabla 5.1. Criterios de evaluación del modelo de medición reflectivo.

Criterio	Umbral	
Fiabilidad individual del indicador	Carga externa: $\lambda \geq 0.707$	
Fiabilidad del constructo (consistencia interna)	- Cronbach Alpha ≥ 0.70 - Fiabilidad compuesta (ρ_c) ≥ 0.70 - Dijkstra-Henseler rho_A ≥ 0.70	- IC 95% bootstrap 2.5% > 0.70 97.5% < 0.95
Validez convergente	Varianza extraída media (AVE) ≥ 0.50	
Validez discriminante	- Cargas cruzadas - Fornell-Larcker: $\sqrt{AVE} >$ correlación con el resto de constructos - HTMT ratio ≤ 0.85; ≤ 0.90 como criterio aceptable - HTMT inference: ¿0.90 dentro del IC 95%? → No	

Fuente: Elaboración propia con base en Hair et al. (2021), Fornell y Larcker (1981), Henseler et al. (2015) y Cepeda-Carrion et al. (2017).

Las cargas entre **0.40 y 0.70** requieren revisión antes de tomar una decisión, mientras que las cargas menores a **0.40** suelen considerarse problemáticas. En el modelo **CAITIZEN**, la evaluación no debe limitarse a observar los colores del reporte de **SmartPLS4**; debe verificarse qué indicadores presentan cargas aceptables, cuáles requieren revisión y cuáles podrían afectar la validez del constructo. Esta revisión es necesaria antes de interpretar confiabilidad, **AVE** o rutas estructurales (Hair et al., 2021). **Los indicadores con cargas externas menores a 0.70 no deben eliminarse automáticamente.** Cuando la carga se encuentra cercana al umbral, especialmente entre 0.60 y 0.70, la decisión debe considerar tanto la evidencia estadística como la pertinencia conceptual del indicador. En el modelo **CAITIZEN**, los indicadores con cargas moderadas deben revisarse en función de su efecto sobre la **AVE**, la confiabilidad interna y la validez discriminante. Solo se recomienda su eliminación cuando la depuración mejora la calidad del modelo sin debilitar la cobertura teórica del constructo. Ver **Figura 5.1**.

Figura 5.1. Vista general de cargas externas del modelo CAITIZEN en SmartPLS4.

	AFDJ	CAIL	CAITZEN	EAR	HAIIC	MTTP
AFDJ1	0.678					
AFDJ10	0.781					
AFDJ2	0.747					
AFDJ3	0.771					
AFDJ4	0.720					
AFDJ5	0.745					
AFDJ6	0.785					
AFDJ7	0.812					
AFDJ8	0.768					
AFDJ9	0.803					
CAIL1		0.676				
CAIL10		0.699				
CAIL2		0.555				
CAIL3		0.602				
CAIL4		0.763				
CAIL5		0.753				
CAIL6		0.710				
CAIL7		0.726				
CAIL8		0.790				
CAIL9		0.741				
CAITZEN1			0.712			
CAITZEN10			0.795			
CAITZEN2			0.740			
CAITZEN3			0.719			
CAITZEN4			0.742			
CAITZEN5			0.731			
CAITZEN6			0.658			

Fuente: Elaboración propia con SmartPLS4.

Criterios para conservar o eliminar indicadores

La decisión de conservar o eliminar indicadores no debe tomarse de forma automática. En modelos reflectivos, una carga externa menor a **0.70** no implica necesariamente que el indicador deba eliminarse. Primero debe revisarse si su eliminación mejora la **AVE**, la **confiabilidad interna**, la **validez discriminante** y la claridad conceptual del constructo. En el modelo **CAITIZEN**, esta decisión es especialmente importante porque los indicadores no son simples variables numéricas; representan dimensiones conceptuales vinculadas con alfabetización crítica en **IA**, responsabilidad ética, justicia de datos, colaboración humano-**IA**, transparencia metacognitiva y ciudadanía sostenible asistida por **IA**. Por ejemplo, un indicador de **CAIL** con carga inferior a **0.70** podría conservarse si expresa un contenido teóricamente relevante para la alfabetización crítica. En cambio, podría eliminarse si presenta carga baja, no mejora la interpretación del constructo y afecta la validez convergente. La depuración debe equilibrar evidencia estadística y justificación teórica; eliminar ítems solo para mejorar cifras puede empobrecer el contenido del modelo (Hair et al., 2019). Ver **Tabla 5.2**.

Tabla 5.2. Decisión metodológica sobre indicadores con cargas bajas.

Situación del indicador	Acción recomendada
Carga ≥ 0.70	Conservar
Carga entre 0.40 y 0.70	Revisar efecto en AVE , confiabilidad y teoría
Carga < 0.40	Considerar eliminación con justificación
Indicador teóricamente relevante	No eliminar solo por criterio numérico
Indicador redundante o ambiguo	Revisar posible eliminación

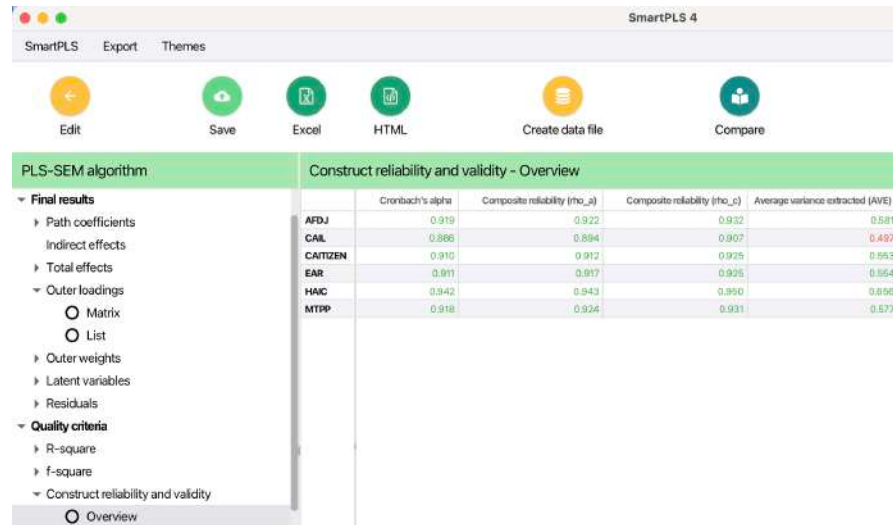
Fuente: Elaboración propia con base en Hair et al. (2021), Fornell y Larcker (1981), Henseler et al. (2015) y Cepeda-Carrion et al. (2017).

Evaluación de confiabilidad interna

Después de revisar las cargas externas, se evalúa la **confiabilidad interna** de cada constructo. Esta revisión permite determinar si los indicadores asignados a **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** funcionan de manera consistente como manifestaciones de una misma variable latente. En **SmartPLS4**, la confiabilidad se revisa mediante **Cronbach Alpha**, **rho_A** y **Composite Reliability** o **rho_c**. Como criterio general, se esperan valores iguales o superiores a **0.70**. En los resultados del modelo **CAITIZEN**, todos los constructos superan este umbral, lo cual indica consistencia interna adecuada. Sin embargo, también debe observarse si algún valor se aproxima o supera **0.95**, porque una confiabilidad demasiado alta puede sugerir redundancia de indicadores. En el caso mostrado, **HAIC** presenta **rho_c = 0.950**, valor que debe reportarse con cautela, aunque no invalida el constructo. Por tanto, la confiabilidad interna del modelo es favorable, pero debe interpretarse junto con cargas externas y **AVE** (Gefen et al., 2000).

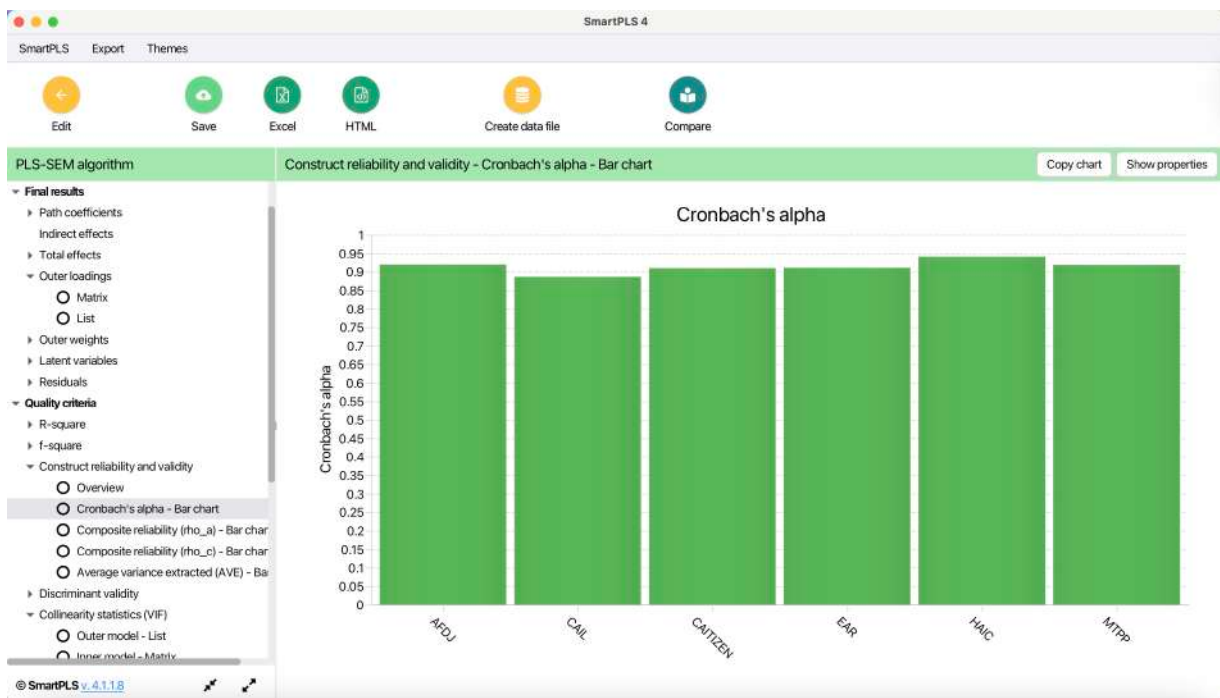
Lo anterior se muestra en la visión conjunta de la confiabilidad y validez del constructo se presenta en la **Figura 5.2**, el comportamiento del alfa de Cronbach se muestra en la **Figura 5.3**, los valores de confiabilidad compuesta **rho_A** se presentan en la **Figura 5.4** y La confiabilidad compuesta **rho_c** se observa en la **Figura 5.5**.

Figura 5.2. Construct Reliability and Validity – Overview en SmartPLS4.



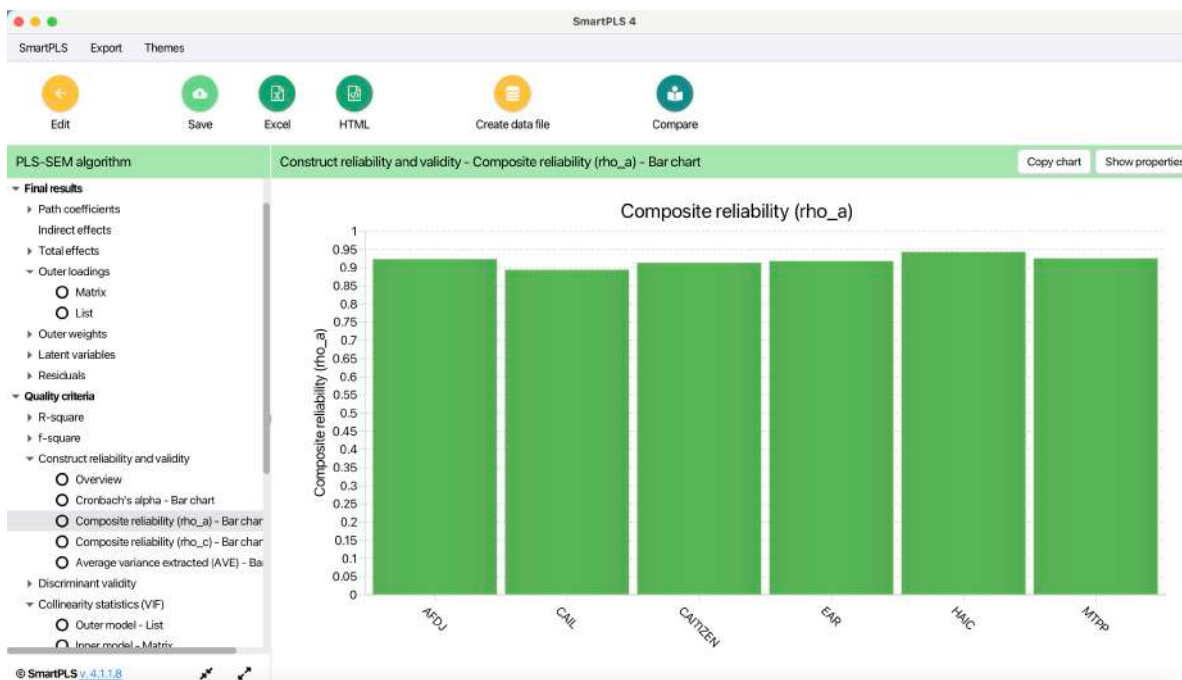
Fuente: Elaboración propia con SmartPLS4.

Figura 5.3. Gráfico de Cronbach Alpha.



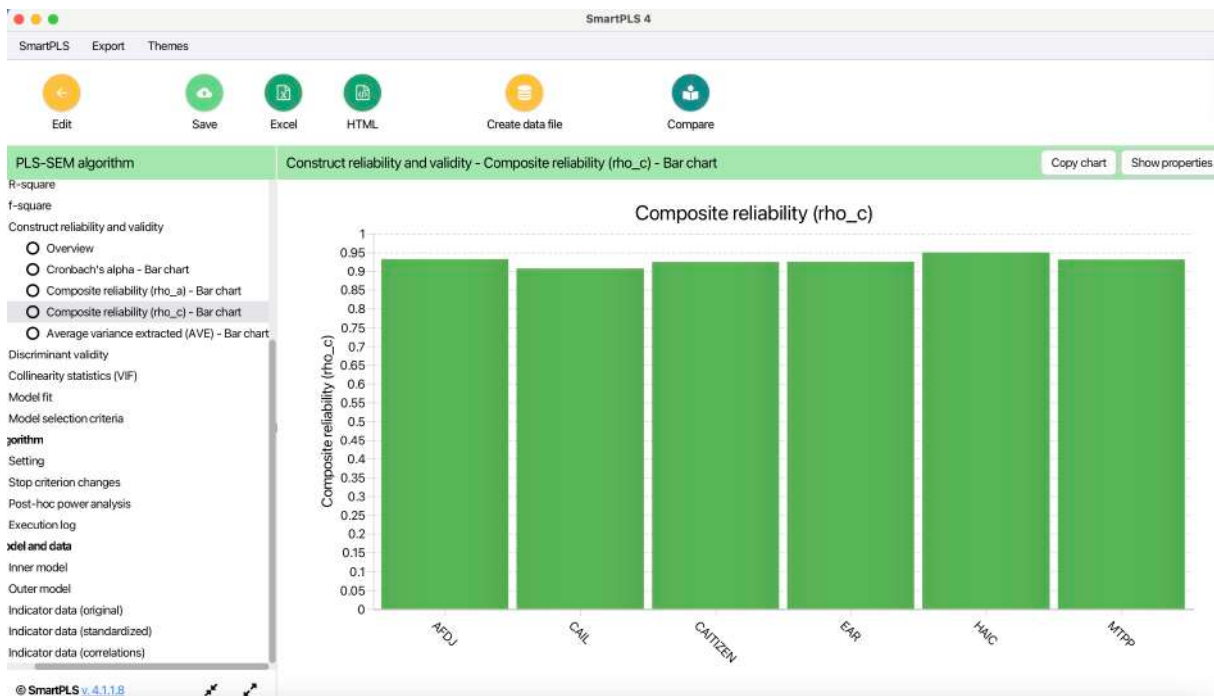
Fuente: Elaboración propia con SmartPLS4.

Figura 5.4. Gráfico de *Composite Reliability rho_a*.



Fuente: Elaboración propia con SmartPLS4.

Figura 5.5. Gráfico de *Composite Reliability rho_c*.



Fuente: Elaboración propia con SmartPLS4.

Juan Mejía Tréjo

Evaluación de validez convergente mediante AVE

La **validez convergente** se evalúa mediante la **Average Variance Extracted** o **AVE**. Este criterio indica qué proporción de la varianza de los indicadores es explicada por el constructo latente. Como regla general, un valor de **AVE ≥ 0.50** indica que el constructo explica al menos la mitad de la varianza de sus indicadores. En el modelo **CAITIZEN**, los resultados muestran valores adecuados para **AFDJ** = 0.581, **CAITIZEN** = 0.553, **EAR** = 0.554, **HAIC** = 0.656 y **MTPP** = 0.577. Sin embargo, **CAIL** presenta **AVE = 0.497**, valor ligeramente inferior al umbral de 0.50. Esta diferencia mínima **no invalida automáticamente el constructo**, pero sí exige revisión. Deben observarse los indicadores **CAIL1–CAIL10**, especialmente aquellos con cargas inferiores a **0.70**, para valorar si conviene conservarlos por relevancia conceptual o si su eliminación mejora el **AVE** sin afectar el contenido del constructo. La decisión debe ser prudente, porque **CAIL** representa una dimensión central del modelo (Fornell & Larcker, 1981). Ver Figura 5.6.

Figura 5.6. Gráfico de *Average Variance Extracted* (AVE) del modelo CAITIZEN.



Fuente: Elaboración propia con SmartPLS4.

Evaluación de validez discriminante

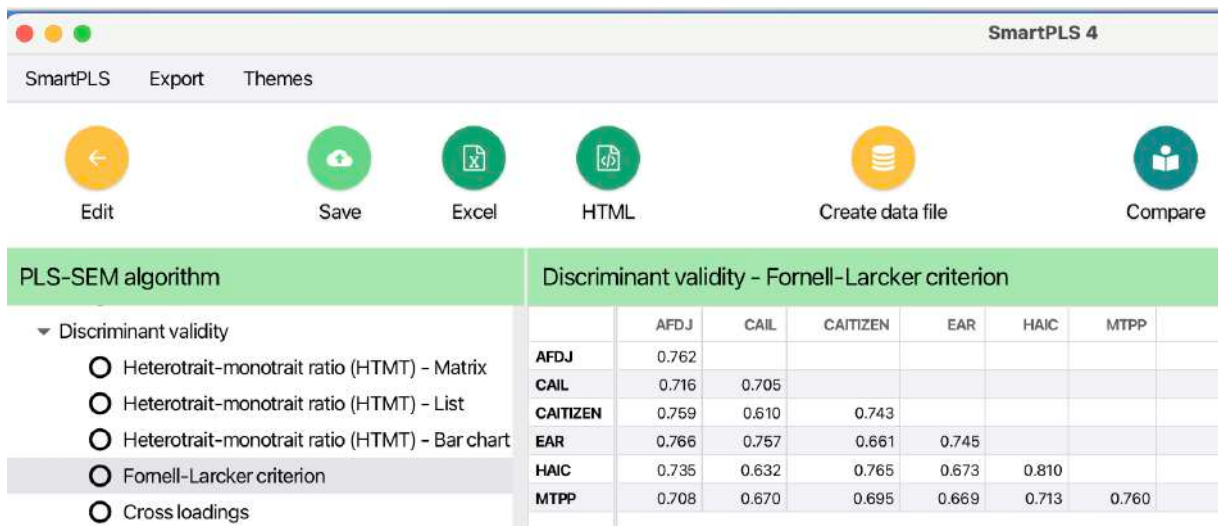
La **validez discriminante** permite verificar si los constructos del modelo **CAITIZEN** son empíricamente distintos entre sí. Esta revisión es necesaria porque **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** están conceptualmente relacionados, pero no deben medir exactamente el mismo fenómeno. En un modelo de ciudadanía sostenible asistida por **IA**, es esperable que la alfabetización crítica, la responsabilidad ética, la justicia de datos, la colaboración humano-IA y la transparencia metacognitiva se encuentren asociadas. Sin embargo, una asociación alta no debe confundirse con equivalencia conceptual. Por ello, la validez discriminante busca confirmar que cada constructo conserva identidad propia dentro del modelo de medición reflectivo. Esta evaluación se realiza antes de interpretar el modelo estructural, porque si los constructos no se distinguen adecuadamente, las rutas estructurales podrían reflejar solapamiento de medición y no relaciones sustantivas entre dimensiones diferenciadas (Hair et al., 2021).

Para evaluar la validez discriminante se emplean **tres criterios complementarios: Fornell-Larcker, Cross Loadings y HTMT**. El criterio **Fornell-Larcker** compara la raíz cuadrada del **AVE** de cada constructo con sus correlaciones frente a los demás constructos. En términos generales, la raíz cuadrada del **AVE** debe ser mayor que las correlaciones del constructo con otros constructos. Las **Cross Loadings** permiten revisar si cada indicador carga más alto en su propio constructo que en los demás. Finalmente, el **HTMT** o **Heterotrait-Monotrait Ratio** evalúa el grado de solapamiento entre constructos y es considerado uno de los criterios más sensibles para detectar problemas de validez discriminante en modelos basados en varianza. En este libro se adopta el criterio estricto de **HTMT < 0.85** y el criterio aceptable de **HTMT < 0.90** (Henseler et al., 2015).

Criterio Fornell-Larcker

En la salida **Fornell-Larcker, SmartPLS4** muestra en la diagonal la raíz cuadrada del **AVE** de cada constructo: **AFDJ = 0.762**, **CAIL = 0.705**, **CAITIZEN = 0.743**, **EAR = 0.745**, **HAIC = 0.810** y **MTPP = 0.760**. Estos valores deben compararse con las correlaciones entre constructos. El resultado muestra una situación que requiere interpretación cuidadosa. Por ejemplo, la correlación entre **AFDJ** y **EAR** es **0.766**, ligeramente superior a la raíz cuadrada del **AVE** de **AFDJ = 0.762** y también superior a la de **EAR = 0.745**. Asimismo, la correlación entre **CAIL** y **AFDJ** es **0.716**, ligeramente superior a la raíz cuadrada del **AVE** de **CAIL = 0.705**. También se observa que la correlación entre **HAIC** y **CAITIZEN** es **0.765**, superior a la raíz cuadrada del **AVE** de **CAITIZEN = 0.743**. Estos resultados no invalidan automáticamente el modelo, pero indican que algunos constructos presentan cercanía empírica y deben revisarse con los otros criterios de validez discriminante (Fornell & Larcker, 1981). Ver **Figura 5.7**.

Figura 5.7. Discriminant Validity – Fornell-Larcker Criterion en SmartPLS4.



Fuente: Elaboración propia con **SmartPLS4**.

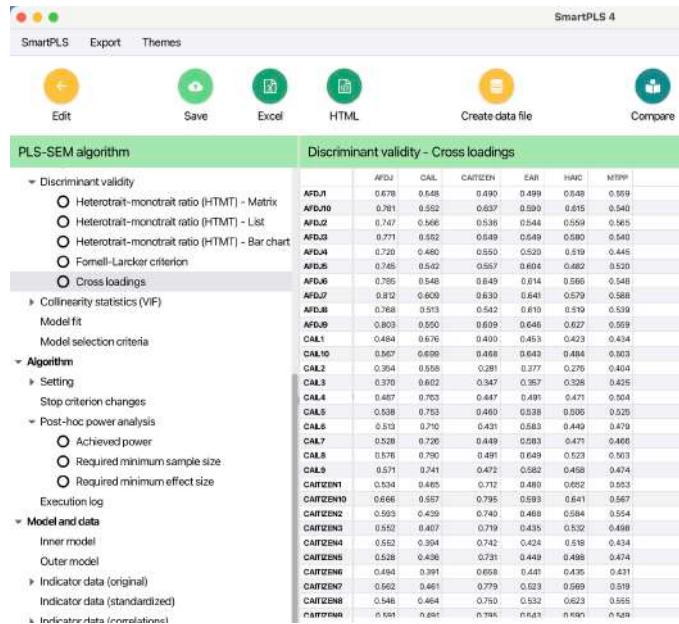
La interpretación del criterio **Fornell-Larcker** debe vincularse con los resultados previos de validez convergente. En particular, **CAIL** presentó un **AVE = 0.497**, ligeramente inferior al umbral recomendado de **0.50**. Esta situación reduce la raíz cuadrada del **AVE** de **CAIL** y puede explicar por qué algunas correlaciones con otros constructos aparecen cercanas o superiores a su valor diagonal. Por ello, el posible problema no necesariamente se encuentra en todo el modelo, sino en la necesidad de revisar con mayor detalle los indicadores de **CAIL**, especialmente aquellos con cargas externas menores a **0.70**. De igual manera, la cercanía entre **AFDJ** y **EAR** resulta conceptualmente comprensible, porque ambos constructos comparten una base ética: **EAR** se orienta a la responsabilidad ética general ante la **IA**, mientras que **AFDJ** se enfoca en justicia de datos, equidad y sesgo algorítmico. Esta cercanía debe reportarse, pero no implica que ambos constructos sean idénticos (Hair et al., 2019).

Criterio Cross Loading

Las **Cross Loadings** permiten realizar una revisión indicador por indicador. El criterio esperado es que cada ítem cargue más alto en su propio constructo que en los constructos restantes. En la salida visible de **SmartPLS4** se observa que los indicadores de **AFDJ** cargan más alto en **AFDJ** que en **CAIL**, **CAITIZEN**, **EAR**, **HAIC** o **MTPP**. Por ejemplo, **AFDJ7 = 0.812**, **AFDJ9 = 0.803**, **AFDJ6 = 0.785** y **AFDJ10 = 0.781** cargan principalmente en su constructo. En el caso de **CAIL**, aunque algunos indicadores presentan cargas externas moderadas, como **CAIL2 = 0.558** y **CAIL3 = 0.602**, siguen cargando más alto en **CAIL** que en los demás constructos visibles. Lo mismo ocurre con los indicadores visibles de **CAITIZEN**, donde **CAITIZEN10 = 0.795**, **CAITIZEN7 = 0.779**, **CAITIZEN8 = 0.750** y **CAITIZEN4 = 0.742** presentan su carga principal en el constructo dependiente. Este patrón respalda la asignación de indicadores a sus constructos correspondientes (Hair et al., 2021). Ver **Figura 5.8**.

Juan Mejía Trejo

Figura 5.8. Discriminant Validity – Cross Loadings en SmartPLS4.



Fuente: Elaboración propia con SmartPLS4.

El análisis de **Cross Loadings** complementa la lectura de **Fornell-Larcker** porque permite distinguir entre dos tipos de problemas. El primero ocurre cuando un constructo presenta correlaciones elevadas con otros constructos; el segundo ocurre cuando un indicador carga más fuerte en un constructo diferente al que fue asignado. En el modelo **CAITIZEN**, la salida visible sugiere que, aunque existen correlaciones altas entre algunos constructos, los indicadores siguen mostrando mayor asociación con su propio constructo. Esto es importante porque indica que el problema no parece ser una asignación incorrecta masiva de ítems, sino una cercanía conceptual y empírica entre algunas dimensiones. Por tanto, no se recomienda eliminar indicadores únicamente por la lectura de **Fornell-Larcker**. La decisión debe considerar también cargas externas, **AVE**, confiabilidad y sentido teórico de cada indicador. En especial, **CAIL** requiere revisión fina por sus cargas moderadas y por su **AVE** ligeramente inferior al umbral (Gefen et al., 2000).

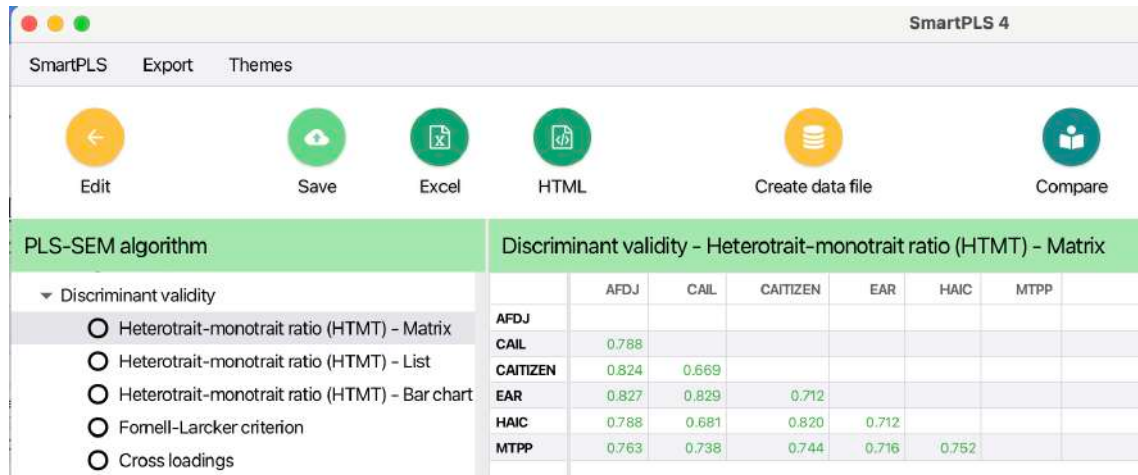
Criterio HTMT

El criterio **HTMT** ofrece una lectura más robusta de la validez discriminante. En la matriz de **SmartPLS4**, todos los valores observados se encuentran por debajo del criterio estricto de **0.85**. Los valores más altos corresponden a **CAIL–EAR = 0.829**, **AFDJ–EAR = 0.827**, **AFDJ–CAITIZEN = 0.824** y **CAITIZEN–HAIC = 0.820**. Aunque estos valores son elevados, permanecen por debajo de **0.85**, por lo que no indican un problema crítico de validez discriminante bajo el criterio estricto. Los demás valores también son aceptables: **AFDJ–CAIL = 0.788**, **AFDJ–HAIC = 0.788**, **AFDJ–MTPP = 0.763**, **CAIL–CAITIZEN = 0.669**, **CAIL–HAIC = 0.681**, **CAIL–MTPP = 0.738**, **EAR–**

Juan Mejía Trejo

HAIC = 0.712, EAR-MTPP = 0.716 y HAIC-MTPP = 0.752. Estos resultados respaldan que los constructos del modelo **CAITIZEN** son empíricamente distinguibles (Henseler et al., 2015). Ver **Figura 5.9**.

Figura 5.9. Discriminant Validity – Heterotrait-Monotrait Ratio (HTMT) – Matrix en SmartPLS4.



Fuente: Elaboración propia con **SmartPLS4**.

La lectura integrada de los tres criterios permite formular una conclusión metodológica equilibrada. El criterio **Fornell-Larcker** muestra algunos puntos de atención, especialmente en las relaciones **AFDJ-EAR**, **CAIL-AFDJ** y **CAITIZEN-HAIC**. Sin embargo, las **Cross Loadings** visibles muestran que los indicadores cargan más alto en sus constructos asignados, y el **HTMT** confirma que todos los pares de constructos se mantienen por debajo de **0.85**. Por tanto, la evidencia conjunta permite sostener que el modelo presenta validez discriminante aceptable, aunque con observaciones que deben reportarse. La decisión metodológica recomendada es conservar los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, revisar con especial atención los indicadores de **CAIL** por su **AVE = 0.497**, y reportar la cercanía conceptual entre **AFDJ** y **EAR** como un resultado esperado dentro de un modelo de ciudadanía asistida por **IA** (Hair et al., 2021).

En consecuencia, la validez discriminante del modelo **CAITIZEN** puede considerarse suficientemente respaldada para continuar con la evaluación del modelo estructural. No obstante, el reporte debe evitar una afirmación absoluta de cumplimiento sin matices. La redacción recomendada es señalar que el criterio **HTMT** respalda la distinción empírica entre constructos, que las **Cross Loadings** no evidencian cargas cruzadas dominantes en los indicadores visibles y que **Fornell-Larcker** muestra algunas relaciones cercanas que requieren cautela interpretativa. Esta conclusión es metodológicamente más sólida que declarar simplemente “cumple”, porque reconoce la complejidad del modelo: las dimensiones de ciudadanía asistida por **IA** están relacionadas entre sí, pero conservan identidad empírica suficiente para ser analizadas como constructos diferenciados. Con esta revisión, el modelo de

Juan Mejía Trejo

medición reflectivo queda en condiciones de avanzar hacia la evaluación estructural, siempre que se mantenga la decisión de revisar puntualmente **CAIL** y reportar las relaciones conceptualmente próximas (Hair et al., 2019). Ver **Tabla 5.3**.

Tabla 5.3. Criterios para evaluar validez discriminante

Método	Qué revisa	Criterio general	Resultado en CAITIZEN
Fornell-Larcker	Compara $\sqrt{\text{AVE}}$ con correlaciones entre constructos	$\sqrt{\text{AVE}}$ debe ser mayor que las correlaciones	Presenta puntos de atención en AFDJ-EAR , CAIL-AFDJ y CAITIZEN-HAIC
Cross Loadings	Verifica si cada indicador carga más alto en su constructo	Cada indicador debe cargar más alto en su propio constructo	En los indicadores visibles, las cargas principales corresponden al constructo asignado
HTMT	Evalúa solapamiento entre constructos	< 0.85 estricto; < 0.90 aceptable	Todos los valores son < 0.85 ; respalda validez discriminante
Decisión integrada	Integra los tres criterios	Conservar si la evidencia conjunta es aceptable	Se conserva el modelo; revisar CAIL y reportar cautela en relaciones cercanas

Fuente: Elaboración propia con base en Fornell y Larcker (1981), Henseler et al. (2015) y Hair et al. (2021).

Validez discriminante conjunta

Para complementar la evaluación de la **validez discriminante**, los resultados pueden integrarse en una sola matriz que combine dos criterios: **Fornell-Larcker** y **HTMT**. En la diagonal se reporta la raíz cuadrada de la **AVE** de cada constructo; en el triángulo inferior se presentan las correlaciones entre constructos utilizadas para el criterio **Fornell-Larcker**; y en el triángulo superior se reportan los valores **HTMT**. Esta forma de presentación permite observar, en una misma tabla, si cada constructo comparte más varianza con sus propios indicadores que con otros constructos y si el solapamiento entre constructos permanece dentro de umbrales aceptables. En el modelo **CAITIZEN**, esta tabla es especialmente útil porque los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** están conceptualmente relacionados, pero deben conservar identidad empírica propia. Ver **Tabla 5.4**.

Tabla 5.4. Validez discriminante del modelo de medición CAITIZEN

Constructo	AFDJ	CAIL	CAITIZEN	EAR	HAIC	MTPP
AFDJ	0.762	0.794	0.824	0.827	0.788	0.763
CAIL	0.718	0.722	0.679	0.834	0.694	0.732
CAITIZEN	0.759	0.614	0.743	0.712	0.820	0.744
EAR	0.766	0.759	0.661	0.745	0.712	0.716
HAIC	0.735	0.637	0.765	0.673	0.810	0.752
MTPP	0.708	0.665	0.695	0.669	0.713	0.760

Fuente: Elaboración propia

Criterio Fornell-Larcker: los valores de la diagonal representan la raíz cuadrada de la **AVE**. Los valores ubicados debajo de la diagonal representan las correlaciones entre constructos. Los valores ubicados por encima de la diagonal representan los coeficientes **HTMT**. Los valores de la diagonal representan la raíz cuadrada de la **AVE**. Los valores ubicados debajo de la diagonal corresponden a las correlaciones entre constructos utilizadas en el criterio **Fornell-Larcker**. Los valores ubicados por encima de la diagonal corresponden a los coeficientes **HTMT**. El **HTMT** evalúa la **validez discriminante** al determinar si los constructos son empíricamente distintos. Valores menores a **0.85** representan un criterio conservador, mientras que valores menores a **0.90** pueden considerarse aceptables cuando los constructos están conceptualmente relacionados. En este modelo, todos los valores **HTMT** se mantienen por debajo de **0.85**, lo que respalda la validez discriminante. Sin embargo, el criterio **Fornell-Larcker** muestra relaciones cercanas entre algunos constructos, especialmente **AFDJ–EAR**, **CAIL–AFDJ** y **CAITIZEN–HAIC**; por ello, estas relaciones deben reportarse con cautela metodológica (Fornell & Larcker, 1981; Henseler et al., 2015; Franke & Sarstedt, 2019).

Decisión metodológica sobre el modelo de medición

Una vez revisadas las **cargas externas**, la **confiabilidad interna**, la **validez convergente mediante AVE** y la **validez discriminante**, corresponde formular una decisión metodológica integrada sobre el modelo de medición reflectivo **CAITIZEN**. Esta decisión no debe reducirse a verificar umbrales de manera aislada, porque la evaluación de un modelo reflectivo requiere interpretar conjuntamente la calidad de los indicadores, la consistencia interna de los constructos, la varianza explicada por cada variable latente y la distinción empírica entre constructos. En este sentido, el modelo **CAITIZEN** muestra una base estadística favorable: los constructos presentan confiabilidad interna adecuada, los valores de **Cronbach Alpha**, **rho_A** y **rho_c** superan el umbral de **0.70**, y cinco de los seis constructos cumplen con el criterio de **AVE ≥ 0.50**. Por tanto, el modelo no requiere una depuración general, sino una revisión puntual de algunos indicadores y constructos específicos (Hair et al., 2021).

El principal punto de atención se encuentra en **CAIL**, cuyo **AVE = 0.497** se ubica apenas por debajo del criterio recomendado de **0.50**. Esta diferencia es mínima, por lo que no debe interpretarse como falla grave ni como motivo automático para eliminar el constructo. Sin embargo, sí justifica una revisión cuidadosa de los indicadores **CAIL1–CAIL10**, especialmente aquellos con cargas externas menores a **0.70**. En este caso, la decisión metodológica más prudente es conservar inicialmente el constructo **CAIL**, revisar si la eliminación de algún indicador mejora el **AVE** y la confiabilidad interna, y verificar que cualquier ajuste no debilite la cobertura conceptual de la **alfabetización crítica en inteligencia artificial**. La eliminación de indicadores solo debe realizarse si mejora la calidad estadística sin empobrecer el contenido teórico del constructo (Hair et al., 2019).

Respecto a la validez discriminante, el modelo presenta una situación que debe reportarse con cautela metodológica. El criterio **Fornell-Larcker** muestra algunos puntos de atención, principalmente en las relaciones **AFDJ–EAR**, **CAIL–AFDJ** y **CAITIZEN–HAIC**, donde algunas correlaciones se aproximan o superan ligeramente la raíz cuadrada del **AVE** correspondiente. No obstante, esta situación no invalida por sí misma el modelo, porque los constructos analizados pertenecen a un mismo campo conceptual: ciudadanía sostenible asistida por **IA**. Es esperable que exista cercanía entre dimensiones como ética, justicia de datos, alfabetización crítica, colaboración humano–**IA** y ciudadanía. Por ello, el resultado debe entenderse como una advertencia de proximidad conceptual y no como evidencia definitiva de ausencia de validez discriminante (Fornell & Larcker, 1981).

La revisión mediante **Cross Loadings** y **HTMT** permite sostener una decisión más favorable. Las **Cross Loadings** visibles muestran que los indicadores cargan principalmente en sus constructos asignados, lo cual respalda la estructura del modelo de medición. Además, los valores **HTMT** se mantienen por debajo del criterio estricto de **0.85**, lo que indica que no existe solapamiento crítico entre constructos. Este resultado es especialmente importante porque el **HTMT** se considera un criterio sensible para evaluar validez discriminante en modelos **PLS-SEM**. Por tanto, aunque **Fornell-Larcker** muestra relaciones cercanas, la evidencia conjunta de **Cross Loadings** y **HTMT** permite conservar los constructos del modelo **CAITIZEN** como dimensiones empíricamente distinguibles (Henseler et al., 2015).

Con base en la evaluación integrada, la decisión metodológica recomendada es conservar el modelo de medición reflectivo **CAITIZEN**, sin eliminar constructos en esta etapa. La confiabilidad interna es adecuada en todos los casos; la validez convergente cumple en cinco constructos y presenta solo una desviación mínima en **CAIL**; y la validez discriminante resulta aceptable cuando se interpreta de manera conjunta mediante **Fornell-Larcker**, **Cross Loadings** y **HTMT**. En consecuencia, el modelo puede avanzar hacia la evaluación del modelo estructural, siempre que se reporten explícitamente dos cautelas: primero, la necesidad de revisar los indicadores de **CAIL** por su **AVE = 0.497**; segundo, la cercanía conceptual entre **AFDJ** y **EAR**, esperable en un modelo donde la justicia de datos y la responsabilidad ética se encuentran relacionadas (Hair et al., 2021). Ver **Tabla 5.5**.

Juan Mejía Tréjo

Tabla 5.5. Decisión metodológica sobre el modelo de medición CAITIZEN

Criterio	Resultado general	Decisión metodológica
Cargas externas	Existen indicadores adecuados y algunos indicadores con cargas menores a 0.70	Revisar indicadores con cargas moderadas antes de eliminar; no depurar automáticamente
Confiabilidad interna	Todos los constructos superan 0.70 en Cronbach Alpha , rho_A y rho_c	Conservar los constructos por consistencia interna adecuada
Validez convergente	Cinco constructos cumplen AVE ≥ 0.50 ; CAIL = 0.497	Conservar CAIL con revisión puntual de indicadores
Fornell-Larcker	Presenta puntos de atención en AFDJ-EAR , CAIL-AFDJ y CAITIZEN-HAIC	Reportar cercanía conceptual y contrastar con otros criterios
Cross Loadings	Los indicadores visibles cargan principalmente en su constructo asignado	Mantener la asignación original de indicadores
HTMT	Todos los valores se mantienen por debajo de 0.85	Respaldar validez discriminante aceptable
Decisión global	Modelo confiable, viable y con validez discriminante aceptable, aunque con cautelas puntuales	Conservar el modelo reflectivo y avanzar al modelo estructural

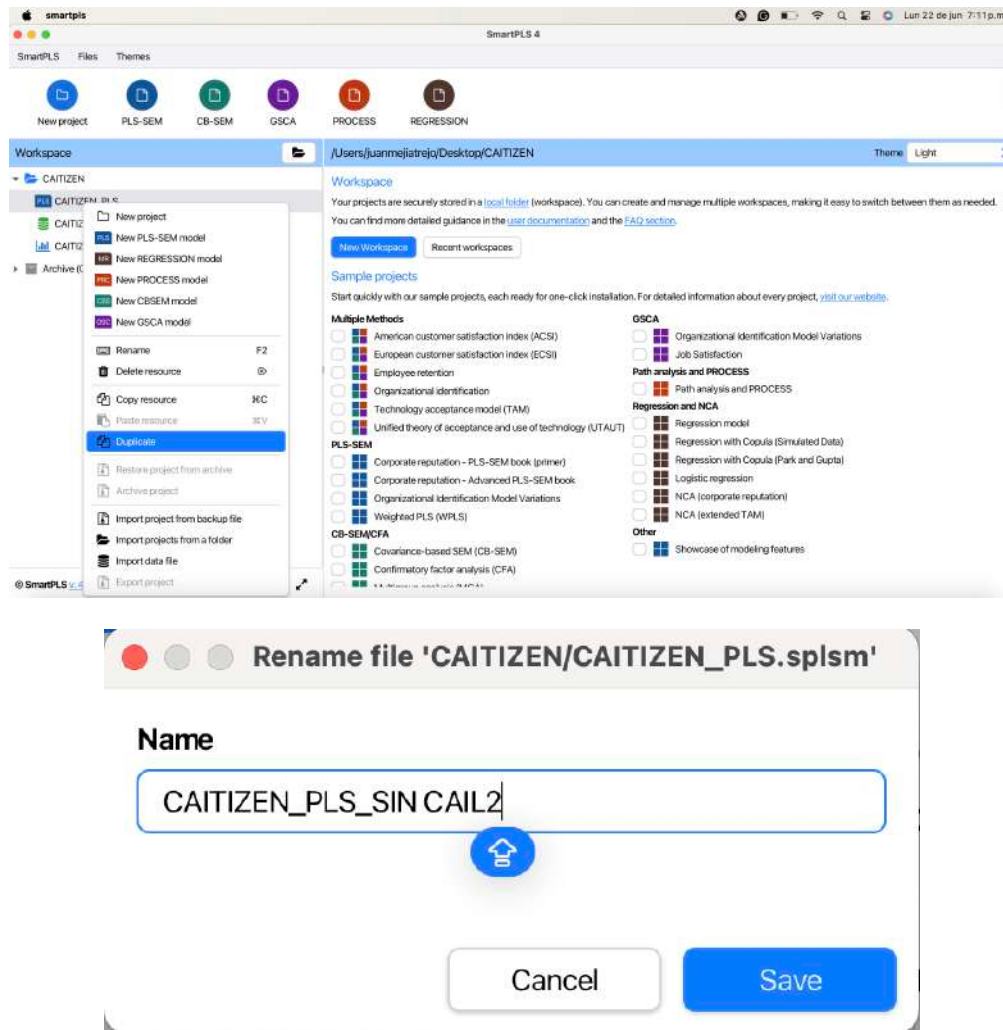
Fuente: Elaboración propia con base en Hair et al. (2021), Hair et al. (2019), Fornell y Larcker (1981) y Henseler et al. (2015).

Copias del modelo original

Cuando se identifican indicadores con cargas externas que no cumplen plenamente el criterio recomendado, no es conveniente eliminarlos directamente del modelo original. Antes de retirar cualquier ítem, se debe generar una copia del modelo en **SmartPLS4**. Esta práctica permite conservar intacta la versión inicial, mantener la trazabilidad metodológica y comparar los resultados antes y después de la depuración. De esta manera, el investigador puede justificar si la eliminación de un indicador mejora realmente la calidad del modelo de medición sin perder la configuración original.

Para ello, en el panel izquierdo de **SmartPLS4** se debe ubicar el modelo actual, por ejemplo, **CAITIZEN_PLS**. Después, se hace clic derecho sobre el modelo y se selecciona la opción disponible para duplicarlo, como **Duplicate**, **Copy**, **Clone** o **SAVE as**, según la versión del software. La copia debe renombrarse de manera clara, por ejemplo: **CAITIZEN_PLS_Depurado_1** o, de forma más específica, **CAITIZEN_PLS_sin_CAIL2**. Este nombre permite identificar que se trata de una versión alternativa del modelo en la que se evaluará el efecto de retirar el indicador **CAIL2**. Ver Figura 5.10

Figura 5.10. Guardado del modelo depurado y conservación de reportes en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Una vez creada la copia, se abre únicamente el modelo duplicado y en esa versión se elimina el indicador **CAIL2**. El modelo original debe permanecer sin cambios. Es importante considerar que **SmartPLS4** suele conservar los cambios realizados en el modelo de trabajo; por ello, cualquier modificación debe hacerse únicamente en la copia depurada. Sin embargo, los resultados generados por cada corrida no deben depender solo de la visualización en pantalla. Después de ejecutar el algoritmo, conviene guardar o exportar el reporte correspondiente para documentar qué versión del modelo produjo esos resultados.

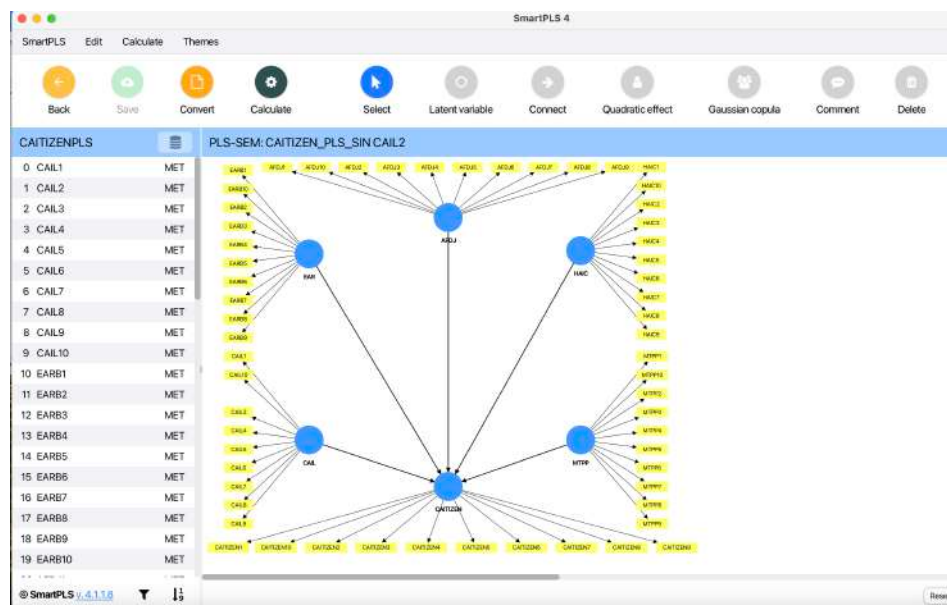
Posteriormente, se ejecuta nuevamente el algoritmo desde **Calculate** → **PLS-SEM Algorithm**. Con los nuevos resultados, se revisa si la eliminación del indicador mejora la **AVE** de **CAIL**, la **rho_A**, la **rho_c**, las cargas externas restantes, el criterio **Fornell-Larcker**, las **Cross Loadings** y el **HTMT**. El reporte de esta corrida debe guardarse

Juan Mejía Trejo

con un nombre claro, por ejemplo, **Reporte_PLS_CAITIZEN_PLS_sin_CAIL2**. De esta forma, se conserva evidencia documental del modelo depurado y se puede comparar contra el reporte del modelo original.

La eliminación solo debe conservarse si mejora la calidad estadística del modelo sin debilitar el significado conceptual del constructo. Después de comparar el modelo original con el modelo depurado, se toma la decisión metodológica. Si el modelo sin **CAIL2** mejora de manera suficiente y mantiene coherencia teórica, puede adoptarse como versión depurada. Si la mejora es mínima o afecta la cobertura conceptual de **CAIL**, conviene conservar el indicador. El procedimiento de **Bootstrapping** debe ejecutarse solo cuando se haya definido el modelo final, ya que sus resultados de significancia, **t-value** y **p-value** deben corresponder a la versión definitiva del modelo. Ver Figura 5.11.

Figura 5.11. Modelo con ajuste final SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

La evaluación del modelo de medición reflectivo debe cerrarse con una síntesis metodológica que integre los criterios revisados y permita tomar una decisión clara antes de avanzar al modelo estructural. En el caso del modelo **CAITIZEN**, la evaluación se realizó en una secuencia ordenada: primero se ejecutó el algoritmo **PLS-SEM**, después se abrió el reporte completo en **SmartPLS4**, posteriormente se revisaron las **cargas externas**, la **confiabilidad interna**, la **validez convergente mediante AVE** y la **validez discriminante**. Esta secuencia es importante porque evita interpretar relaciones estructurales o hipótesis sin haber demostrado antes que los constructos están correctamente medidos. En modelos reflectivos, la calidad del modelo no depende de un solo criterio, sino de la lectura conjunta de indicadores, confiabilidad, varianza explicada y distinción entre constructos (Hair et al., 2021).

Juan Mejía Trejo

El primer paso consistió en revisar las **cargas externas** de los indicadores. Esta revisión permitió identificar qué ítems representan adecuadamente a sus constructos y cuáles requieren análisis adicional. En términos generales, las cargas iguales o superiores a **0.70** se consideran adecuadas; las cargas entre **0.40** y **0.70** deben revisarse antes de decidir su eliminación; y las cargas menores a **0.40** suelen considerarse problemáticas. En el modelo **CAITIZEN**, se observaron indicadores con cargas adecuadas y algunos indicadores con cargas moderadas, especialmente dentro de **CAIL**. Sin embargo, la decisión metodológica no debe consistir en eliminar automáticamente todos los indicadores menores a **0.70**, sino en revisar si su permanencia conserva la cobertura conceptual del constructo y si su eliminación mejora realmente la calidad estadística del modelo (Hair et al., 2019).

El segundo paso fue evaluar la **confiabilidad interna** mediante **Cronbach Alpha**, **rho_A** y **rho_c**. Los resultados muestran que todos los constructos superan el umbral de **0.70**, por lo que existe consistencia interna suficiente en **AFDJ**, **CAIL**, **CAITIZEN**, **EAR**, **HAIC** y **MTPP**. Este resultado permite sostener que los indicadores de cada constructo funcionan de manera coherente como manifestaciones de una misma variable latente. No obstante, también debe observarse si algún valor se aproxima a **0.95**, porque una confiabilidad muy alta puede sugerir redundancia entre indicadores. En el modelo **CAITIZEN**, **HAIC** presenta **rho_c = 0.950**, lo cual no invalida el constructo, pero debe reportarse con cautela como posible indicio de elevada similitud entre algunos ítems (Gefen et al., 2000).

El tercer paso consistió en revisar la **validez convergente** mediante la **Average Variance Extracted** o **AVE**. Este criterio permite determinar si cada constructo explica una proporción suficiente de la varianza de sus indicadores. En el modelo **CAITIZEN**, cinco constructos cumplen el criterio de **AVE ≥ 0.50**: **AFDJ = 0.581**, **CAITIZEN = 0.553**, **EAR = 0.554**, **HAIC = 0.656** y **MTPP = 0.577**. El único caso que requiere atención es **CAIL**, con **AVE = 0.497**, valor apenas inferior al umbral recomendado. Esta diferencia mínima no obliga a eliminar el constructo ni invalida el modelo, pero sí justifica revisar los indicadores **CAIL1–CAIL10**, principalmente aquellos con cargas externas moderadas, para valorar si algún ajuste mejora el **AVE** sin afectar el significado de la alfabetización crítica en inteligencia artificial (Fornell & Larcker, 1981).

El cuarto paso fue analizar la **validez discriminante** mediante **Fornell-Larcker**, **Cross Loadings** y **HTMT**. La evidencia debe interpretarse de manera integrada. El criterio **Fornell-Larcker** mostró algunos puntos de atención en las relaciones **AFDJ–EAR**, **CAIL–AFDJ** y **CAITIZEN–HAIC**, lo cual sugiere cercanía empírica entre ciertos constructos. Sin embargo, las **Cross Loadings** visibles indican que los indicadores cargan principalmente en sus constructos asignados, y los valores **HTMT** se mantienen por debajo del criterio estricto de **0.85**. Por tanto, la evidencia conjunta respalda una validez discriminante aceptable. La conclusión metodológica no debe ser que el modelo “cumple sin observaciones”, sino que presenta distinción empírica suficiente, con algunas relaciones conceptualmente cercanas que deben reportarse con cautela (Henseler et al., 2015).

Con base en la revisión integrada, la decisión metodológica es conservar el modelo de medición reflectivo **CAITIZEN** y avanzar hacia la evaluación del modelo estructural. Esta decisión se justifica porque la confiabilidad interna es adecuada en todos los constructos, la validez convergente cumple en cinco de seis constructos, la validez discriminante es aceptable al considerar conjuntamente **Fornell-Larcker**, **Cross Loadings** y **HTMT**, y no se identifica una razón suficiente para eliminar constructos completos. No obstante, deben mantenerse dos cautelas: revisar puntualmente **CAIL** por su **AVE = 0.497** y reportar la cercanía conceptual entre **AFDJ** y **EAR**, así como entre **CAITIZEN** y **HAIC**, porque estas relaciones reflejan proximidad temática dentro del campo de ciudadanía sostenible asistida por **IA** (Hair et al., 2021). Se recomienda hacer los ajustes de los ítems con cargas que no cumplan a fin de definir el modelo ajustado que cumpla con los requerimientos.

En consecuencia, el modelo de medición reflectivo **CAITIZEN** puede considerarse suficientemente respaldado para continuar con el análisis. La evaluación realizada permite afirmar que los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** cuentan con una base de medición razonable para avanzar hacia la evaluación del modelo estructural. Sin embargo, el reporte final debe ser metodológicamente prudente: no debe ocultar que **CAIL** se encuentra ligeramente por debajo del umbral de **AVE**, ni que **Fornell-Larcker** muestra relaciones cercanas entre algunos constructos. Esta forma de reportar fortalece la transparencia del análisis, porque reconoce tanto los resultados favorables como los puntos de atención que deben acompañar la interpretación posterior del modelo estructural (Hair et al., 2019).

Exportación de reportes a Excel

Después de ejecutar el algoritmo **PLS-SEM** en **SmartPLS4**, es recomendable exportar los resultados a **Excel** para conservar evidencia documental de cada corrida del modelo. Esta práctica es especialmente importante cuando se realizan ajustes, como retirar indicadores con cargas externas bajas, porque permite comparar el modelo original con las versiones depuradas. Aunque **SmartPLS4** conserva los cambios realizados en el modelo de trabajo, los reportes deben guardarse por separado para documentar qué resultados corresponden a cada versión estimada.

Para exportar el reporte, se debe abrir el reporte completo después de ejecutar **Calculate** → **PLS-SEM Algorithm**. En la barra superior se selecciona la opción **Excel**. **SmartPLS4** muestra una ventana de exportación donde puede configurarse el formato numérico, por ejemplo **US (example: 1,000.23)**. También pueden activarse opciones como **Render model image**, **Render graphs** y **Plain layout**, según se desee conservar imágenes del modelo, gráficos y tablas en un formato más claro. Posteriormente, se presiona **Export** y se guarda el archivo con un nombre identificable.

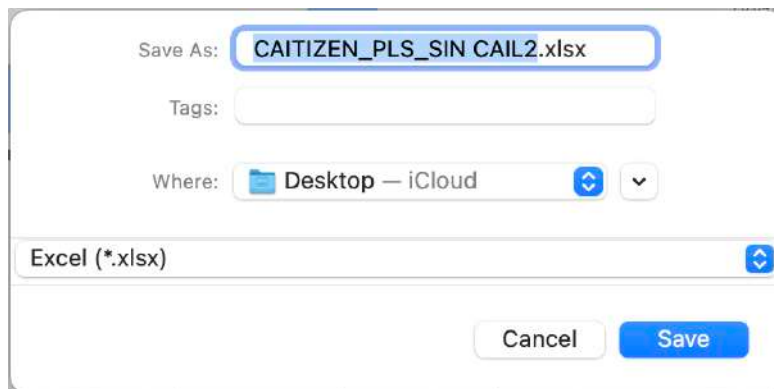
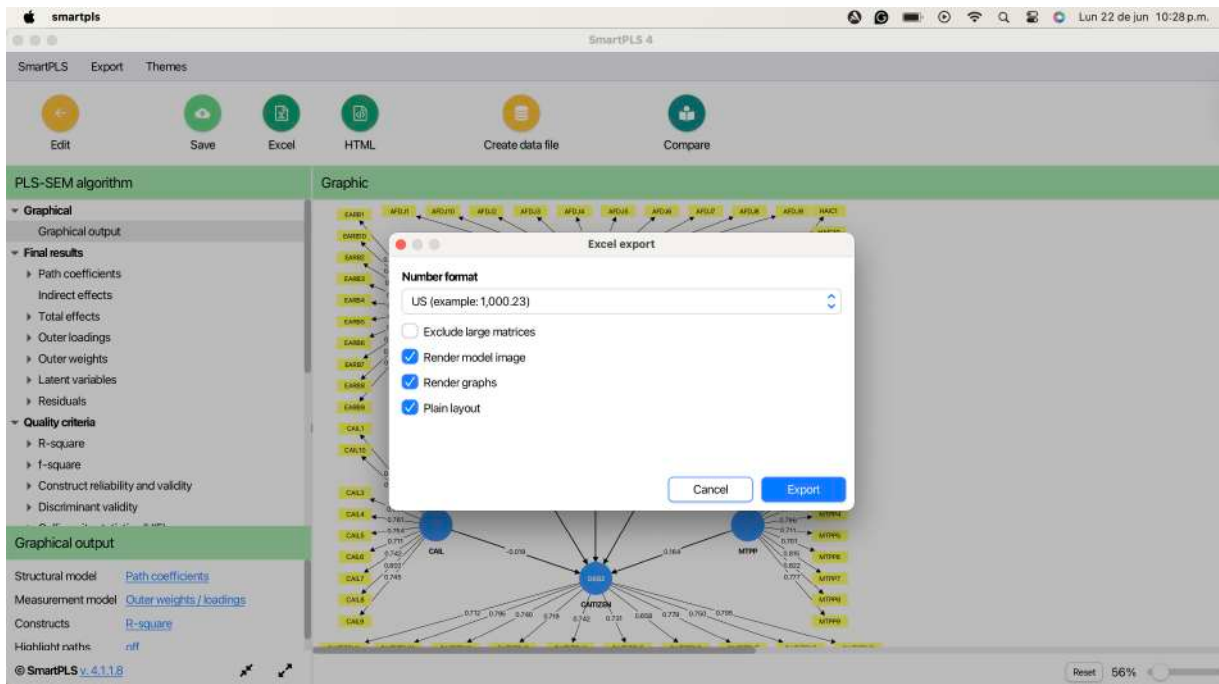
El nombre del archivo exportado debe permitir reconocer la versión del modelo y el procedimiento aplicado. Por ejemplo, para el modelo original puede utilizarse **Reporte_PLS_CAITIZEN_PLS_Original.xlsx**. Si se genera una versión depurada sin el indicador **CAIL2**, puede guardarse como

Juan Mejía Tréjo

Reporte_PLS_CAITIZEN_PLS_sin_CAIL2.xlsx. Esta nomenclatura facilita la trazabilidad metodológica, evita confundir resultados y permite justificar si la depuración mejoró la **AVE**, la **rho_A**, la **rho_c**, las cargas externas, el criterio **Fornell-Larcker**, las **Cross Loadings** o el **HTMT**.

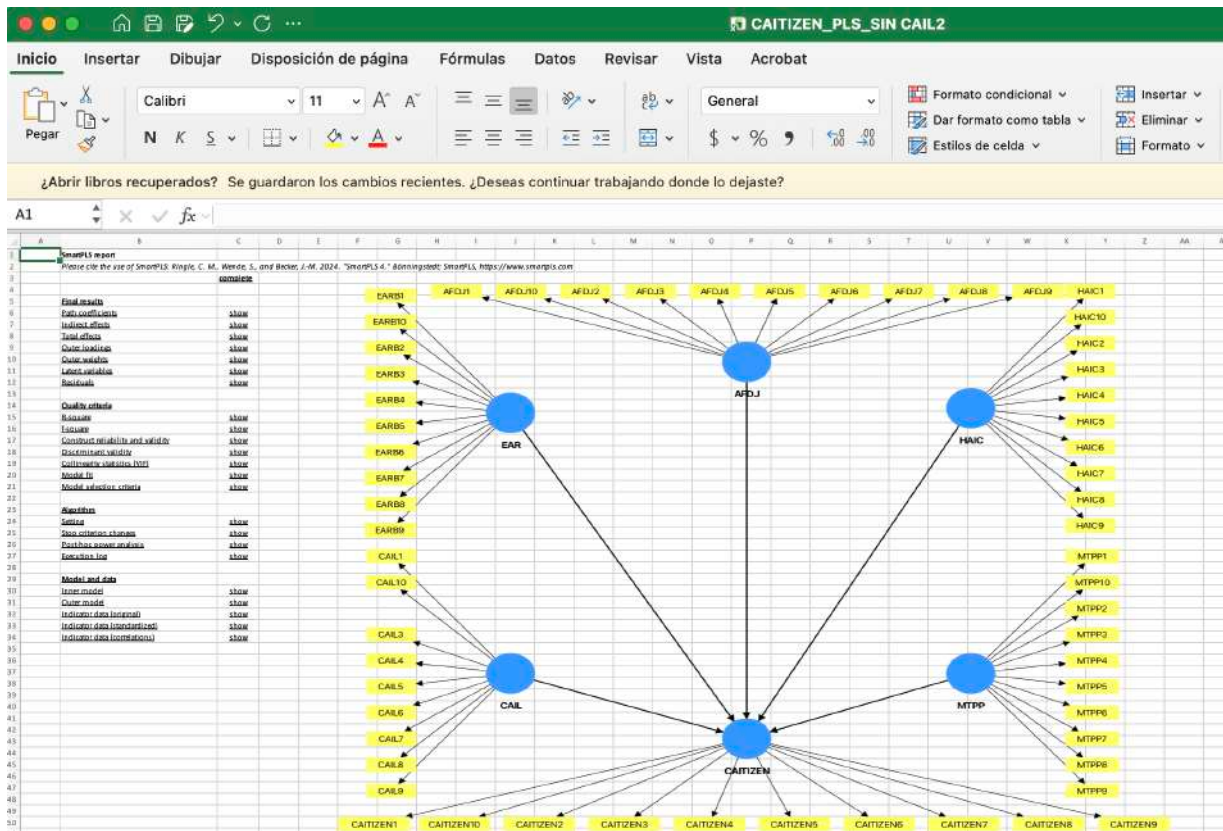
La exportación a **Excel** no modifica el modelo ni altera los resultados; únicamente guarda una copia del reporte generado por **SmartPLS4**. Por ello, debe realizarse después de cada corrida relevante: modelo original, modelo depurado y modelo final. Esta práctica fortalece la transparencia del análisis, porque permite conservar evidencia verificable de las decisiones metodológicas tomadas durante la evaluación del modelo de medición reflectivo. Ver **Figura 5.12.** y **Figura 5.13**

Figura 5.12. Exportación del reporte del algoritmo PLS-SEM a Excel en SmartPLS4



Juan Mejía Tréjo

Figura 5.13. Muestra del archivo Excel en SmartPLS4



Fuente: Elaboración propia desde Excel.

Síntesis del procedimiento

- Ejecutar el algoritmo **PLS-SEM** desde **Calculate** → **PLS-SEM Algorithm**.
- Abrir **Open Full Report** para consultar los resultados completos.
- Ingresar a **Outer Loadings** para revisar las cargas externas.
- Clasificar los indicadores en tres grupos: cargas ≥ 0.70 , cargas entre 0.40 y 0.70 , y cargas < 0.40 .
- Revisar si los indicadores con cargas moderadas deben conservarse por relevancia teórica.
- Ingresar a **Construct Reliability and Validity – Overview**.
- Interpretar **Cronbach Alpha**, **rho_A** y **rho_c**.
- Confirmar que los valores de confiabilidad interna sean superiores a **0.70**.
- Revisar si algún valor se aproxima o supera **0.95**, por posible redundancia.
- Revisar la **AVE** como criterio de validez convergente.
- Identificar constructos con **AVE ≥ 0.50** y constructos cercanos al umbral.
- Registrar que **CAIL = 0.497** requiere revisión puntual
- Todo retiro de cargas se debe copiar el archivo para conservar el modelo original.
- Ingresar a **Discriminant Validity**.

Juan Mejía Trejo

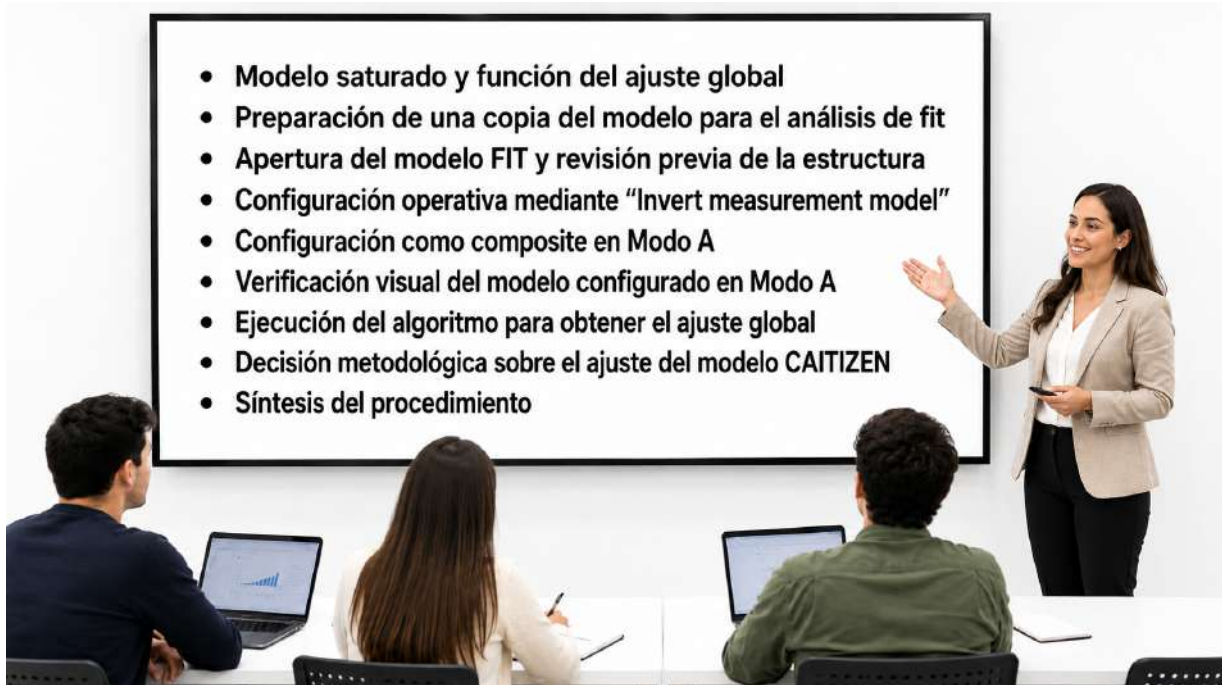
- Revisar **Fornell-Larcker Criterion**.
- Revisar **Cross Loadings** para confirmar que los indicadores cargan más alto en su propio constructo.
- Revisar **HTMT** y confirmar si los valores son menores a **0.85** o, al menos, menores a **0.90**.
- Integrar los resultados de cargas, confiabilidad, **AVE** y validez discriminante.
- Decidir si se conservan o eliminan indicadores con base estadística y teórica.
- Conservar el modelo reflectivo si la evidencia conjunta es suficiente.
- Exportar y guardar en archivo **Excel**.
- Avanzar al modelo estructural reportando cautelas metodológicas. Ver **Tabla 5.6**.

Tabla 5.6. Síntesis del procedimiento de evaluación del modelo de medición reflectivo

Paso	Criterio evaluado	Ruta o salida en SmartPLS4	Resultado en CAITIZEN
1	Ejecución del algoritmo	Calculate → PLS-SEM Algorithm	Modelo estimado correctamente
2	Cargas externas	Outer Loadings	Indicadores adecuados y algunos con cargas moderadas
3	Confiabilidad interna	Construct Reliability and Validity – Overview	Todos los constructos superan 0.70
4	Posible redundancia	ρ_c	HAIC $\rho_c = 0.950$
5	Validez convergente	AVE	Cinco constructos cumplen; CAIL = 0.497
6	Fornell-Larcker	Discriminant Validity – Fornell-Larcker	Puntos de atención en AFDJ–EAR , CAIL–AFDJ y CAITIZEN–HAIC
7	Cross Loadings	Discriminant Validity – Cross Loadings	Indicadores visibles cargan principalmente en su constructo
8	HTMT	Discriminant Validity – HTMT	Todos los valores < 0.85
9	Decisión integrada	Evaluación conjunta	Modelo confiable y viable con cautelas puntuales
10	Continuidad del análisis	Cierre del modelo de medición	Modelo listo para evaluación estructural

Fuente: Elaboración propia con base en Hair et al. (2021), Hair et al. (2019), Fornell y Larcker (1981) y Henseler et al. (2015).

CAPÍTULO 6. Ajuste global del modelo saturado



El **Capítulo 6. Ajuste global del modelo saturado**, desarrolla una fase metodológica complementaria dentro del análisis **PLS-SEM** aplicado. Después de evaluar el modelo de medición reflectivo, resulta necesario revisar si la estructura general del modelo reproduce de forma razonable las asociaciones observadas entre los constructos. Esta etapa no sustituye la revisión de cargas externas, confiabilidad, **AVE** o validez discriminante; más bien, funciona como un **control adicional de ajuste global** antes de avanzar hacia la interpretación estructural. En este sentido, el capítulo fortalece la ruta analítica al distinguir entre validación de medición, ajuste del modelo saturado y evaluación posterior de relaciones estructurales (Dijkstra & Henseler, 2015).

El capítulo aporta una distinción clave entre **modelo saturado** y **modelo estimado**. El modelo saturado considera las asociaciones posibles entre todos los constructos latentes, mientras que el modelo estimado representa únicamente las rutas planteadas por el investigador. En **CAITIZEN**, el modelo estimado mantiene la lógica de que **CAIL, EAR, AFDJ, HAIC y MTPP predicen a CAITIZEN**; en cambio, el modelo saturado permite valorar la correspondencia general entre la matriz empírica observada y la matriz reproducida por el modelo. Esta diferencia evita confundir el ajuste global con la comprobación de hipótesis, ya que ambas tareas pertenecen a momentos analíticos distintos (Henseler, 2017).

Una contribución relevante del capítulo es la **preparación técnica de una copia específica para el análisis de fit**. En lugar de modificar directamente el modelo principal, se conserva la versión depurada **CAITIZEN_PLS_SIN CAIL2** y se genera una copia denominada **CAITIZEN_PLS_SIN CAIL2_FIT**. Esta práctica asegura **trazabilidad metodológica**, porque permite distinguir la versión usada para el análisis ordinario del modelo de medición y estructural, de la versión operativa destinada al ajuste global. La gestión cuidadosa de versiones resulta especialmente importante cuando se trabaja con modelos que integran múltiples constructos, indicadores y decisiones de depuración (Becker et al., 2012).

El procedimiento también aclara la función de **Invert measurement model** en **SmartPLS4**. Esta opción cambia la orientación de las flechas del modelo de medición para que los indicadores apunten hacia el constructo. Sin embargo, no debe confundirse con una recodificación inversa de los datos ni con una transformación de la escala de respuesta. En este capítulo, la inversión tiene una finalidad operativa: configurar el modelo como **composite** para calcular el ajuste global del modelo saturado. Por tanto, el modelo **CAITIZEN** no se redefine conceptualmente como formativo, sino que se prepara técnicamente para una estimación compatible con el análisis de fit (Schuberth et al., 2020).

La selección de **Weighting Mode A** constituye otro punto metodológico central. **Mode A** se utiliza cuando los indicadores comparten una orientación conceptual común y mantienen una lógica correlacional cercana a la medición reflectiva. Esto resulta congruente con **CAITIZEN**, porque sus indicadores fueron diseñados para representar dimensiones coherentes como alfabetización crítica en **IA**, responsabilidad ética, justicia de datos, colaboración humano-**IA**, transparencia metacognitiva y ciudadanía sostenible asistida por **IA**. En contraste, **Mode B** corresponde mejor a modelos formativos o compuestos heterogéneos, donde los indicadores representan componentes diferenciados del constructo (Sarstedt et al., 2019).

Una vez configurado el modelo como **composite en Modo A**, el capítulo enfatiza la necesidad de una revisión visual rigurosa. Todos los constructos deben mostrar el símbolo A; el indicador **CAIL2** debe permanecer excluido; los indicadores restantes deben seguir asignados al constructo correcto; y las rutas estructurales deben conservar la lógica del modelo **CAITIZEN**. Esta verificación es importante porque cualquier error operativo puede generar resultados de ajuste que no correspondan al modelo depurado. Así, la revisión visual no es un trámite menor, sino una evidencia de **control, consistencia y reproducibilidad del procedimiento** (Gudergan et al., 2008).

Finalmente, el capítulo prepara la transición hacia la valoración del **ajuste global del modelo saturado** mediante indicadores como **SRMR, d_ULS y d_G**, que deberán interpretarse posteriormente con apoyo del procedimiento de Bootstrapping. Su propósito no es declarar que el modelo queda completamente validado, sino dejar una versión técnica, documentada y defendible para evaluar la discrepancia entre la matriz observada y la matriz reproducida. De este modo, el **Capítulo 6** funciona como puente

Juan Mejía Trejo

entre la evaluación del modelo de medición y la evaluación estructural, fortaleciendo la transparencia del análisis **PLS-SEM** aplicado al modelo **CAITIZEN** (Ringle et al., 2012).

Modelo saturado y función del ajuste global

Una vez evaluado el modelo de medición reflectivo del modelo **CAITIZEN** mediante cargas externas, confiabilidad interna, **AVE**, **Fornell-Larcker**, **Cross Loadings** y **HTMT**, resulta pertinente realizar una revisión complementaria del ajuste global del **modelo saturado**. Esta etapa permite valorar si la estructura general del modelo reproduce de manera razonable las relaciones observadas entre los constructos. No se trata todavía de probar hipótesis estructurales ni de interpretar la magnitud de las rutas entre variables latentes; su función principal es revisar si la matriz reproducida por el modelo presenta una discrepancia aceptable frente a la matriz empírica observada. Por ello, el ajuste global debe entenderse como una evaluación adicional de calidad del modelo, posterior al diagnóstico del modelo de medición y previa a la interpretación completa del modelo estructural (Hair et al., 2021). El **modelo saturado** es una especificación en la que se consideran todas las relaciones posibles entre los constructos latentes incluidos en el modelo. A diferencia del modelo estructural estimado, que solo incorpora las relaciones hipotéticas definidas por el investigador, el modelo saturado permite que todos los constructos se relacionen entre sí. Esta característica lo convierte en una referencia útil para valorar el ajuste global del sistema de medición, porque permite analizar la correspondencia general entre las asociaciones empíricas observadas y las asociaciones reproducidas por el modelo. En el caso **CAITIZEN**, el modelo saturado considera las asociaciones globales entre **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, sin limitarse únicamente a las rutas dirigidas hacia la variable dependiente (Henseler et al., 2015).

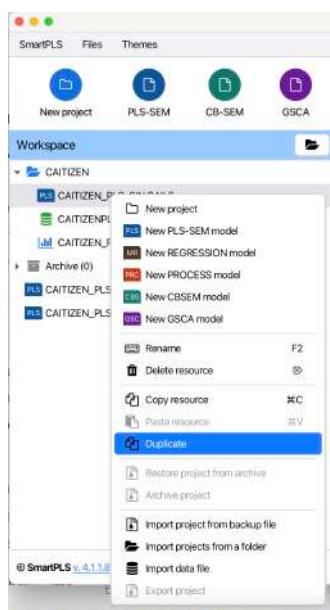
La diferencia entre **modelo saturado** y **modelo estimado** debe quedar clara para evitar interpretaciones equivocadas. El **modelo saturado** sirve para revisar el ajuste global del modelo de medición y de las relaciones posibles entre constructos. En cambio, el **modelo estimado** corresponde al modelo estructural que el investigador propone teóricamente. En el caso **CAITIZEN**, el modelo estimado plantea que **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** predicen a **CAITIZEN**. Por tanto, el modelo saturado no debe utilizarse para aceptar o rechazar hipótesis, sino para revisar si el modelo presenta una estructura global razonablemente compatible con los datos observados. La contrastación de hipótesis corresponde a una etapa posterior, mediante la evaluación del modelo estructural y el procedimiento de **Bootstrapping** (Hair et al., 2019). Desde el enfoque confirmatorio asociado a Henseler, el ajuste global del **modelo saturado** puede valorarse mediante criterios como **SRMR**, **d_ULS** y **d_G**. El **SRMR** permite revisar el ajuste aproximado del modelo, mientras que **d_ULS** y **d_G** permiten valorar discrepancias exactas a partir de límites obtenidos mediante **Bootstrap**. Estos criterios ayudan a determinar si el modelo presenta un ajuste aceptable antes de avanzar a la interpretación estructural. Sin embargo, deben reportarse con prudencia, porque un ajuste global favorable no sustituye la revisión de cargas externas, confiabilidad interna, **AVE**, validez discriminante ni colinealidad

estructural. El ajuste global es una pieza más dentro del proceso integral de evaluación del modelo **PLS-SEM** (Henseler et al., 2015). En consecuencia, la revisión del **modelo saturado** implica tres decisiones metodológicas. Primero, reconocer que se trata de una evaluación complementaria, no de una sustitución del análisis del modelo de medición. Segundo, distinguir que el modelo saturado revisa la estructura global de asociaciones entre constructos, pero no prueba las hipótesis estructurales. Tercero, preparar una versión específica del modelo en **SmartPLS4** para calcular correctamente los criterios de ajuste. Por ello, antes de interpretar los valores de **SRMR**, **d_ULS** y **d_G**, es necesario configurar una copia del modelo como **composite en Modo A**, mediante la opción **Invert measurement model** (Hair et al., 2021).

Preparación de una copia del modelo para el análisis de fit

Antes de modificar la orientación del modelo de medición, es indispensable generar una copia del modelo original. Esta práctica evita alterar la versión principal utilizada para la evaluación del modelo de medición y del modelo estructural. En el caso **CAITIZEN**, el modelo base ya había sido depurado previamente mediante la eliminación del indicador **CAIL2**, por lo que la versión de trabajo se identifica como **CAITIZEN_PLS_SIN CAIL2**. Para el análisis del ajuste global se debe duplicar esta versión y trabajar únicamente sobre la copia, de manera que el modelo original permanezca intacto y pueda consultarse nuevamente si fuera necesario (Hair et al., 2019). Ver **Figura 6.1**.

Figura 6.1. Duplicación del modelo depurado para el análisis de ajuste global.

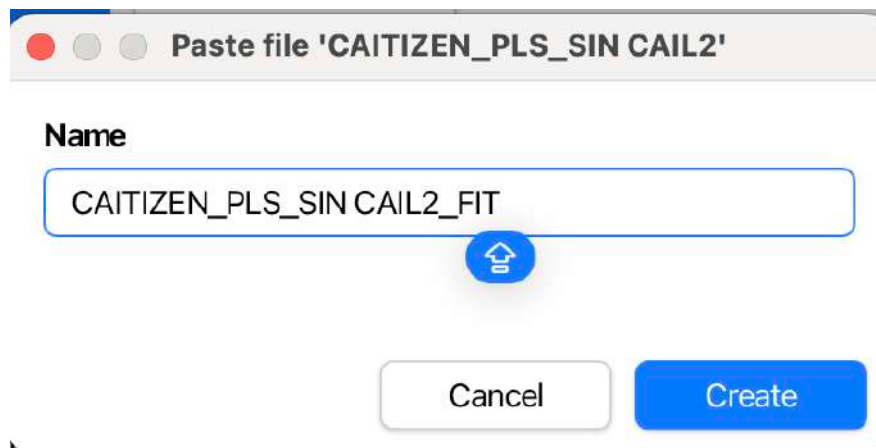


Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

En **SmartPLS4**, la duplicación se realiza desde el panel izquierdo del proyecto. Se selecciona el modelo correspondiente, se hace clic derecho y se elige la opción **Duplicate**. Esta acción crea un nuevo recurso dentro del proyecto, conservando la estructura del modelo original: constructos, indicadores, rutas estructurales y configuración previa. La ventaja de este procedimiento es que permite realizar ajustes técnicos para el cálculo del **fit** sin modificar la versión que será utilizada para reportar el análisis ordinario del modelo de medición y del modelo estructural. La trazabilidad de versiones es una buena práctica en análisis **PLS-SEM**, porque permite saber qué modelo produjo cada reporte (Hair et al., 2021). Al crear la copia, se recomienda asignar un nombre específico que indique su finalidad. En este caso, el nombre **CAITIZEN_PLS_SIN CAIL2_FIT** es adecuado porque comunica que se trata del modelo **CAITIZEN_PLS**, que corresponde a la versión sin **CAIL2**, y que se utilizará para la evaluación del **fit** o ajuste global. Esta denominación evita confusiones entre el modelo original, el modelo depurado y el modelo preparado para valorar el **modelo saturado**. El nombre de la copia debe mantenerse en todas las capturas, reportes exportados y tablas derivadas de esta corrida para conservar coherencia documental (Hair et al., 2019). Ver **Figura 6.2**

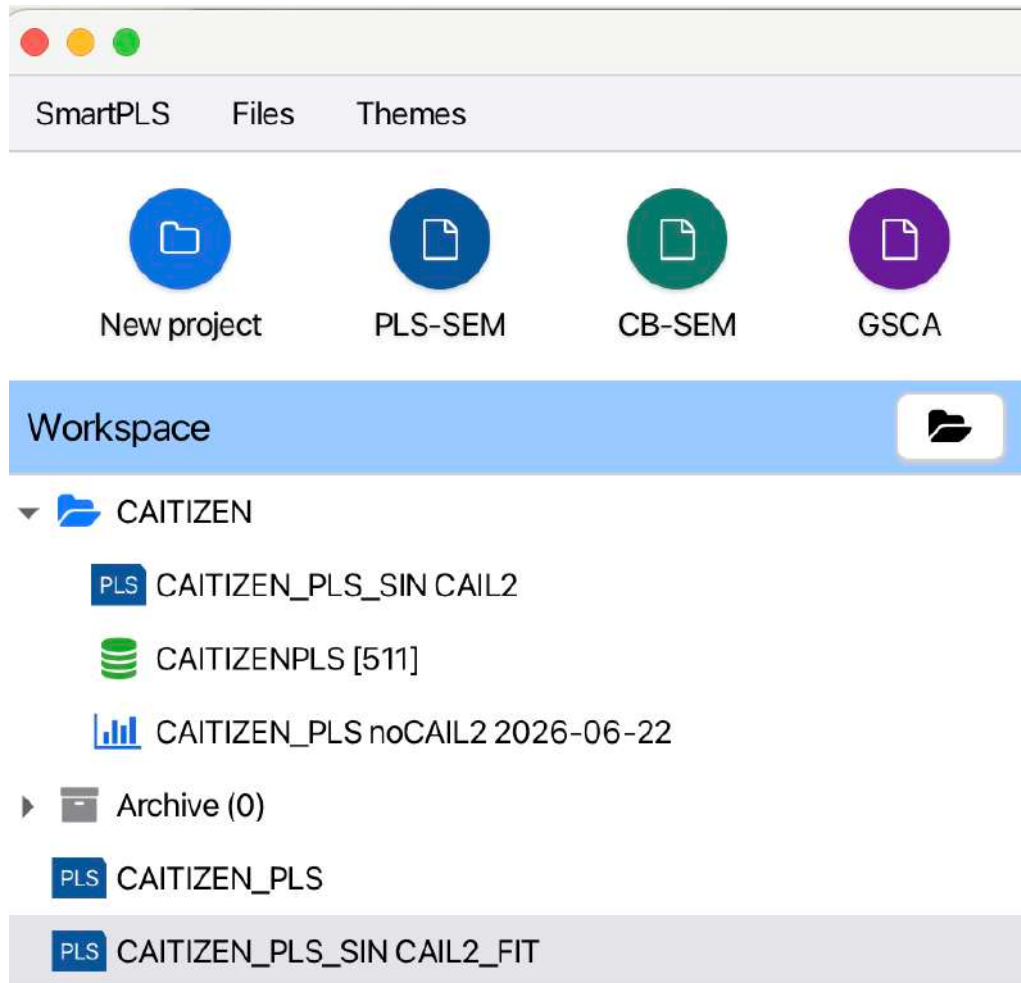
Figura 6.2. Asignación del nombre CAITIZEN_PLS_SIN CAIL2_FIT a la copia del modelo.



Fuente: Elaboración propia con **SmartPLS4**.

Una vez creada la copia, esta aparece en el árbol del proyecto como un modelo independiente. A partir de ese momento, cualquier modificación necesaria para calcular el ajuste global debe realizarse únicamente en la versión terminada en **_FIT**. La versión original o depurada se conserva para el análisis principal, mientras que la versión **_FIT** se reserva para la valoración del **modelo saturado**. Esta separación permite reportar con claridad que el ajuste global se calculó bajo una configuración operativa específica, sin mezclar sus resultados con los del modelo reflectivo evaluado en el capítulo anterior (Hair et al., 2021). Ver **Figura 6.3**.

Figura 6.3. Modelo **CAITIZEN_PLS_SIN CAIL2_FIT** creado dentro del proyecto **SmartPLS4**.



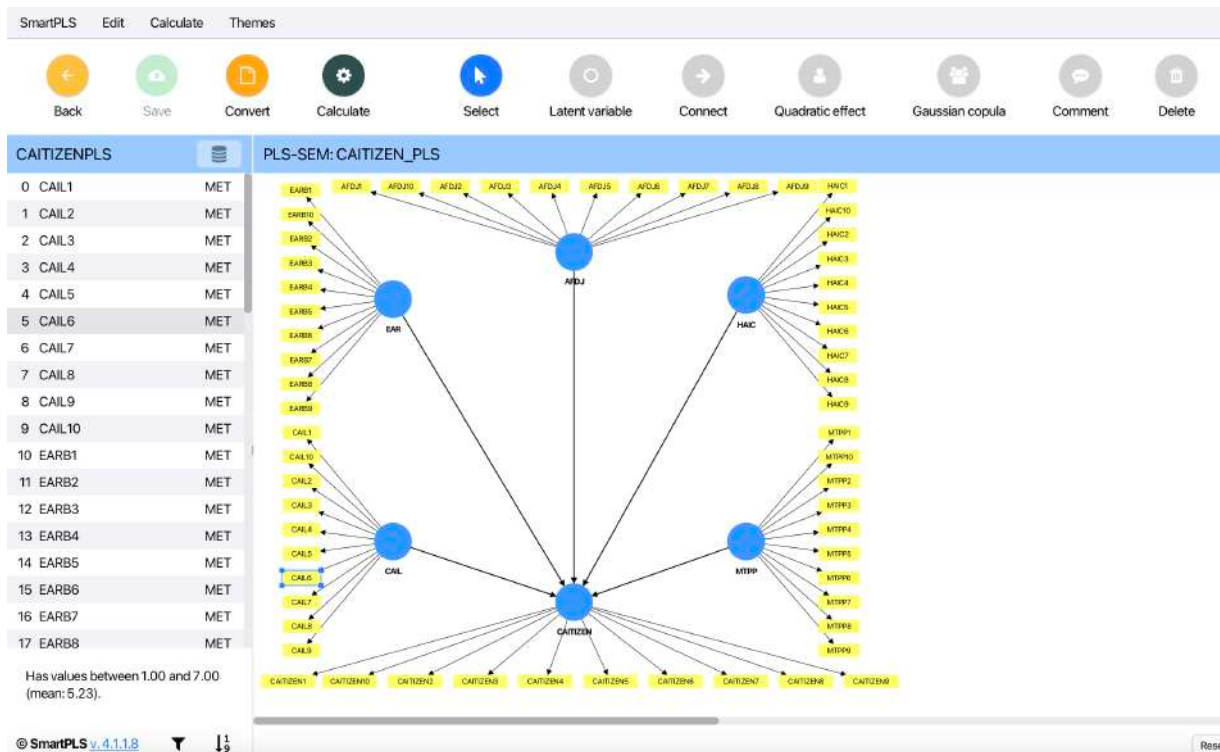
Fuente: Elaboración propia con **SmartPLS4**.

Apertura del modelo FIT y revisión previa de la estructura

Después de crear la copia, se debe abrir el modelo **CAITIZEN_PLS_SIN CAIL2_FIT** en el área gráfica de **SmartPLS4**. Esta revisión inicial permite confirmar que el modelo duplicado conserva la estructura esperada: los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**; los indicadores asignados a cada constructo; y las rutas estructurales dirigidas hacia **CAITIZEN**. Esta verificación es necesaria porque cualquier análisis de ajuste global debe realizarse sobre la versión exacta que se desea valorar. Si el modelo abierto no corresponde a la versión **_FIT**, o si los indicadores no coinciden con el modelo depurado, los resultados posteriores no serán interpretables con claridad (Hair et al., 2021) como se muestra en la **Figura 6.4**.

Juan Mejía Tréjo

Figura 6.4. Apertura del modelo CAITIZEN_PLS_SIN CAIL2_FIT en el área gráfica.



Fuente: Elaboración propia con **SmartPLS4**.

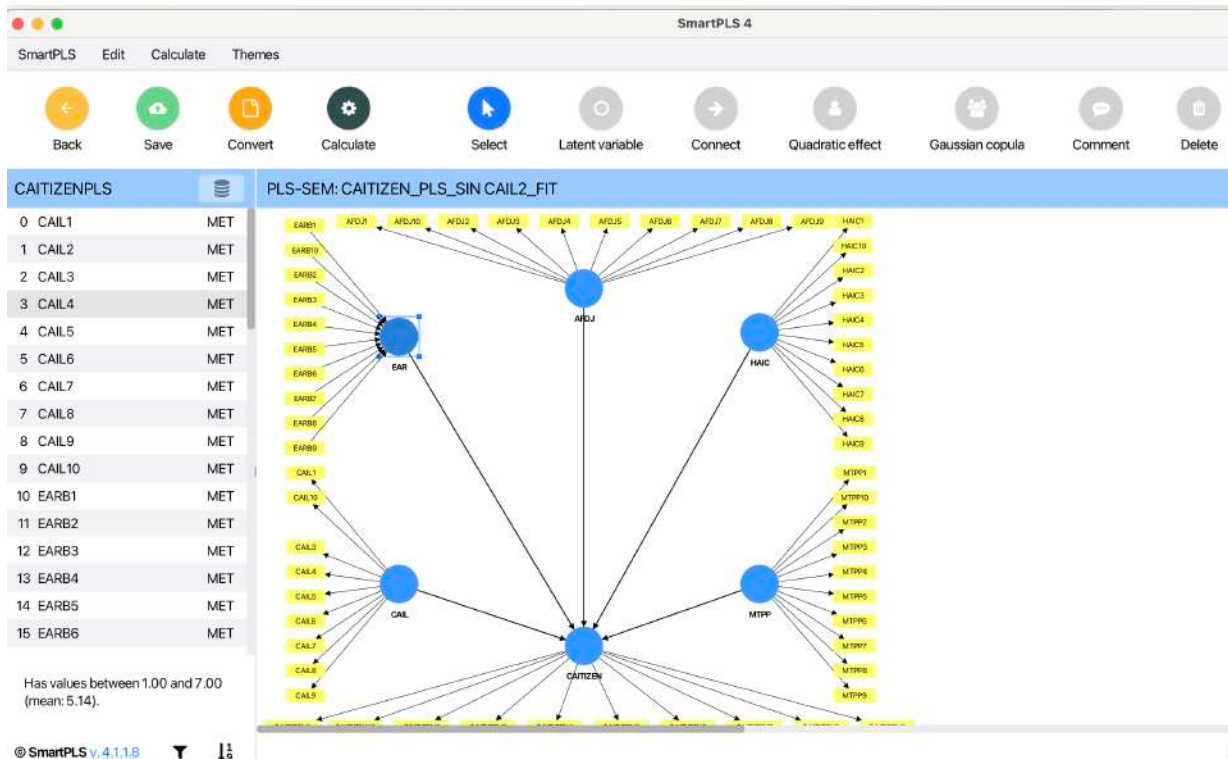
Antes de aplicar cualquier cambio técnico, debe observarse que la estructura del modelo siga siendo congruente con la propuesta conceptual. En el modelo **CAITIZEN**, los constructos antecedentes **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** se mantienen conectados hacia el constructo dependiente **CAITIZEN**. Además, debe revisarse que **CAIL2** no forme parte del constructo **CAIL**, si esta versión fue definida como modelo depurado sin ese indicador. Esta inspección permite distinguir entre una modificación metodológicamente controlada y un error accidental de edición (Hair et al., 2019).

Configuración operativa mediante “Invert measurement model”

Una vez abierta la copia **CAITIZEN_PLS_SIN CAIL2_FIT**, se debe **modificar la orientación** del modelo de medición. En **SmartPLS4 4.1.1.8**, esta operación se realiza seleccionando cada constructo y utilizando la opción **Invert measurement model**. Para ello, se debe hacer clic sobre el círculo azul del constructo correspondiente y posteriormente abrir el menú contextual con clic derecho. Esta operación debe realizarse constructo por constructo, porque cada variable latente debe quedar configurada de forma homogénea para la valoración del ajuste global (Henseler et al., 2015). Ver **Figura 6.5**.

Juan Mejía Trejo

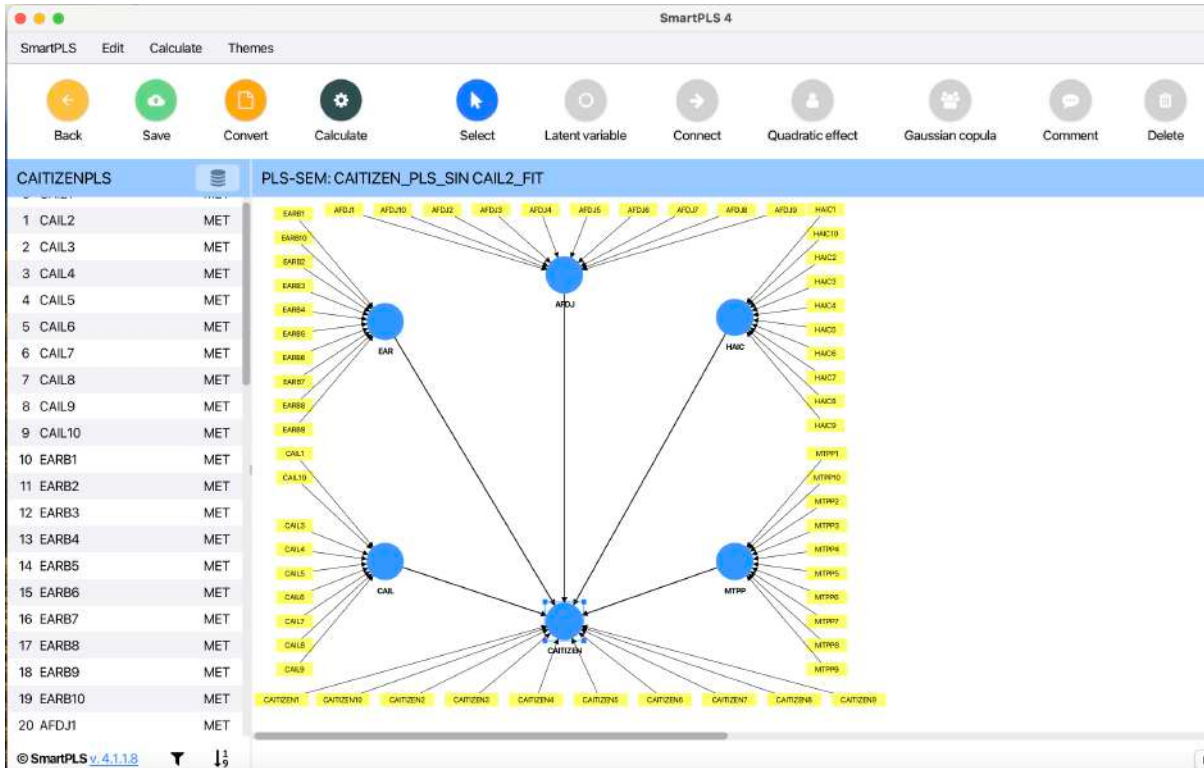
Figura 6.5. Selección del constructo para modificar la orientación del modelo de medición.



Fuente: Elaboración propia con **SmartPLS4**.

La opción ***Invert measurement model*** aparece dentro del menú contextual del constructo. Su función es invertir la dirección de las flechas del modelo de medición para que los indicadores se orienten hacia el constructo. Esta configuración permite tratar operativamente los constructos como ***composites***, lo cual es necesario para la valoración del ajuste global del **modelo saturado** desde el enfoque de **Henseler**. Esta acción debe repetirse en **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITZEN** (Henseler et al., 2015). Ver **Figura 6.6**.

Figura 6.6. Opción Invert measurement model en el menú contextual de SmartPLS4.



Fuente: Elaboración propia con **SmartPLS4**.

La inversión de flechas tiene una razón metodológica específica. En la lógica reflectiva tradicional, el constructo latente se concibe como causa común de sus indicadores; por ello, las flechas suelen representarse desde el constructo hacia los ítems. Sin embargo, para el análisis confirmatorio del **modelo saturado** desde el **enfoque de Henseler**, los constructos pueden configurarse operativamente como **composites**, de modo que los indicadores se orienten hacia el constructo. Esto permite que **SmartPLS4** genere resultados de ajuste global y de calidad del modelo bajo una especificación compatible con la valoración del **fit** (Henseler et al., 2015). Es importante aclarar al lector que **Invert measurement model** no significa invertir la escala de respuesta de los indicadores. Esta opción no transforma los datos, no cambia los valores de 1 a 7 ni recodifica las respuestas. La recodificación inversa de ítems solo procede cuando un indicador está redactado en sentido contrario al constructo que pretende medir. En cambio, **Invert measurement model** cambia la orientación del modelo de medición dentro del diagrama de **SmartPLS4**. Por ello, se trata de una decisión de especificación operativa del modelo, no de una modificación de la base de datos (Hair et al., 2021).

En términos visuales, antes de la inversión las flechas del modelo de medición se interpretan desde el constructo hacia los indicadores, como en la lógica reflectiva tradicional. Después de aplicar **Invert measurement model**, las flechas se orientan desde los indicadores hacia el constructo. Esta representación no debe leerse como una afirmación conceptual de que los indicadores “forman” sustantivamente el constructo, sino como una configuración técnica para calcular el ajuste global en **SmartPLS4**. La distinción entre conceptualización teórica y configuración operativa es clave para evitar contradicciones en el reporte (Hair et al., 2019).

Configuración como *composite* en Modo A

Después de aplicar **Invert measurement model**, al revisar la configuración del constructo en **Edit settings**, el campo **Measurement model** puede aparecer como **formative**. Esta situación puede generar confusión si se interpreta literalmente. En el contexto del análisis de ajuste, “formative” indica que el modelo fue invertido y que los indicadores se orientan hacia el constructo. Sin embargo, el propósito no es transformar conceptualmente el modelo **CAITIZEN** en un modelo formativo puro, sino configurarlo operativamente como **composite** para calcular el ajuste global del **modelo saturado** (Hair et al., 2021). Ver **Figura 6.7**.

Figura 6.7. Ventana de configuración del constructo después de invertir el modelo de medición.



Fuente: Elaboración propia con **SmartPLS4**.

Una vez que el constructo aparece con el modelo de medición invertido, debe revisarse el campo **Weighting mode**. En ese campo se selecciona **Mode A**. Esta selección es central para el procedimiento, porque el **Modo A** trabaja con una lógica basada en correlaciones entre indicadores y constructo, lo que lo hace compatible con

Juan Mejía Trejo

la valoración de **composites** cercanos a una lógica reflectiva. Por ello, para el análisis confirmatorio del **modelo saturado** se recomienda seleccionar **Mode A**, no **Mode B**, ni pesos iguales ni pesos predefinidos (Henseler et al., 2015). Ver **Figura 6.8**.

Figura 6.8. Selección de Weighting mode: Mode A en SmartPLS4.

The screenshot shows the 'Latent variable' settings window in SmartPLS4. The variable name 'EAR' is entered at the top. Under the 'Settings' tab, the 'Caption in report' is also 'EAR'. The 'Measurement model' dropdown menu is open, displaying two options: 'formative' and 'reflective'. The 'reflective' option is currently selected and highlighted with a blue background. Below this, the 'left' variable is listed, and the sorting is set to 'Alphabetic order'. The 'Margin' is set to '100' and the 'Folded' option is set to 'no'. At the bottom of the window, there are four buttons: 'Previous', 'Next', 'Cancel', and 'Apply'.

Fuente: Elaboración propia con **SmartPLS4**.

La selección de **Mode A** debe aplicarse a cada constructo del modelo. En el caso **CAITIZEN**, esto implica configurar en **Modo A** a **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. Después de elegir la opción correspondiente, debe presionarse **Apply** para guardar el cambio. Si el investigador omite este paso en alguno de los constructos, el modelo puede quedar configurado de forma incompleta. Por ello, conviene revisar uno por uno los constructos y confirmar que todos muestren la misma especificación antes de ejecutar nuevamente el algoritmo (Hair et al., 2021).

El **Modo A** sirve para estimar constructos compuestos cuando se espera que sus indicadores estén relacionados entre sí y compartan una orientación conceptual

Juan Mejía Trejo

común. En esta lógica, los indicadores no se tratan como piezas completamente independientes, sino como variables observables que se asocian con una misma dimensión subyacente. Esto resulta congruente con el modelo **CAITIZEN**, porque los indicadores de **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** fueron diseñados para representar dimensiones conceptuales coherentes. Por ello, aunque se configuren operativamente como **composites** para calcular el **fit**, conservan una lógica cercana a la medición reflectiva (Hair et al., 2019). La elección de **Mode A** permite que **SmartPLS4** calcule y reporte criterios necesarios para evaluar el modelo de medición y el ajuste global, como **Cronbach Alpha**, **rho_A**, **rho_c**, **AVE**, **HTMT**, **SRMR**, **d_ULS** y **d_G**. Esta es la razón práctica por la que se realiza la configuración en la versión **_FIT** del modelo. No se trata de alterar la teoría del modelo, sino de generar una especificación operativa compatible con la valoración confirmatoria del **modelo saturado** (Henseler et al., 2015).

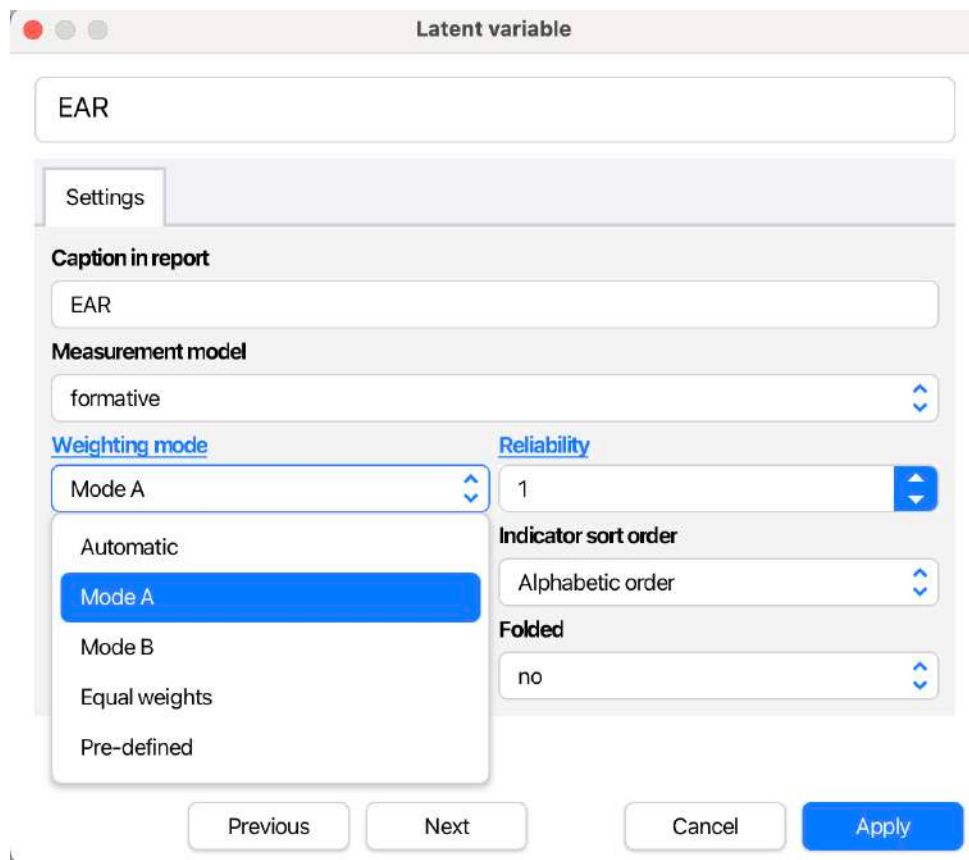
El **Modo B**, por su parte, responde a una lógica distinta. En **Mode B**, los pesos de los indicadores se estiman mediante una lógica similar a la regresión múltiple, donde los indicadores contribuyen a formar el constructo y no necesariamente se espera que estén altamente correlacionados. Esta opción suele ser más pertinente para modelos formativos o compuestos en los que cada indicador representa una dimensión diferente y no intercambiable del constructo. Por ejemplo, si un constructo estuviera formado por componentes heterogéneos como infraestructura, presupuesto, personal especializado y recursos tecnológicos, podría tener sentido usar **Mode B**, porque cada indicador aportaría una parte distinta del concepto total (Hair et al., 2021).

La diferencia entre **Mode A** y **Mode B** puede resumirse de forma metodológica. **Mode A** es apropiado cuando los indicadores se comportan como manifestaciones correlacionadas de una misma dimensión y se desea una estimación cercana a la lógica reflectiva. **Mode B** es apropiado cuando los indicadores funcionan como componentes explicativos diferenciados que forman el constructo y no necesariamente son intercambiables. Por esta razón, para el análisis de ajuste global del **modelo saturado CAITIZEN**, se selecciona **Mode A**, ya que los constructos provienen de una conceptualización reflectiva y se configuran como **composites** únicamente para efectos de estimación y cálculo del **fit** (Henseler et al., 2015). En consecuencia, invertir las flechas y seleccionar **Mode A** sirve para preparar técnicamente el modelo para el análisis de ajuste global. Esta decisión debe aplicarse a todos los constructos incluidos en la valoración del modelo saturado: **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. Si solo algunos constructos fueran configurados en **Modo A** y otros permanecieran con otra especificación, el análisis perdería homogeneidad y podría generar resultados difíciles de interpretar. Por ello, la configuración debe realizarse de manera uniforme, constructo por constructo, siempre dentro de una copia del modelo original o depurado (Hair et al., 2019).

Verificación visual del modelo configurado en Modo A

Después de configurar todos los constructos como **composites en Modo A**, es necesario revisar visualmente el modelo completo. En **SmartPLS4**, una señal útil es la aparición del símbolo **A** dentro del círculo azul de cada constructo. Este símbolo indica que el constructo fue configurado con **Weighting Mode A**. En los screenshots se observa que los constructos **AFDJ**, **EAR**, **CAIL**, **HAIC**, **MTPP** y **CAITIZEN** muestran el símbolo **A**, lo cual confirma que la configuración fue aplicada de manera uniforme (Henseler et al., 2015). Ver **Figura 6.9**.

Figura 6.9. Modelo CAITIZEN con constructos configurados en Modo A.



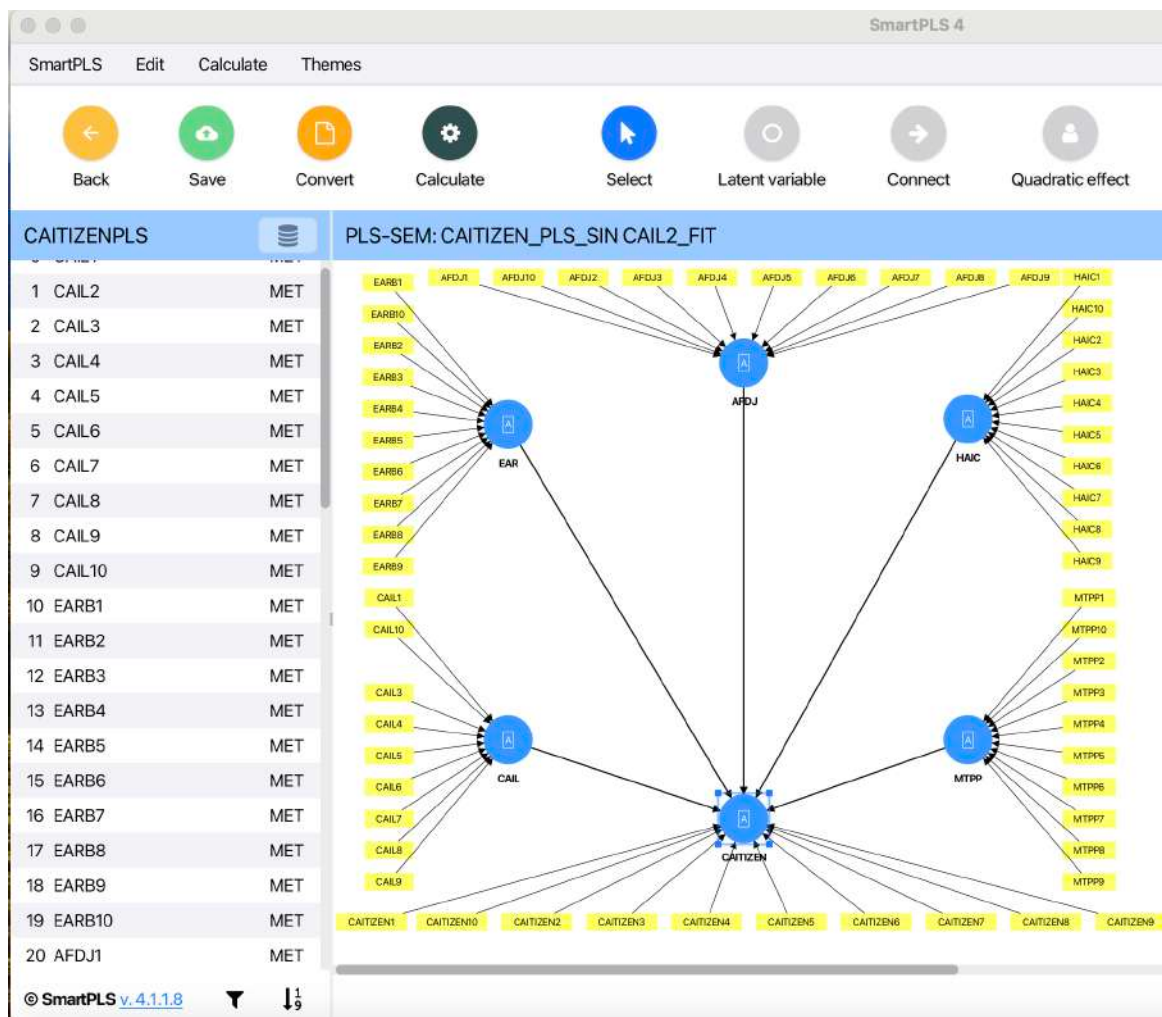
Fuente: Elaboración propia con **SmartPLS4**.

La verificación visual también permite comprobar que los indicadores permanecen asignados al constructo correcto después de haber configurado el modelo para el análisis de ajuste global. Esta revisión es necesaria porque, al duplicar el modelo, invertir el modelo de medición y configurar los constructos como **composites en Modo A**, el investigador puede modificar accidentalmente la estructura de medición. En el caso del modelo **CAITIZEN_PLS_SIN CAIL2_FIT**, debe confirmarse que el indicador

Juan Mejía Tréjo

CAIL2 ya no forme parte del bloque de indicadores de **CAIL**, debido a que esta versión corresponde al modelo depurado. Al mismo tiempo, debe verificarse que los indicadores restantes de **CAIL**, así como los indicadores de **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, continúen vinculados a sus constructos correspondientes. Esta revisión garantiza que el ajuste global se calcule sobre el mismo modelo que fue decidido metodológicamente en la etapa de depuración. Si un indicador eliminado aparece nuevamente, si un indicador queda sin asignar o si un ítem se vincula al constructo incorrecto, el reporte de **SRMR**, **d_ULS** y **d_G** no representaría el modelo previsto, sino una versión alterada del análisis. Por tanto, antes de ejecutar nuevamente el algoritmo, la pantalla del modelo debe revisarse como evidencia de control metodológico, trazabilidad y correspondencia entre la versión depurada y la versión preparada para el análisis de **fit** (Hair et al., 2019). Ver **Figura 6.10**.

Figura 6.10. Verificación general del modelo CAITIZEN_PLS_SIN CAIL2_FIT antes del cálculo.

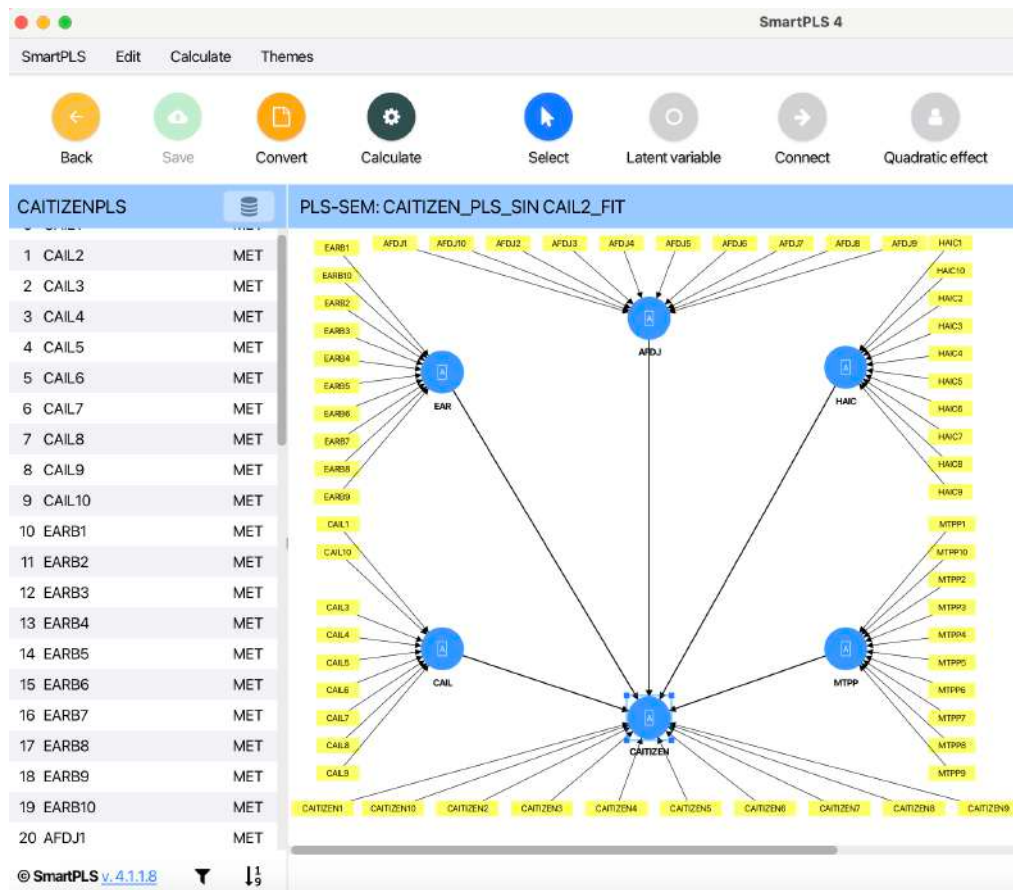


Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Trejo

También debe verificarse que las rutas estructurales del modelo se mantengan de acuerdo con el planteamiento teórico. Aunque el **modelo saturado** se utiliza para valorar el ajuste global, la versión gráfica del modelo debe conservar la estructura básica del caso **CAITIZEN**, donde **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTTP** se dirigen hacia **CAITIZEN**. Esta revisión ayuda a confirmar que la copia **_FIT** no sufrió modificaciones accidentales en sus relaciones estructurales durante el proceso de configuración operativa (Hair et al., 2021). La verificación final debe realizarse antes de ejecutar nuevamente el algoritmo. Si todos los constructos muestran el símbolo **A**, si las flechas de medición se orientan desde los indicadores hacia los constructos, si el indicador **CAIL2** permanece excluido y si las rutas estructurales conservan la lógica del modelo, entonces la versión **CAITIZEN_PLS_SIN CAIL2_FIT** está lista para calcular el ajuste global del **modelo saturado**. Esta revisión reduce errores de procedimiento y fortalece la reproducibilidad del análisis (Hair et al., 2019). Ver **Figura 6.11**

Figura 6.11. Modelo CAITIZEN preparado para ejecutar el algoritmo PLS-SEM y calcular el ajuste global.



Fuente: Elaboración propia con **SmartPLS4**.

Ejecución del algoritmo para obtener el ajuste global

Una vez configurado el modelo en **Modo A**, debe guardarse la versión modificada y ejecutarse nuevamente el algoritmo desde **Calculate** → **PLS-SEM Algorithm**. Esta ejecución no corresponde a la evaluación ordinaria del modelo reflectivo realizada en el Capítulo 5, sino a una corrida específica para valorar el ajuste global del **modelo saturado**. Por esa razón, el reporte generado debe exportarse y guardarse con un nombre que permita reconocer que pertenece a la versión **CAITIZEN_PLS_SIN CAIL2_FIT** (Hair et al., 2021). Después de ejecutar el algoritmo, se debe abrir el reporte completo mediante **Open full report**. Dentro del reporte, la sección relevante para esta etapa es la correspondiente al ajuste del modelo, donde **SmartPLS4** muestra criterios como **SRMR**, **d_ULS** y **d_G**. Estos resultados deben interpretarse como indicadores de discrepancia entre la matriz observada y la matriz reproducida por el modelo. Mientras menor sea la discrepancia, mayor será la evidencia de ajuste global aceptable (Henseler et al., 2015). El **SRMR** o **Standardized Root Mean Square Residual** es el índice de ajuste aproximado más utilizado en esta etapa. Un valor menor a **0.08** suele considerarse evidencia de buen ajuste bajo un criterio estricto. Un valor menor a **0.10** puede considerarse aceptable bajo un criterio más flexible, especialmente en modelos aplicados de ciencias sociales donde los constructos suelen estar conceptualmente relacionados. En el modelo **CAITIZEN**, este criterio debe interpretarse junto con los resultados del Capítulo 5, no de forma aislada (Hu & Bentler, 1999).

Además del **SRMR**, se revisan las medidas de discrepancia **d_ULS** y **d_G**. Estas medidas permiten realizar una valoración de ajuste exacto mediante comparación con límites obtenidos por **Bootstrap**, especialmente **HI95** y **HI99**. La lógica es que el valor original de **d_ULS** y **d_G** debe ser menor que los valores críticos generados por el procedimiento bootstrap. Si esto ocurre, el modelo presenta evidencia favorable de ajuste desde esta perspectiva confirmatoria (Henseler et al., 2015). La interpretación de estos criterios debe ser prudente. Un **SRMR** adecuado y valores aceptables en **d_ULS** y **d_G** respaldan el ajuste global del **modelo saturado**, pero no prueban por sí mismos que las hipótesis estructurales estén apoyadas. Las hipótesis se evaluarán posteriormente mediante coeficientes de trayectoria, tamaños de efecto y **Bootstrapping**. Por tanto, el ajuste global debe entenderse como una condición favorable para continuar el análisis, no como conclusión final sobre el modelo completo (Hair et al., 2019).

Decisión metodológica sobre el ajuste del modelo CAITIZEN

La decisión metodológica sobre el ajuste global debe integrar tres elementos: la configuración operativa del modelo, los criterios obtenidos en el reporte y la evaluación previa del modelo de medición. Si el modelo fue correctamente configurado como **composite en Modo A**, si todos los constructos muestran el símbolo **A** y si los valores de **SRMR**, **d_ULS** y **d_G** se ubican dentro de los umbrales aceptables, puede señalarse que el **modelo saturado** presenta ajuste global razonable. Esta conclusión respalda la continuidad hacia la evaluación del modelo estructural (Henseler et al., 2015). La redacción del reporte debe evitar afirmaciones absolutas. No conviene decir que el modelo queda “validado completamente” solo porque el ajuste global sea aceptable. Una formulación más adecuada es señalar que el ajuste del **modelo saturado** ofrece evidencia complementaria favorable, consistente con la evaluación previa del modelo de medición. Esta forma de reportar mantiene prudencia metodológica y evita convertir un índice de ajuste en una conclusión total sobre la calidad del modelo (Hair et al., 2021).

En el caso **CAITIZEN**, la decisión debe vincularse con lo ya observado en el Capítulo 5. El modelo mostró confiabilidad interna adecuada, validez convergente aceptable en la mayoría de los constructos, **HTMT** por debajo del umbral estricto y algunos puntos de atención en **Fornell-Larcker**. Por ello, si el ajuste global del **modelo saturado** resulta aceptable, la conclusión debe indicar que la evidencia conjunta permite continuar hacia la evaluación estructural, aunque manteniendo las cautelas previamente identificadas en **CAIL**, **AFDJ-EAR**, **CAIL-AFDJ** y **CAITIZEN-HAIC** (Hair et al., 2019). Así, el Capítulo 6 cumple una función de puente metodológico. El Capítulo 5 evaluó la calidad del modelo de medición; el Capítulo 6 revisa el ajuste global del **modelo saturado**; y el Capítulo 7 puede avanzar hacia la evaluación del modelo estructural. Esta secuencia permite que el análisis **PLS-SEM** del modelo **CAITIZEN** se desarrolle de manera ordenada, transparente y defendible. La finalidad no es solo generar resultados en **SmartPLS4**, sino explicar qué configuración se utilizó, por qué se utilizó y cómo debe interpretarse dentro de la ruta completa del análisis estructural aplicado (Hair et al., 2021).

Síntesis del procedimiento

El procedimiento para preparar el **modelo saturado** inicia después de evaluar el modelo de medición reflectivo. Primero, se conserva intacta la versión principal del modelo **CAITIZEN_PLS_SIN CAIL2**, ya que esta contiene la estructura depurada del modelo de medición. Para evitar alterar el análisis original, se genera una copia específica destinada al ajuste global, denominada **CAITIZEN_PLS_SIN CAIL2_FIT**. Esta decisión permite separar el modelo utilizado para la evaluación reflectiva ordinaria del modelo configurado técnicamente para valorar el **fit** del **modelo saturado**. Una vez creada la copia, se abre el modelo **CAITIZEN_PLS_SIN CAIL2_FIT** en el área gráfica

Juan Mejía Trejo

de **SmartPLS4** y se revisa que conserve la estructura esperada. Deben verificarse los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITIZEN**; sus indicadores correspondientes; la exclusión del indicador **CAIL2**; y las rutas estructurales dirigidas hacia **CAITIZEN**. Esta revisión garantiza que el análisis de ajuste se realice sobre la versión correcta del modelo depurado.

Posteriormente, se modifica la orientación del modelo de medición mediante la opción **Invert measurement model**. Esta acción se aplica constructo por constructo y permite que los indicadores se orienten hacia el constructo. Es importante aclarar que esta operación no recodifica los datos ni invierte la escala de respuesta de los ítems; únicamente cambia la especificación operativa del modelo de medición en **SmartPLS4**. Su finalidad es preparar los constructos para ser tratados como **composites** dentro de la valoración del **modelo saturado**. Después de invertir el modelo de medición, cada constructo puede aparecer en la configuración como **formative**. Esta etiqueta debe interpretarse como una configuración operativa del software, no como una redefinición conceptual del modelo **CAITIZEN**. El propósito no es transformar teóricamente los constructos en formativos, sino estimarlos como **composites** para calcular el ajuste global del modelo. Por ello, en el campo **Weighting mode** debe seleccionarse **Mode A** para todos los constructos.

La elección de **Mode A** se justifica porque permite estimar **composites** bajo una lógica correlacional, cercana a la medición reflectiva. En el caso **CAITIZEN**, los indicadores de cada constructo fueron diseñados para representar dimensiones conceptualmente coherentes; por tanto, se espera que compartan una base común. En contraste, **Mode B** se reserva para constructos formativos o compuestos heterogéneos, donde los indicadores representan componentes distintos y no necesariamente correlacionados. Por esta razón, para valorar el **modelo saturado CAITIZEN**, se selecciona **Mode A** y no **Mode B**. Después de configurar todos los constructos en **Modo A**, se realiza una verificación visual del modelo. Debe confirmarse que cada constructo muestre el símbolo **A**, que las flechas de medición se orienten desde los indicadores hacia los constructos, que **CAIL2** permanezca excluido y que los demás indicadores sigan vinculados a sus constructos correctos. Esta revisión evita errores operativos y asegura que el modelo preparado para el análisis de **fit** corresponda exactamente a la versión depurada decidida metodológicamente.

Finalmente, cuando el modelo **CAITIZEN_PLS_SIN CAIL2_FIT** ha sido duplicado, revisado, invertido, configurado en **Modo A** y verificado visualmente, queda listo para ejecutar el análisis inferencial correspondiente. El siguiente paso ya no pertenece a la preparación del modelo, sino al **Complete Bootstrapping**, que debe desarrollarse en el Capítulo 7. Mediante ese procedimiento se podrán calcular e interpretar los indicadores de bondad de ajuste del **modelo saturado**, particularmente **SRMR**, **d_ULS** y **d_G**, junto con sus límites **HI95** y **HI99**. Así, se tiene:

- **Conservar el modelo original.** Antes de realizar cualquier ajuste técnico, debe mantenerse intacta la versión principal del modelo **CAITIZEN_PLS_SIN CAIL2**, porque esta representa el modelo depurado que fue evaluado previamente en

Juan Mejía Trejo

términos de cargas externas, confiabilidad interna, **AVE**, **Fornell-Larcker**, **Cross Loadings** y **HTMT**.

- **Duplicar el modelo para el análisis de fit.** La valoración del **modelo saturado** no debe realizarse directamente sobre el modelo principal. Por ello, se debe crear una copia específica mediante la opción **Duplicate** en **SmartPLS4**.
- **Nombrar la copia con trazabilidad.** La copia debe identificarse claramente como **CAITIZEN_PLS_SIN CAIL2_FIT**, ya que este nombre indica que se trata del modelo **CAITIZEN_PLS**, sin el indicador **CAIL2**, preparado exclusivamente para el análisis de fit.
- **Abrir la versión correcta del modelo.** Antes de modificar cualquier elemento, debe confirmarse que el modelo abierto sea efectivamente **CAITIZEN_PLS_SIN CAIL2_FIT** y no la versión original o depurada destinada al análisis ordinario.
- **Revisar la estructura inicial del modelo.** Deben verificarse los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, así como la correcta asignación de sus indicadores y la permanencia de las rutas estructurales dirigidas hacia **CAITIZEN**.
- **Confirmar la exclusión de CAIL2.** En la versión **CAITIZEN_PLS_SIN CAIL2_FIT**, el indicador **CAIL2** no debe aparecer dentro del bloque de indicadores de **CAIL**, porque esta versión corresponde al modelo depurado previamente.
- **Aplicar *Invert measurement model*.** Cada constructo debe seleccionarse individualmente y configurarse mediante la opción **Invert measurement model**, disponible al hacer clic derecho sobre el círculo azul del constructo.
- **Entender qué significa invertir el modelo de medición.** La opción ***Invert measurement model*** cambia la dirección de las flechas del modelo de medición para que los indicadores se orienten hacia el constructo. Esta acción permite configurar operativamente los constructos como **composites**.
- **No confundir con recodificación inversa.** Invertir el modelo de medición no significa cambiar la escala de respuesta. No transforma 1 en 7, ni 2 en 6. La recodificación inversa solo se aplica a ítems redactados en sentido contrario al constructo.
- **Aplicar la inversión a todos los constructos.** La configuración debe realizarse en **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, para mantener una especificación homogénea del modelo saturado.
- **Interpretar correctamente la etiqueta formative.** Después de invertir el modelo de medición, **SmartPLS4** puede mostrar el constructo como **formative**. Esta etiqueta debe entenderse como una configuración operativa del software, no como una redefinición teórica del modelo **CAITIZEN**.
- **Seleccionar *Weighting mode: Mode A*.** En la ventana de configuración de cada constructo debe seleccionarse **Mode A**, porque este modo estima **composites** bajo una lógica correlacional, cercana a la medición reflectiva.

- **Evitar Mode B para este caso.** **Mode B** se reserva para constructos formativos o compuestos heterogéneos, donde los indicadores representan dimensiones diferentes y no necesariamente correlacionadas. Por ello, no es la opción recomendada para el modelo **CAITIZEN**.
- **Guardar los cambios en cada constructo.** Después de seleccionar **Mode A**, debe presionarse **Apply** para asegurar que la configuración quede registrada correctamente en cada constructo.
- **Verificar el símbolo A.** Una vez configurados los constructos, debe revisarse visualmente que cada uno muestre el símbolo **A** dentro del círculo azul. Este símbolo confirma que el constructo fue configurado con **Weighting Mode A**.
- **Revisar la asignación de indicadores.** Debe comprobarse que los indicadores restantes de **CAIL**, así como los indicadores de **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**, permanezcan vinculados a sus constructos correspondientes.
- **Evitar errores de estructura.** Si un indicador eliminado reaparece, si un indicador queda sin asignar o si un ítem se vincula al constructo incorrecto, los resultados del ajuste global no corresponderán al modelo previsto.
- **Confirmar las rutas estructurales.** Aunque el análisis se enfoca en el **modelo saturado**, la representación gráfica debe conservar la lógica teórica del modelo **CAITIZEN**, donde **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** se dirigen hacia **CAITIZEN**.
- **Usar la pantalla como evidencia metodológica.** La revisión visual del modelo configurado en **Modo A** debe utilizarse como evidencia de control, trazabilidad y correspondencia entre el modelo depurado y la versión preparada para el análisis de *fit*.
- **No interpretar todavía los indicadores de bondad de ajuste.** En el **Capítulo 6** solo se prepara el modelo. La interpretación de **SRMR**, **d_ULS**, **d_G**, **HI95** y **HI99** debe desarrollarse en el **Capítulo 7** mediante **Complete Bootstrapping**.
- **Transición metodológica.** Una vez que el modelo **CAITIZEN_PLS_SIN_CAIL2_FIT** ha sido duplicado, revisado, invertido, configurado en **Modo A** y verificado visualmente, queda listo para ejecutar **Complete Bootstrapping** en el siguiente capítulo.
- **Decisión final del procedimiento.** Si todos los constructos muestran el símbolo **A**, si **CAIL2** permanece excluido, si los indicadores están correctamente asignados y si las rutas estructurales conservan la lógica del modelo, entonces el modelo está preparado para calcular el ajuste global del **modelo saturado**.

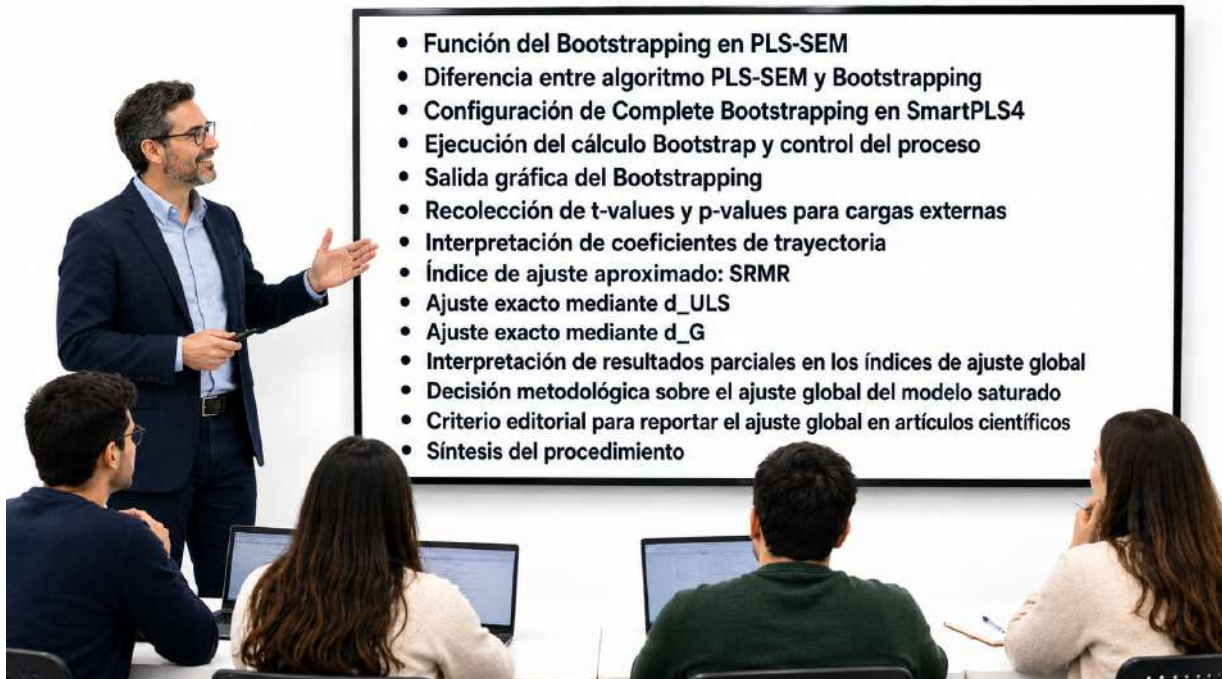
Resultado esperado del Capítulo 6. Al finalizar este capítulo, el lector debe comprender que la preparación del modelo saturado es una etapa técnica previa al análisis inferencial. Su propósito es dejar el modelo listo para que, en el **Capítulo 7**, el **Complete Bootstrapping** permita valorar la bondad de ajuste mediante **SRMR**, **d_ULS**, **d_G**, **HI95** y **HI99**.

CAPÍTULO 6. Ajuste global del modelo saturado

De esta forma, en el modelo **CAITIZEN_PLS_SIN CAIL2_FIT** se utiliza únicamente para evaluar el **ajuste global del modelo saturado** mediante **SRMR**, **d_ULS** y **d_G**. **No debe utilizarse para reportar cargas externas, confiabilidad, AVE, HTMT, VIF, coeficientes de trayectoria, R^2 , f^2 , hipótesis ni *PLSpredict*.** Todos esos resultados deben obtenerse del modelo normal final depurado **CAITIZEN_PLS_SIN CAIL2**.

Juan Mejía Tréjo

CAPÍTULO 7. *Bootstrapping* e inferencia del modelo



El **Capítulo 7. *Bootstrapping* e inferencia del modelo**, introduce una fase inferencial decisiva dentro del análisis **PLS-SEM** aplicado. Después de preparar el modelo saturado en **SmartPLS4**, el ***Bootstrapping*** permite evaluar la **estabilidad estadística de los parámetros estimados** mediante remuestreo con reemplazo. A diferencia del algoritmo **PLS-SEM**, que calcula los valores originales del modelo a partir de la muestra disponible, el **Bootstrapping** genera distribuciones empíricas para obtener **errores estándar, t-values, p-values e intervalos de confianza**. Esta etapa resulta indispensable porque permite distinguir entre estimaciones descriptivas y evidencia inferencial para valorar cargas, rutas y criterios de ajuste (Hair et al., 2021).

El capítulo establece una distinción metodológica importante entre el uso del **Bootstrapping** para **contrastar hipótesis** y su uso para **evaluar el ajuste global del modelo saturado**. Para contrastar cargas externas, coeficientes de trayectoria y relaciones estructurales, el procedimiento debe ejecutarse sobre el modelo normal final depurado, identificado como **CAITIZEN_PLS_SIN CAIL2**. En cambio, para valorar el ajuste global mediante **SRMR, d_ULS y d_G**, el **Bootstrapping** debe realizarse sobre la copia **CAITIZEN_PLS_SIN CAIL2 FIT**, previamente configurada como composite en Modo A. Esta separación evita mezclar resultados de versiones distintas y fortalece la **trazabilidad metodológica** del análisis (Henseler et al., 2015)

Juan Mejía Trejo

La diferencia entre **PLS-SEM Algorithm** y **Bootstrapping** es central para comprender el capítulo. El algoritmo responde cuál es el valor estimado de un parámetro en la muestra original; el **Bootstrapping** responde qué tan estable, significativo o confiable resulta ese parámetro frente a variaciones muestrales. Por ejemplo, una ruta estructural puede mostrar un coeficiente positivo, pero solo el Bootstrapping permite conocer si dicho efecto es estadísticamente significativo mediante su t-value, p-value e intervalo de confianza. Por ello, el capítulo enseña que la interpretación rigurosa del modelo no debe apoyarse únicamente en los coeficientes originales, sino en su estabilidad inferencial (Hair et al., 2019).

La configuración del **Complete Bootstrapping** en **SmartPLS4** constituye otro aporte operativo del capítulo. Se recomienda utilizar **5,000 submuestras**, resultados completos, prueba de dos colas y nivel de significancia de **0.05**. La opción Complete permite obtener una salida amplia que incluye rutas, cargas, pesos, confiabilidad, **HTMT**, **R²** y criterios de ajuste del modelo. Esta configuración es especialmente pertinente en un libro metodológico aplicado, porque no solo interesa probar hipótesis, sino documentar con profundidad el comportamiento inferencial del modelo **CAITIZEN**. La prueba de dos colas fortalece además una lectura conservadora de los parámetros estimados (Ringle et al., 2012).

Los resultados de Bootstrapping permiten interpretar los **coeficientes de trayectoria** del modelo **CAITIZEN**. En la salida reportada, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN** presentan relaciones positivas y estadísticamente significativas. En contraste, **CAIL** → **CAITIZEN** y **EAR** → **CAITIZEN** no muestran significancia directa bajo el criterio de $p < 0.05$. Esta lectura no implica que **CAIL** o **EAR** carezcan de valor conceptual, sino que, controlando por los demás predictores, no presentan un efecto directo significativo sobre **CAITIZEN** en esta estimación. Por tanto, el capítulo enseña a diferenciar entre **relevancia teórica** y **significancia estadística directa** (Guenther et al., 2023).

El capítulo también examina el **ajuste global del modelo saturado** mediante **SRMR**, **d_ULS** y **d_G**. El **SRMR** = 0.047 se ubica por debajo de los umbrales convencionales de 0.08 y 0.10, lo que sugiere un ajuste aproximado favorable. Sin embargo, al compararlo con los límites Bootstrap **HI95** y **HI99**, así como con los resultados de **d_ULS** y **d_G**, se observa que el modelo no cumple plenamente los criterios de ajuste exacto. Esta diferencia obliga a reportar una conclusión prudente: el modelo presenta **ajuste aproximado favorable**, pero **ajuste exacto limitado**, por lo que no debe afirmarse que el modelo queda validado de manera plena (Henseler, 2017).

En síntesis, el capítulo aporta una ruta metodológica para ejecutar, interpretar y reportar el **Complete Bootstrapping** en **SmartPLS4** con rigor académico. Su contribución consiste en mostrar que el **Bootstrapping** no es un trámite final, sino una herramienta para evaluar la **estabilidad inferencial**, la significancia de rutas e indicadores y los criterios de ajuste global del modelo saturado. En el caso **CAITIZEN**, los resultados permiten continuar hacia la evaluación estructural, pero bajo una postura

Juan Mejía Trejo

metodológica equilibrada: reconocer los efectos significativos, reportar las rutas no significativas y advertir que el ajuste exacto exige cautela. Esta forma de análisis fortalece la transparencia, la reproducibilidad y la defensa metodológica del modelo (Kock & Hadaya, 2018).

En este capítulo se distinguen dos usos del **Bootstrapping**. Para contrastar hipótesis, cargas externas y coeficientes de trayectoria, el **Bootstrapping** debe ejecutarse sobre el modelo normal final depurado **CAITIZEN_PLS_SIN CAIL2**. Para evaluar el ajuste global del modelo saturado, el **Bootstrapping** debe ejecutarse sobre la copia **CAITIZEN_PLS_SIN CAIL2_FIT**.

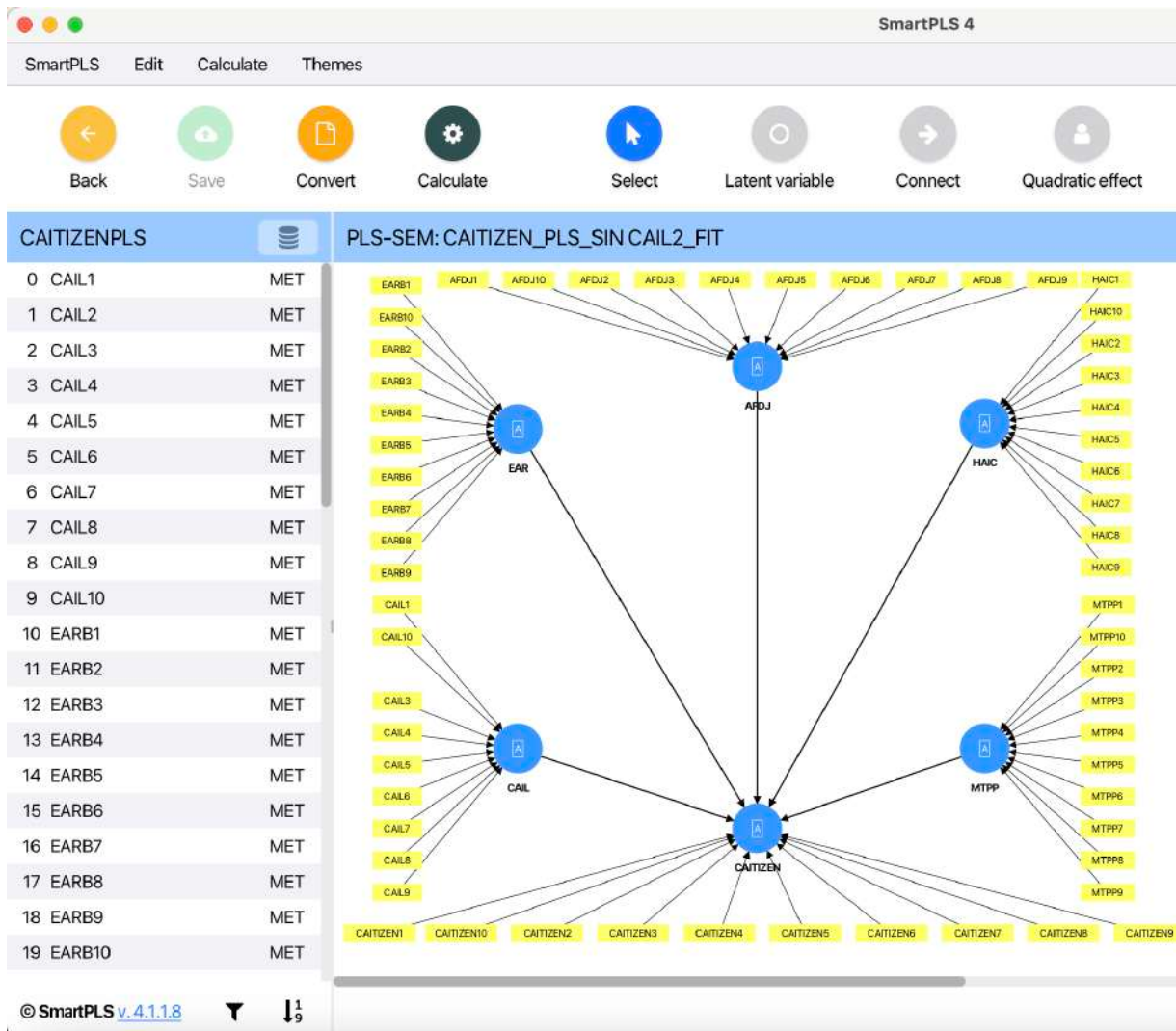
Función del **Bootstrapping** en PLS-SEM

Después de preparar el modelo **CAITIZEN_PLS_SIN CAIL2_FIT** como **composite en Modo A**, el siguiente paso consiste en ejecutar el procedimiento de **Bootstrapping**. Mientras el algoritmo **PLS-SEM** estima los valores originales del modelo con la muestra disponible, el **Bootstrapping** permite evaluar la estabilidad estadística de esos resultados mediante remuestreo. Esto significa que **SmartPLS4** genera múltiples submuestras con reemplazo para construir una distribución empírica de los parámetros estimados. A partir de esa distribución se obtienen errores estándar, **t-values**, **p-values** e intervalos de confianza, los cuales permiten valorar si las cargas, pesos, rutas y criterios de ajuste son estadísticamente estables. En este sentido, el **Bootstrapping** complementa la estimación original del modelo al generar evidencia inferencial sobre la estabilidad de sus parámetros (Hair et al., 2021). En el modelo **CAITIZEN**, el **Bootstrapping** cumple varias funciones. Primero, permite revisar la significancia de los coeficientes de trayectoria entre **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. Segundo, permite valorar la significancia de indicadores y pesos externos. Tercero, permite obtener límites **Bootstrap** para criterios de ajuste global del **modelo saturado**, como **SRMR**, **d_ULS** y **d_G**. Cuarto, facilita la revisión inferencial de diversos resultados del modelo, incluidos coeficientes de cargas, pesos, intervalos de confianza y, cuando el reporte lo permite, criterios complementarios como **HTMT**, **Cronbach Alpha**, **rho_A**, **rho_c** y **R²**, cuando se genera el reporte completo. Por ello, el **Bootstrapping** no debe verse únicamente como una técnica para contrastar hipótesis, sino como un procedimiento general de inferencia dentro de **PLS-SEM** (Hair et al., 2019).

Antes de iniciar el cálculo, debe verificarse que el modelo abierto sea la versión preparada en el capítulo anterior: **CAITIZEN_PLS_SIN CAIL2_FIT**. Esta versión conserva la exclusión de **CAIL2**, mantiene las rutas estructurales del modelo y presenta los constructos configurados en **Modo A**. La pantalla inicial del modelo permite confirmar que los constructos muestran el símbolo **A**, que los indicadores se orientan hacia los constructos y que el modelo está listo para calcular resultados inferenciales. Esta verificación es importante porque el **Bootstrapping** debe ejecutarse sobre la misma versión preparada para el ajuste global del **modelo saturado** (Henseler et al., 2015). Ver **Figura 7.1**

Juan Mejía Trejo

Figura 7.1. Modelo CAITIZEN_PLS_SIN CAIL2_FIT preparado para ejecutar *Bootstrapping*.



Fuente: Elaboración propia con **SmartPLS4**.

Diferencia entre algoritmo PLS-SEM y *Bootstrapping*

El algoritmo **PLS-SEM** y el procedimiento de **Bootstrapping** cumplen funciones distintas. El algoritmo **PLS-SEM** estima los parámetros del modelo con la muestra original y genera los valores base, como cargas externas, pesos, coeficientes de trayectoria, R^2 , f^2 , confiabilidad, **AVE**, **HTMT** y criterios de ajuste. Posteriormente, el **Bootstrapping** permite evaluar la estabilidad inferencial de esos resultados. Sin embargo, esos valores no indican por sí mismos si los parámetros son estadísticamente significativos ni si las diferencias observadas son estables frente a

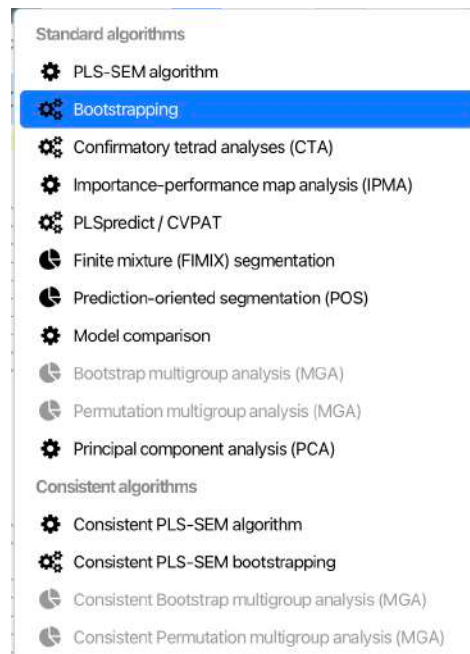
Juan Mejía Trejo

variaciones muestrales. Para ello se requiere **Bootstrapping**, ya que este procedimiento genera una distribución empírica de los resultados y permite estimar la incertidumbre asociada a cada parámetro (Hair et al., 2021).

La diferencia puede resumirse así: el algoritmo **PLS-SEM** responde “cuál es el valor estimado del parámetro”, mientras que el **Bootstrapping** responde “qué tan estable o significativo es ese parámetro”. Por ejemplo, una ruta estructural puede tener un coeficiente positivo en la muestra original, pero solo el **Bootstrapping** permite conocer su **t-value**, **p-value** e intervalo de confianza. Del mismo modo, el algoritmo puede mostrar el valor original de **SRMR**, **d_ULS** y **d_G**, pero el **Bootstrapping** permite contrastar esos valores con límites como **HI95** y **HI99**, necesarios para valorar el ajuste exacto del **modelo saturado** (Henseler et al., 2015).

En **SmartPLS4**, el procedimiento se inicia desde la opción **Calculate**, seleccionando **Bootstrapping** dentro del menú de algoritmos estándar. Esta ruta debe utilizarse una vez que el modelo se encuentra correctamente configurado. En el caso **CAITIZEN**, el análisis se ejecuta sobre la versión **CAITIZEN_PLS_SIN CAIL2_FIT**, ya preparada como **composite en Modo A**. Esta decisión asegura que los resultados **Bootstrap** correspondan al modelo configurado para valorar el ajuste global y no a una versión previa o conceptualmente distinta del modelo (Hair et al., 2021). Ver **Figura 7.2**.

Figura 7.2. Selección de *Bootstrapping* desde el menú Calculate en SmartPLS4.



Fuente: Elaboración propia con **SmartPLS4**.

Configuración de Complete *Bootstrapping* en SmartPLS4

Al seleccionar **Bootstrapping**, SmartPLS4 abre una ventana de configuración donde se definen los parámetros principales del procedimiento. En la pestaña **BT setup** se especifica el número de submuestras, el tipo de resultados a generar, el método de intervalo de confianza, el tipo de prueba, el nivel de significancia y el generador de números aleatorios. Cada uno de estos elementos afecta la forma en que se producirán los resultados inferenciales. Por esta razón, el investigador debe reportar con claridad la configuración utilizada y conservar la captura correspondiente como evidencia metodológica del procedimiento (Hair et al., 2021).

En la configuración mostrada, se utilizan **5,000 submuestras**, lo cual es una práctica recomendada para obtener estimaciones **Bootstrap** estables en análisis **PLS-SEM**. También se activa la opción **Do parallel processing**, lo que permite que el cálculo se realice con mayor rapidez si el equipo lo permite. En **Amount of results** se selecciona **Complete (slower)**, opción que genera una salida más amplia que la alternativa **Most important (faster)**. Esta decisión es adecuada porque el capítulo requiere obtener resultados completos para rutas, cargas, pesos, **HTMT**, **R²**, confiabilidad y, especialmente, criterios de ajuste del modelo como **SRMR**, **d_ULS** y **d_G** (Hair et al., 2019). La opción **Complete (slower)** es especialmente importante en este capítulo porque el objetivo no se limita a probar hipótesis estructurales. También se busca valorar el ajuste global del **modelo saturado** mediante límites bootstrap. Si se eligiera una salida reducida, podrían no estar disponibles todos los apartados necesarios para interpretar **SRMR**, **d_ULS**, **d_G**, intervalos de confianza o resultados complementarios. Por ello, para un libro metodológico aplicado, es preferible utilizar la opción completa, aunque el tiempo de procesamiento sea mayor, ya que permite documentar con mayor profundidad el análisis inferencial del modelo (Hair et al., 2021).

En la misma ventana se selecciona **Test type: Two tailed**, con un nivel de significancia de **0.05000**. La prueba de **dos colas** es una opción prudente cuando se desea evaluar si un parámetro es estadísticamente distinto de cero, sin limitar la inferencia a una sola dirección. Aunque las hipótesis del modelo **CAITIZEN** puedan formularse con una dirección esperada, la prueba de dos colas o **two-tailed test** permite evaluar si el parámetro difiere significativamente de cero en cualquiera de los dos sentidos lo que fortalece la exigencia inferencial. Esta configuración es especialmente útil en investigaciones aplicadas de ciencias sociales y administración, donde se busca mantener una interpretación conservadora (Hair et al., 2019). Ver **Figura 7.3**.

Figura 7.3. Configuración de **Complete Bootstrapping** con 5,000 submuestras y prueba de dos colas.

The screenshot shows the 'Bootstrapping' window in SmartPLS4. It features a title bar with standard window controls. Below the title bar, a descriptive text explains bootstrapping as a nonparametric procedure for testing statistical significance of PLS-SEM results. A 'More info' button is located to the right of this text. The main area is divided into three tabs: 'BT setup' (selected), 'PLS setup', and 'Data'. Under the 'BT setup' tab, several configuration options are listed on the left, each with a corresponding input field or dropdown on the right: 'Subsamples' is set to 5000; 'Do parallel processing' is checked; 'Amount of results' has two options, 'Complete (slower)' and 'Most important (faster)', with 'Complete (slower)' selected; 'Confidence interval method' is set to 'Complete (slower)'; 'Test type' is set to 'Two tailed'; 'Significance level' is set to 0.05000; and 'Random number generator' is set to 'Fixed seed'. At the bottom right, there is an 'Open report' checkbox which is unchecked. At the bottom center, there is a 'Close' button. At the bottom right, there is a blue 'Start calculation' button. A 'Default settings' link is located at the bottom left.

Fuente: Elaboración propia con **SmartPLS4**.

Ejecución del cálculo *Bootstrap* y control del proceso

Una vez configurados los parámetros, se presiona **Start calculation** para iniciar el procedimiento. Durante el cálculo, **SmartPLS4** muestra la pestaña **Progress**, donde se observa el avance del procesamiento. En la captura se identifican etapas como lectura de la matriz de puntuaciones, cálculo de la muestra original, ejecución de las **5,000** submuestras, cálculo del ajuste del **modelo saturado**, cálculo del ajuste del modelo estimado, pesos internos, efectos indirectos, efectos totales, pesos externos y cargas externas. Esta pantalla es importante porque confirma que el procedimiento no solo está estimando rutas, sino también criterios asociados al ajuste global del modelo (Hair et al., 2021).

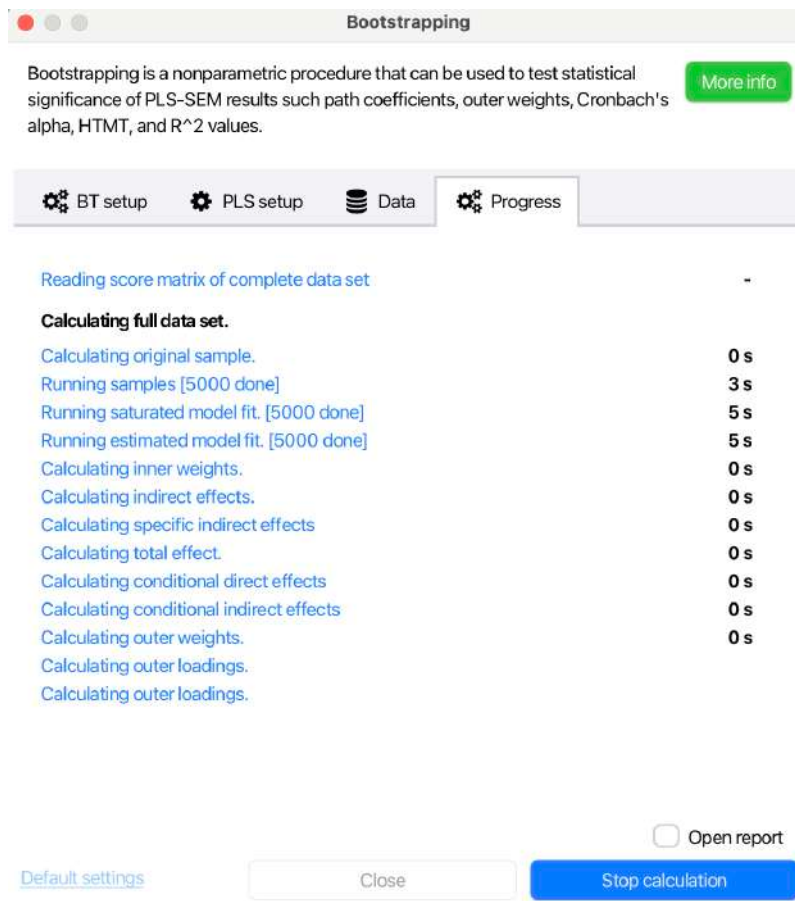
La línea **Running saturated model fit [5000 done]** confirma que se calcularon los criterios de ajuste del **modelo saturado** en las submuestras bootstrap. Esta

Juan Mejía Tréjo

información es necesaria porque los criterios **SRMR**, **d_ULS** y **d_G** deben compararse contra límites **Bootstrap** como **HI95** y **HI99**. Sin esta ejecución completa, el investigador solo tendría valores originales, pero no los límites necesarios para evaluar el ajuste exacto del modelo. Por tanto, esta etapa documenta que la valoración del **fit** se generó bajo un procedimiento inferencial y no solo descriptivo (Henseler et al., 2015).

También se observa la línea **Running estimated model fit [5000 done]**, lo que indica que **SmartPLS4** calcula resultados tanto para el **modelo saturado** como para el **modelo estimado**. Esta distinción es metodológicamente importante. El **modelo saturado** permite revisar la estructura global de asociaciones entre constructos, mientras que el **modelo estimado** corresponde a las relaciones teóricas especificadas por el investigador. En el caso **CAITIZEN**, ambos resultados pueden aparecer en el reporte, pero la valoración del ajuste global del modelo de medida debe concentrarse principalmente en el **modelo saturado**, sin confundirlo con la interpretación del **modelo estimado** (Henseler et al., 2015). Ver Figura 7.4.

Figura 7.4. Progreso del Complete *Bootstrapping* y cálculo del saturated model



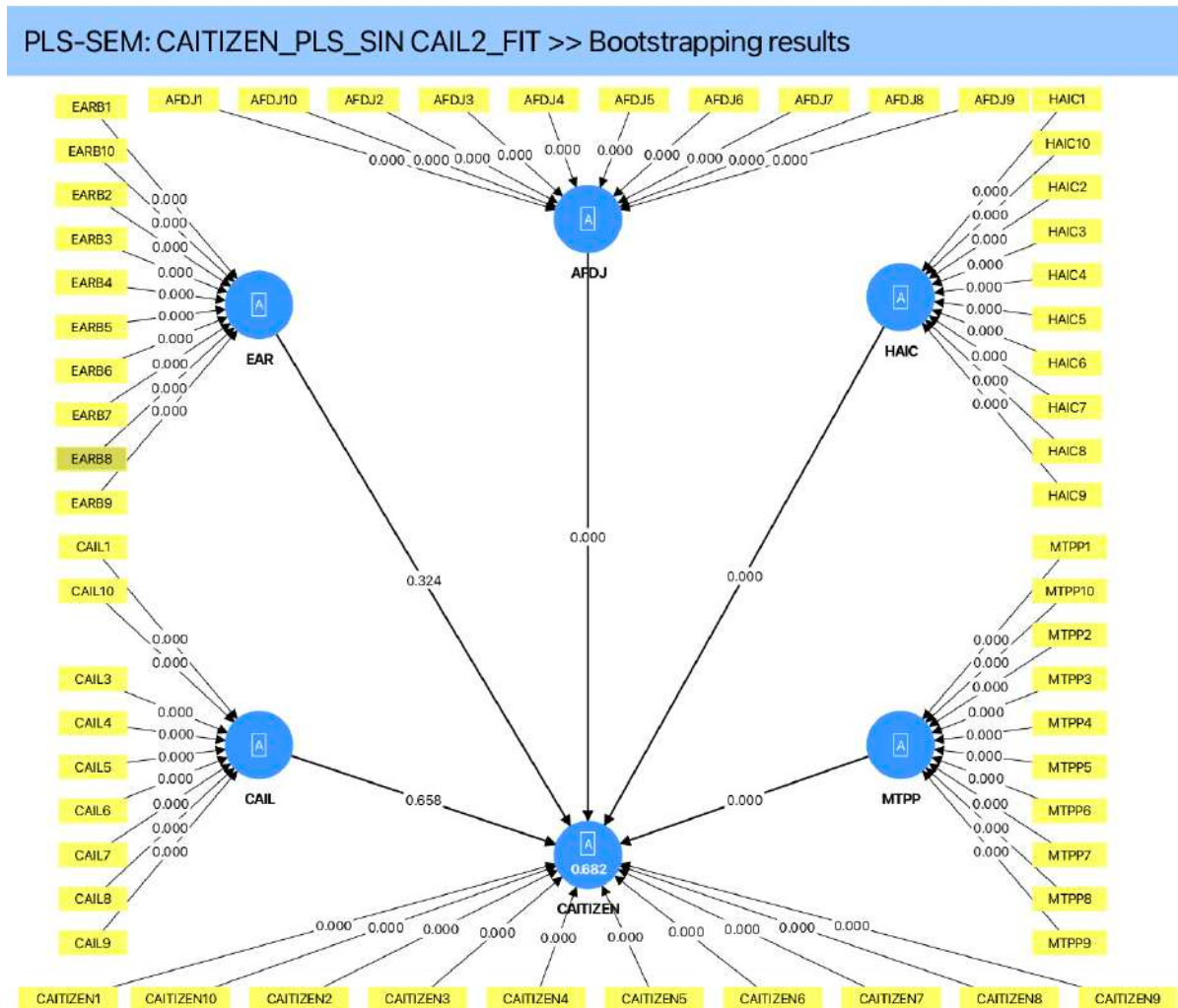
Fuente: Elaboración propia con **SmartPLS4**.

Salida gráfica del **Bootstrapping**

Después de concluir el cálculo, **SmartPLS4** muestra los resultados del **Bootstrapping** en una salida gráfica. En esta pantalla pueden observarse valores asociados a rutas estructurales e indicadores, dependiendo del tipo de resultado seleccionado. En el caso mostrado, aparecen **p-values** asociados a indicadores y rutas, lo que permite una lectura visual preliminar de la significancia estadística. Esta visualización permite una primera lectura preliminar de la significancia estadística, aunque la interpretación formal debe realizarse en las tablas del reporte, donde se encuentran **Original Sample**, **Sample Mean**, **STDEV**, **t-value** y **p-value** (Hair et al., 2021).

La salida gráfica permite identificar rápidamente cuáles relaciones parecen significativas y cuáles no. Por ejemplo, en la figura se observan valores **p** iguales a **0.000** en varias cargas e indicadores, lo que sugiere alta significancia estadística. Sin embargo, también puede observarse una ruta con valor **p = 0.324** y otra con **p = 0.658**, lo que anticipa que no todas las relaciones estructurales del modelo son estadísticamente significativas. Esta primera lectura no debe sustituir la tabla de resultados, pero ayuda al lector a ubicar visualmente dónde aparecen los resultados inferenciales en el modelo (Hair et al., 2019).

Debe explicarse al lector que un **p-value** significativo no es suficiente, por sí solo, para conservar un indicador o aceptar la relevancia sustantiva de una ruta. En el caso de las cargas externas, se debe revisar simultáneamente la magnitud de la carga, la confiabilidad, la **AVE** y la coherencia conceptual. Del mismo modo, en el caso de las rutas estructurales, se debe revisar el coeficiente original, el **t-value**, el **p-value**, el intervalo de confianza y la relevancia sustantiva de la relación. La salida gráfica es útil para orientación inicial, pero la decisión metodológica debe basarse en el reporte completo (Hair et al., 2021). Ver **Figura 7.5**.

Figura 7.5. Salida gráfica del *Bootstrapping* del modelo CAITIZEN

Fuente: Elaboración propia con **SmartPLS4**..

Recolección de *t-values* y *p-values* para cargas externas

Una vez ejecutado el **Complete Bootstrapping**, los *t-values* y *p-values* de los indicadores deben recolectarse desde el reporte de **Outer loadings**. Estos valores permiten complementar la tabla del modelo de medición, ya que la carga externa muestra la magnitud de la relación entre el indicador y su constructo, mientras que el *t-value* y el *p-value* permiten valorar su significancia estadística. Para reportar cada indicador, se debe conservar la carga externa como **Original Sample (O)** y añadir, entre paréntesis, el *t-value* obtenido mediante **Bootstrapping**; posteriormente, en una columna separada, se registra el *p-value* correspondiente. Por ejemplo, la tabla puede reportarse como **0.761 (t = X.XXX)** y, en la columna siguiente, **p = X.XXX**. Esta forma

Juan Mejía Trejo

de presentación permite distinguir entre la magnitud de la carga y su significancia inferencial, evitando asumir que una carga es adecuada solo porque resulta estadísticamente significativa. La decisión final sobre conservar o revisar un indicador debe considerar conjuntamente la carga externa, la confiabilidad interna, la **AVE**, la validez discriminante y la coherencia conceptual del constructo (Hair et al., 2021). Ver **Figura 7.6**.

Figura 7.6. Recolección de t-values y p-values de las cargas externas del modelo CAITIZEN

SmartPLS

Export

Themes

Edit

Save

Excel

HTML

Create data file

Compare

Bootstrapping

Outer loadings - Mean, STDEV, T values, p values

▼ Graphical

Graphical output

▼ Final results

▶ Path coefficients

Intercepts

Total indirect effects

Specific indirect effects

▶ Total effects

▼ Outer loadings

☐ Mean, STDEV, T values, p values

☐ Confidence intervals

☐ Confidence intervals bias corrected

▶ Outer weights

Conditional direct effects

Conditional indirect effects

▼ Quality criteria

▶ R-square

▶ R-square adjusted

▶ f-square

Average variance extracted (AVE)

Composite reliability (rho_c)

▶ Composite reliability (rho_a)

Cronbach's alpha

	Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O/STDEV)	P values
AFDJ1 -> AFDJ	0.678	0.677	0.031	22.066	0.000
AFDJ10 -> AFDJ	0.781	0.780	0.022	36.238	0.000
AFDJ2 -> AFDJ	0.747	0.746	0.027	27.486	0.000
AFDJ3 -> AFDJ	0.771	0.770	0.023	33.435	0.000
AFDJ4 -> AFDJ	0.720	0.720	0.027	26.639	0.000
AFDJ5 -> AFDJ	0.745	0.745	0.024	30.743	0.000
AFDJ6 -> AFDJ	0.785	0.784	0.024	32.474	0.000
AFDJ7 -> AFDJ	0.812	0.812	0.019	42.974	0.000
AFDJ8 -> AFDJ	0.768	0.766	0.028	27.375	0.000
AFDJ9 -> AFDJ	0.803	0.802	0.020	40.303	0.000
CAIL1 -> CAIL	0.662	0.661	0.034	19.219	0.000
CAIL10 -> CAIL	0.706	0.706	0.029	24.044	0.000
CAIL3 -> CAIL	0.593	0.592	0.038	15.583	0.000
CAIL4 -> CAIL	0.761	0.760	0.023	32.789	0.000
CAIL5 -> CAIL	0.754	0.753	0.025	29.681	0.000
CAIL6 -> CAIL	0.711	0.710	0.032	22.521	0.000
CAIL7 -> CAIL	0.742	0.740	0.031	24.226	0.000
CAIL8 -> CAIL	0.802	0.801	0.020	40.435	0.000
CAIL9 -> CAIL	0.745	0.744	0.024	30.415	0.000
CAITZEN1 -> CAITZEN	0.712	0.712	0.028	25.121	0.000
CAITZEN10 -> CAITZEN	0.795	0.795	0.020	38.911	0.000
CAITZEN2 -> CAITZEN	0.740	0.740	0.029	25.553	0.000
CAITZEN3 -> CAITZEN	0.719	0.719	0.030	24.357	0.000
CAITZEN4 -> CAITZEN	0.742	0.742	0.027	27.366	0.000
CAITZEN5 -> CAITZEN	0.731	0.731	0.025	29.067	0.000
CAITZEN6 -> CAITZEN	0.658	0.658	0.036	18.382	0.000
CAITZEN7 -> CAITZEN	0.779	0.779	0.023	34.150	0.000
CAITZEN8 -> CAITZEN	0.750	0.749	0.031	24.017	0.000
CAITZEN9 -> CAITZEN	0.795	0.795	0.021	37.658	0.000

Fuente: Elaboración propia con **SmartPLS4**.

Los valores **t-value** y **p-value** reportados en la tabla proceden del procedimiento de **Bootstrapping** y no del algoritmo **PLS-SEM**. El algoritmo proporciona las cargas originales del modelo, mientras que el **Bootstrapping** permite estimar la significancia estadística de dichas cargas mediante remuestreo. Por ello, la tabla debe distinguir entre la magnitud de la carga externa, expresada como **Original Sample (O)**, y la evidencia inferencial, expresada mediante **t-value** y **p-value**. Bajo una prueba de **dos colas** con $\alpha = 0.05$, un indicador se considera estadísticamente significativo cuando $t \geq 1.96$ y $p < 0.05$. Sin embargo, la significancia estadística no sustituye la evaluación de la magnitud de la carga ni la coherencia teórica del indicador.

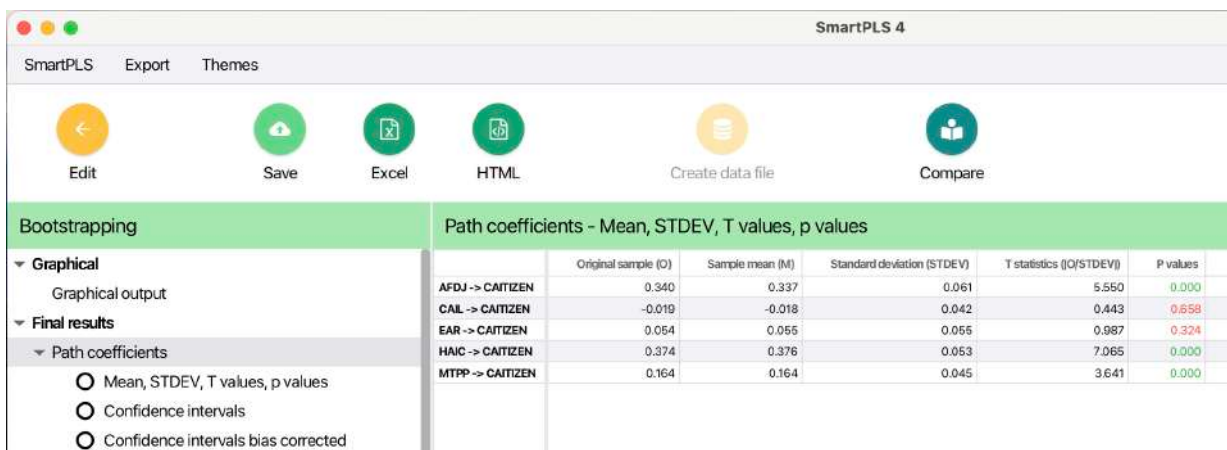
Interpretación de coeficientes de trayectoria

La tabla **Path coefficients – Mean, STDEV, T values, p values** permite interpretar la significancia de las rutas estructurales del modelo. En la captura se muestran cinco relaciones: **AFDJ** → **CAITIZEN**, **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTTP** → **CAITIZEN**. Para cada relación, **SmartPLS4** presenta el valor original de la muestra, la media **Bootstrap**, la desviación estándar, el estadístico **t** y el valor **p**. Estos elementos permiten valorar si cada predictor tiene un efecto estadísticamente significativo sobre **CAITIZEN** (Hair et al., 2021).

En los resultados mostrados, **AFDJ** → **CAITIZEN** presenta un coeficiente original de **0.340**, **t** = **5.550** y **p** = **0.000**, lo que indica una relación positiva y estadísticamente significativa. **HAIC** → **CAITIZEN** muestra un coeficiente de **0.374**, **t** = **7.065** y **p** = **0.000**, siendo la relación con mayor coeficiente directo entre las rutas observadas. **MTTP** → **CAITIZEN** presenta un coeficiente de **0.164**, **t** = **3.641** y **p** = **0.000**, lo que también indica significancia estadística. Estos resultados sugieren que la conciencia de justicia de datos, la colaboración humano-IA y la transparencia metacognitiva en el uso de prompts contribuyen significativamente a **CAITIZEN** (Hair et al., 2019).

En contraste, **CAIL** → **CAITIZEN** presenta un coeficiente original de **-0.019**, **t** = **0.443** y **p** = **0.658**, por lo que no muestra significancia estadística. De forma similar, **EAR** → **CAITIZEN** presenta un coeficiente de **0.054**, **t** = **0.987** y **p** = **0.324**, también sin significancia estadística bajo una prueba de dos colas con $\alpha = 0.05$. Estos resultados no implican que **CAIL** o **EAR** carezcan de relevancia conceptual; indican que, en este modelo y controlando por los demás predictores, no presentan efecto directo significativo sobre **CAITIZEN** (Hair et al., 2021). Ver **Figura 7.7**.

Figura 7.7. Coeficientes de trayectoria del Bootstrapping: media, desviación estándar, t-values y p-values.



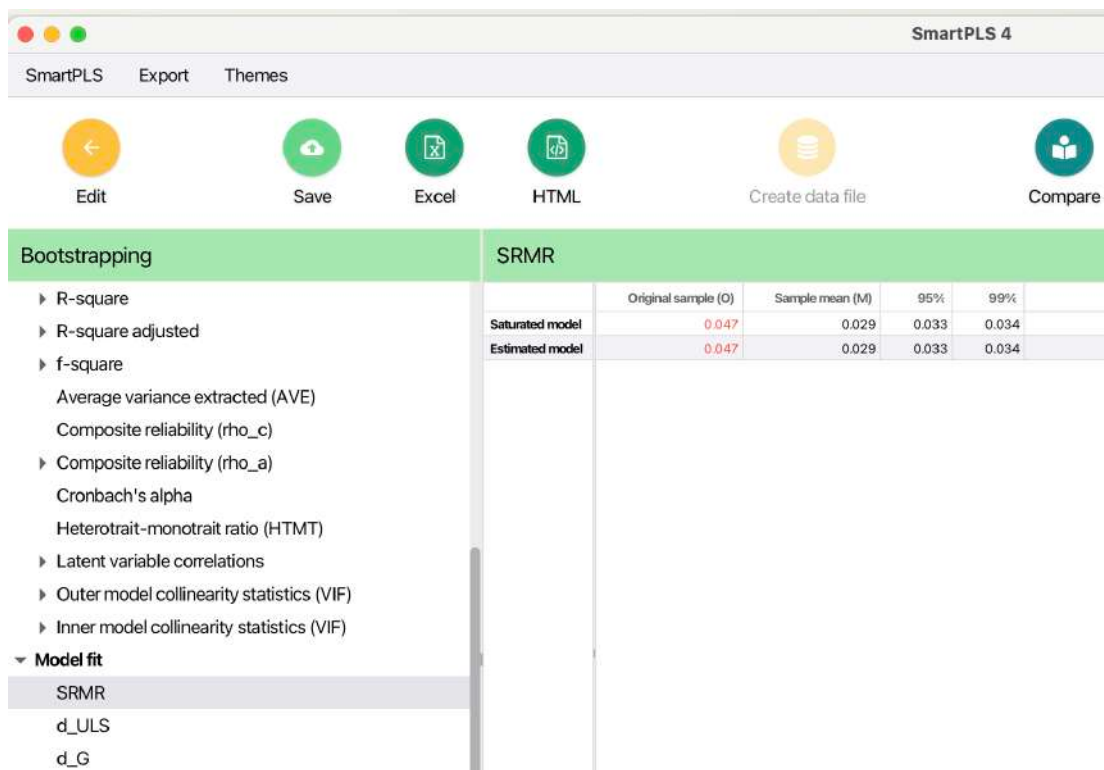
Bootstrapping		Path coefficients - Mean, STDEV, T values, p values				
		Original sample (O)	Sample mean (M)	Standard deviation (STDEV)	T statistics (O /STDEV)	P values
AFDJ → CAITIZEN		0.340	0.337	0.061	5.550	0.000
CAIL → CAITIZEN		-0.019	-0.018	0.042	0.443	0.658
EAR → CAITIZEN		0.054	0.055	0.055	0.987	0.324
HAIC → CAITIZEN		0.374	0.376	0.053	7.065	0.000
MTTP → CAITIZEN		0.164	0.164	0.045	3.641	0.000

Fuente: Elaboración propia con **SmartPLS4**.

Índice de ajuste aproximado: SRMR

El **SRMR** o **Standardized Root Mean Square Residual** es un índice de ajuste aproximado que expresa la discrepancia promedio estandarizada entre la matriz observada y la matriz reproducida por el modelo. En términos prácticos, cuanto menor sea el **SRMR**, menor será la discrepancia global entre los datos y el modelo. En la tabla de resultados **Bootstrap** se observa que el **SRMR** del **modelo saturado** presenta un valor original de **0.047**, con media **Bootstrap** de **0.029**, límite **95% = 0.033** y límite **99% = 0.034**. Desde los umbrales convencionales de ajuste aproximado, el valor original **0.047** se encuentra por debajo de **0.08** y **0.10**, lo que respalda una lectura favorable del ajuste aproximado del modelo (Hu & Bentler, 1999). Ver **Figura 7.8**.

Figura 7.8. Resultado *Bootstrap* del índice SRMR para el modelo saturado y estimado.



Fuente: Elaboración propia con **SmartPLS4**.

Aunque el **SRMR = 0.047** se encuentra dentro de los umbrales convencionales aceptables, debe explicarse con cautela la comparación con los límites bootstrap. En algunos reportes de **SmartPLS4**, el valor original puede aparecer en color rojo cuando supera el límite **Bootstrap** mostrado para **95%** o **99%**, aun cuando sea menor a los umbrales convencionales de **0.08** o **0.10**. Por ello, la interpretación debe distinguir entre ajuste aproximado por umbral convencional y ajuste exacto basado en bootstrap. En este caso, el **SRMR** permite sostener una evidencia favorable de ajuste

Juan Mejía Trejo

aproximado, pero no debe interpretarse de forma aislada ni como prueba definitiva de ajuste global perfecto (Henseler et al., 2015).

La tabla de criterios utilizada en el capítulo debe presentar de forma clara dos niveles de lectura. El primero corresponde a los índices de ajuste aproximado: **SRMR < 0.08** como criterio estricto y **SRMR < 0.10** como criterio aceptable. El segundo corresponde al test de ajustes exactos basados en **Bootstrap**, donde los valores originales de **SRMR**, **d_ULS** y **d_G** se comparan con **HI95** y **HI99**. Esta tabla ayuda al lector a entender que no todos los criterios se interpretan de la misma forma, aunque todos contribuyen a valorar el ajuste global del **modelo saturado** (Williams et al., 2009).

Después de revisar el valor **SRMR**, conviene sintetizar los criterios de interpretación del ajuste global del **modelo saturado**. Para ello, se utiliza una tabla de referencia que distingue entre dos tipos de valoración: los **índices de ajuste aproximado**, representados principalmente por **SRMR**, y las **pruebas de ajuste exacto basadas en Bootstrap**, representadas por **SRMR**, **d_ULS** y **d_G** en comparación con los límites **HI95** y **HI99**. Esta distinción es importante porque un modelo puede cumplir el criterio convencional de **SRMR < 0.08** o **SRMR < 0.10**, pero no cumplir plenamente los criterios exactos basados en límites bootstrap. Por ello, la tabla debe utilizarse como guía para interpretar los resultados de ajuste de manera diferenciada y prudente. Ver **Tabla 7.1**

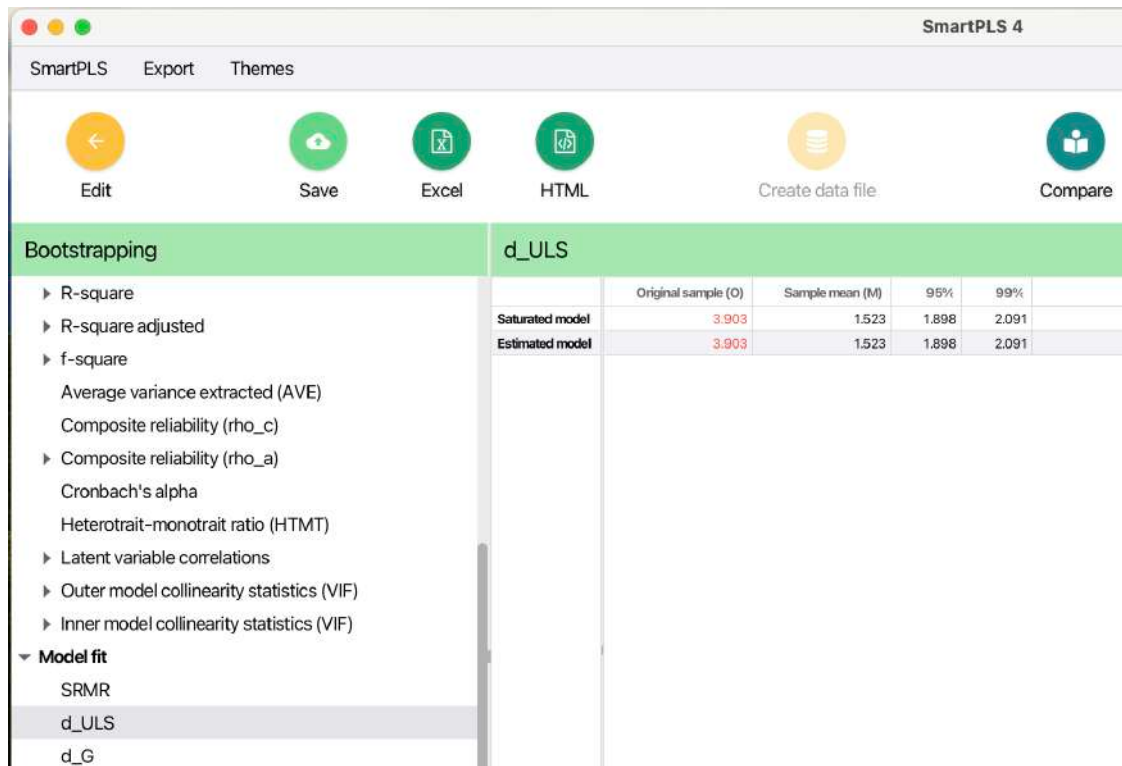
Tabla 7.1. Criterios para la valoración del modelo de medida mediante el modelo saturado.

Criterio	Umbral
Índices de ajuste (valoración aproximada)	• SRMR < 0.08 (Hu & Bentler, 1999). • SRMR < 0.10 (Williams et al., 2009).
Test de ajustes exactos basados en bootstrap	• SRMR ≤ HI95 ≤ HI99 • d_{ULS} ≤ HI95 ≤ HI99 • d_G ≤ HI95 ≤ HI99

Fuente: Elaboración propia con base en Hu y Bentler (1999); Williams et al. (2009)

Ajuste exacto mediante d_ULS

La medida **d_ULS** o **Squared Euclidean Distance**, corresponde a una discrepancia basada en la distancia euclidiana al cuadrado entre la matriz observada y la matriz reproducida por el modelo. En la captura se observa que **d_ULS** presenta un valor original de **3.903**, con media **Bootstrap** de **1.523**, límite **95% = 1.898** y límite **99% = 2.091**, tanto para el modelo saturado como para el modelo estimado. Bajo la regla de ajuste exacto, el valor original debería ser menor o igual que los límites bootstrap; sin embargo, en este caso el valor original supera tanto **HI95** como **HI99** (Henseler et al., 2015). Ver **Figura 7.9**.

Figura 7.9. Resultado *Bootstrap* de d_{ULS} para el modelo saturado y estimado.


Bootstrapping		d_ULS				
		Original sample (O)	Sample mean (M)	95%	99%	
▶ R-square						
▶ R-square adjusted						
▶ f-square						
Average variance extracted (AVE)						
Composite reliability (rho_c)						
▶ Composite reliability (rho_a)						
Cronbach's alpha						
Heterotrait-monotrait ratio (HTMT)						
▶ Latent variable correlations						
▶ Outer model collinearity statistics (VIF)						
▶ Inner model collinearity statistics (VIF)						
▼ Model fit						
SRMR						
d_ULS						
d_G						

Fuente: Elaboración propia con SmartPLS4.

Cuando d_{ULS} supera los límites **HI95** y **HI99**, la interpretación debe ser prudente. Esto no significa automáticamente que el modelo no sirva o que deba descartarse, pero sí indica que el ajuste exacto basado en esta discrepancia no es plenamente satisfactorio. En términos metodológicos, puede afirmarse que el modelo presenta evidencia favorable de ajuste aproximado mediante **SRMR**, pero no cumple completamente el criterio de ajuste exacto según d_{ULS} . Esta distinción es importante para evitar conclusiones excesivas y para reportar con transparencia los resultados del modelo **CAITIZEN** (Hair et al., 2021).

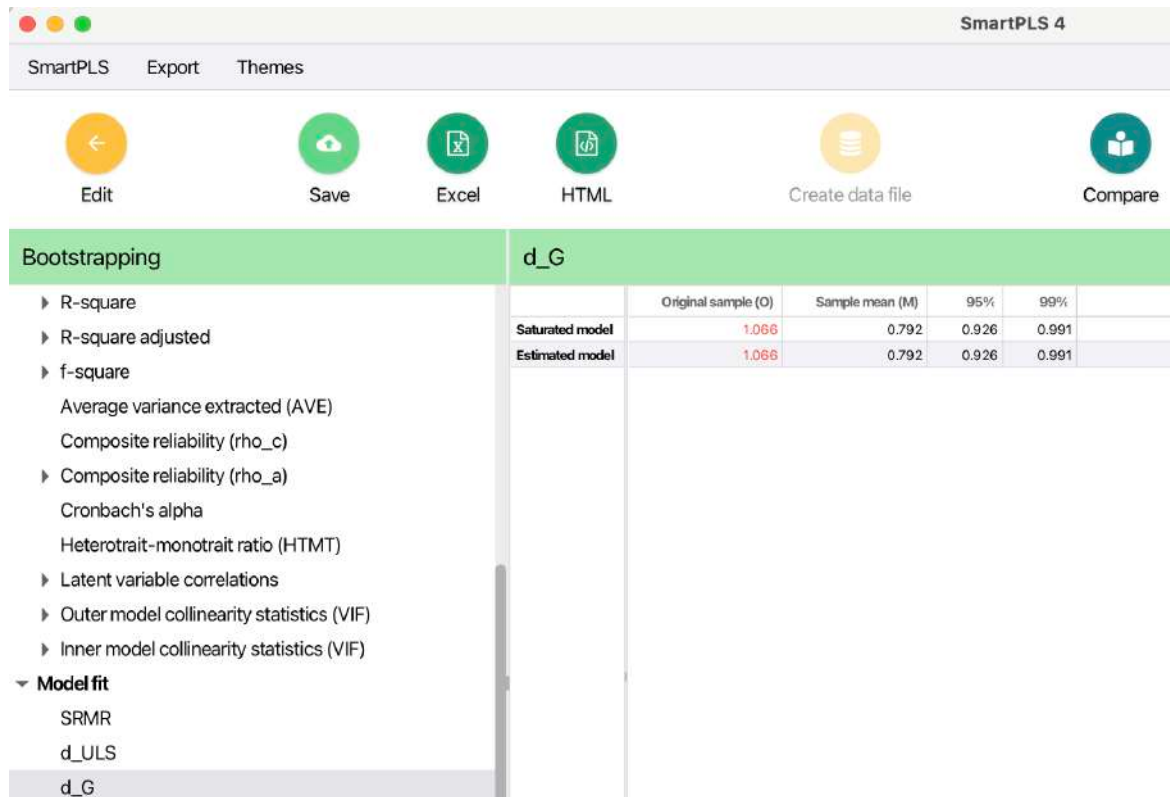
Ajuste exacto mediante d_G

La medida d_G o **Geodesic Distance** es una medida de discrepancia exacta que estima la separación entre la matriz empírica observada y la matriz implicada por el modelo mediante una distancia geodésica. Al igual que d_{ULS} , se interpreta comparando el valor original con los límites **HI95** y **HI99**. Si el valor original de d_G es menor o igual que los límites **Bootstrap**, existe evidencia favorable de ajuste exacto. Si el valor original es mayor que esos límites, el modelo no cumple plenamente ese criterio y debe reportarse con cautela. Se interpreta comparando el valor original con los límites obtenidos mediante **Bootstrapping**. En la captura se observa que d_G

Juan Mejía Trejo

presenta un valor original de **1.066**, con media **Bootstrap** de **0.792**, límite **95% = 0.926** y límite **99% = 0.991**. En este caso, el valor original también supera los límites **Bootstrap**, por lo que el ajuste exacto según **d_G** tampoco se cumple plenamente (Henseler et al., 2015). Ver **Figura 7.10**.

Figura 7.10. Resultado *Bootstrap* de d_G para el modelo saturado y estimado.



d_G		Original sample (O)	Sample mean (M)	95%	99%
Saturated model		1.066	0.792	0.926	0.991
Estimated model		1.066	0.792	0.926	0.991

Fuente: Elaboración propia con **SmartPLS4**.

La lectura conjunta de **SRMR**, **d_ULS** y **d_G** permite formular una decisión metodológica equilibrada. El **SRMR** muestra un valor favorable bajo criterios convencionales de ajuste aproximado; sin embargo, **d_ULS** y **d_G** superan los límites **Bootstrap** de **95%** y **99%**, lo que sugiere que el ajuste exacto del modelo saturado debe reportarse con cautela. Por tanto, no sería adecuado afirmar que el modelo presenta ajuste global perfecto. Una redacción más precisa sería señalar que el modelo presenta ajuste aproximado aceptable, pero evidencia parcial o limitada en las pruebas de ajuste exacto basadas en **Bootstrap** (Hair et al., 2021).

Interpretación de resultados parciales en los índices de ajuste global

La interpretación de los índices **SRMR**, **d_ULS** y **d_G** debe realizarse con prudencia, porque no todos evalúan el ajuste global de la misma manera. El **SRMR** se utiliza principalmente como índice de **ajuste aproximado**, mientras que **d_ULS** y **d_G** funcionan como medidas de **ajuste exacto** basadas en comparación con límites obtenidos mediante **Bootstrapping**, como **HI95** y **HI99**. Por ello, que uno o más criterios no se cumplan no significa automáticamente que el modelo deba descartarse. Significa que el investigador debe distinguir qué tipo de ajuste está fallando, revisar posibles causas metodológicas y decidir si el modelo conserva valor teórico, explicativo y empírico dentro del análisis **PLS-SEM** (Henseler et al., 2015).

Cuando los tres criterios presentan evidencia favorable, es decir, cuando el **SRMR** se encuentra por debajo de los umbrales convencionales y **d_ULS** y **d_G** no superan sus límites **HI95** y **HI99**, puede afirmarse que el modelo presenta un ajuste global sólido. En este escenario, el modelo cuenta con evidencia favorable tanto de ajuste aproximado como de ajuste exacto. Sin embargo, incluso en este caso, no debe afirmarse que el modelo queda “validado totalmente”, porque el ajuste global no sustituye la evaluación de cargas externas, confiabilidad interna, **AVE**, validez discriminante, colinealidad, coeficientes de trayectoria, **R²**, **f²** y significancia estadística. La conclusión correcta sería que el ajuste global respalda la continuidad del análisis estructural (Hair et al., 2021).

Cuando falla únicamente uno de los tres criterios, la conclusión debe ser moderada. Por ejemplo, si **SRMR** cumple, pero **d_ULS** o **d_G** no cumple, puede decirse que el modelo presenta evidencia favorable de ajuste aproximado, aunque una de las pruebas exactas sugiere cautela.

Si, por el contrario, **d_ULS** y **d_G** cumplen, pero **SRMR** supera los umbrales convencionales, el problema se concentra en el ajuste aproximado y debe revisarse con especial atención la discrepancia promedio entre la matriz observada y la matriz reproducida por el modelo. En cualquiera de los dos casos, el modelo no se descarta automáticamente; se revisan posibles fuentes de discrepancia y se reporta una conclusión parcial, no absoluta (Hair et al., 2019).

Cuando fallan dos de los tres criterios, la cautela debe ser mayor. Esta situación indica que el ajuste global no es plenamente satisfactorio y que existe evidencia de discrepancia relevante entre el modelo y los datos. Sin embargo, la decisión todavía no debe ser mecánica. Si el **SRMR** cumple, pero **d_ULS** y **d_G** no cumplen, el modelo conserva evidencia de ajuste aproximado, pero no de ajuste exacto. Esta situación es frecuente en modelos aplicados con constructos cercanos, muchos indicadores o alta complejidad. La redacción adecuada sería señalar que el modelo presenta ajuste aproximado aceptable, pero evidencia limitada en las pruebas exactas basadas en **Bootstrapping** (Henseler et al., 2015).

Juan Mejía Trejo

Si fallan **SRMR** y uno de los dos índices exactos, pero el otro índice exacto cumple, la situación también debe interpretarse con reserva. En ese caso, el modelo presenta problemas en el ajuste aproximado y una señal parcial de discrepancia exacta. Esto sugiere revisar con mayor detenimiento el modelo de medición, las cargas externas, la validez discriminante, la cercanía conceptual entre constructos, la posible redundancia de indicadores o la especificación de las rutas estructurales. No obstante, si el modelo tiene una justificación teórica fuerte y los resultados estructurales son coherentes, puede conservarse como modelo exploratorio o aplicado, siempre que se reporte explícitamente la limitación en el ajuste global (Hair et al., 2021).

Cuando fallan los tres criterios, es decir, cuando **SRMR** supera los umbrales convencionales y además **d_ULS** y **d_G** superan los límites **Bootstrap HI95** y **HI99**, el modelo presenta evidencia clara de falta de ajuste global. En este escenario, no es recomendable reportar el modelo como ajustado. Tampoco conviene avanzar hacia conclusiones estructurales fuertes sin revisar antes la especificación del modelo. Deben analizarse posibles problemas como indicadores mal asignados, constructos excesivamente similares, cargas externas bajas, problemas de validez discriminante, relaciones estructurales omitidas, redundancia entre constructos o errores en la configuración del modelo de medición. Aquí sí se recomienda reexaminar el modelo antes de sostener conclusiones teóricas fuertes (Hair et al., 2019).

En modelos **PLS-SEM** aplicados a ciencias sociales y administración, la falta de ajuste perfecto puede explicarse por varias razones. Una causa frecuente es la cercanía conceptual entre constructos, especialmente cuando las dimensiones del modelo pertenecen a un mismo campo teórico. Otra causa puede ser la presencia de indicadores con cargas moderadas, los cuales pueden ser conceptualmente útiles, pero generar discrepancia estadística. También influyen la complejidad del modelo, el número elevado de indicadores, la existencia de relaciones no especificadas o la sensibilidad de las pruebas exactas basadas en **Bootstrapping**. Por ello, la interpretación de **SRMR**, **d_ULS** y **d_G** debe integrarse con el resto de la evidencia metodológica (Hair et al., 2021).

En el caso del modelo **CAITIZEN**, los resultados muestran una situación mixta en la valoración del ajuste global del modelo saturado. El **SRMR = 0.047** se ubica por debajo de los umbrales convencionales de **0.08** y **0.10**, por lo que sugiere ajuste aproximado favorable de **ajuste aproximado**; sin embargo, supera los límites **Bootstrap HI95 = 0.033** y **HI99 = 0.034**, razón por la cual **SmartPLS4** lo marca en rojo. Esto indica que, aunque el modelo presenta una lectura aceptable bajo criterios aproximados, **no cumple el criterio Bootstrap** de ajuste exacto para **SRMR**, aunque sí cumple los umbrales convencionales de ajuste aproximado. La cautela aumenta al observar **d_ULS = 3.903** y **d_G = 1.066**, ya que ambos valores superan sus respectivos límites **HI95** y **HI99**. El valor **d_ULS = 3.903** supera sus límites **Bootstrap**, y el valor **d_G = 1.066** también se encuentra por encima de sus respectivos límites **HI95** y **HI99**. Por tanto, el modelo no satisface plenamente las pruebas exactas de ajuste global. No obstante, esta situación no implica rechazar automáticamente el modelo ni concluir que carece de utilidad teórica o empírica. La interpretación metodológicamente adecuada

es señalar que el modelo **CAITIZEN** presenta evidencia favorable de ajuste aproximado mediante **SRMR**, pero evidencia limitada en los criterios de ajuste exacto basados en **SRMR Bootstrap**, **d_ULS** y **d_G**.

En consecuencia, el modelo debe reportarse con **cautela metodológica**. No conviene afirmar que el ajuste global es pleno o perfecto; tampoco debe descartarse de manera automática. La conclusión más equilibrada es que el modelo conserva valor analítico para continuar hacia la evaluación estructural, siempre que los resultados posteriores de colinealidad, coeficientes de trayectoria, **R²**, **f²**, significancia estadística y relevancia sustantiva sean interpretados de manera integrada. Esta formulación evita tanto la sobrevalidación del modelo como su descarte injustificado, y mantiene una lectura coherente con la evaluación progresiva recomendada en **PLS-SEM** (Henseler et al., 2015).

La falta de ajuste exacto en **d_ULS** y **d_G** puede relacionarse con características ya observadas en el modelo **CAITIZEN**. En capítulos previos se identificaron relaciones conceptualmente cercanas entre algunos constructos, como **AFDJ-EAR**, **CAIL-AFDJ** y **CAITIZEN-HAIC**. Esta cercanía puede generar mayor dificultad para que el modelo reproduzca perfectamente la matriz empírica bajo criterios exactos. Además, algunos indicadores con cargas moderadas pueden contribuir a incrementar la discrepancia global. Por tanto, los resultados de **d_ULS** y **d_G** deben interpretarse como una señal de cautela metodológica, no como una invalidación automática del modelo (Hair et al., 2019).

En consecuencia, el modelo puede continuar hacia la evaluación estructural, siempre que se mantenga una redacción prudente. El investigador debe revisar colinealidad, coeficientes de trayectoria, **R²**, **f²**, significancia estadística, relevancia sustantiva y coherencia teórica de las relaciones. Si estos resultados son consistentes, el modelo puede conservar valor explicativo y utilidad académica, aun cuando las pruebas exactas de ajuste global no sean plenamente satisfactorias. La interpretación adecuada para **CAITIZEN** es que el modelo muestra ajuste aproximado favorable mediante **SRMR**, pero requiere cautela al afirmar ajuste global exacto debido a los resultados de **d_ULS** y **d_G** (Hair et al., 2021).

La decisión final debe formularse en términos de grados de evidencia, no en términos absolutos de aceptación o rechazo. **Un modelo con tres índices favorables puede reportarse como sólido en ajuste global; un modelo con dos índices favorables puede reportarse como aceptable con observaciones; un modelo con un solo índice favorable debe reportarse como parcialmente respaldado y sujeto a cautela; y un modelo con los tres índices desfavorables requiere revisión sustantiva antes de avanzar.** Esta lógica permite que la evaluación del ajuste global sea más realista, especialmente en investigaciones aplicadas donde los modelos suelen ser teóricamente complejos y empíricamente sensibles (Hair et al., 2019).

Por tanto, la pregunta no debe formularse como “¿el modelo sirve o no sirve?”, sino como “¿qué tipo de evidencia de ajuste ofrece el modelo y qué cautelas deben reportarse?”. En el caso **CAITIZEN**, el modelo no debe descartarse automáticamente

Juan Mejía Trejo

porque cuenta con un **SRMR** favorable y con resultados estructurales que pueden ser interpretados posteriormente. Sin embargo, tampoco debe presentarse como un modelo con ajuste global perfecto, ya que **d_ULS** y **d_G** no cumplen los criterios exactos. La redacción recomendada es señalar que el modelo presenta **ajuste aproximado favorable**, pero **ajuste exacto limitado**, por lo que se continúa hacia la evaluación estructural con cautela metodológica (Henseler et al., 2015).

Decisión metodológica sobre el ajuste global del modelo saturado

Con base en los resultados del **Bootstrapping**, la decisión metodológica sobre el **modelo saturado CAITIZEN_PLS_SIN CAIL2_FIT** debe distinguir entre **ajuste aproximado** y **ajuste exacto**. El valor **SRMR = 0.047** se ubica por debajo de los umbrales convencionales de **0.08** y **0.10**, por lo que ofrece evidencia favorable de **ajuste aproximado** del modelo saturado. Sin embargo, este mismo valor supera los límites **Bootstrap HI95 = 0.033** y **HI99 = 0.034**, razón por la cual **SmartPLS4** lo marca en rojo. Esto significa que el **SRMR** cumple bajo criterios convencionales de ajuste aproximado, pero no cumple el criterio **Bootstrap** de ajuste exacto. A esta cautela se suman las discrepancias **d_ULS = 3.903** y **d_G = 1.066**, las cuales también superan sus respectivos límites **HI95** y **HI99**. En consecuencia, el **modelo saturado CAITIZEN_PLS_SIN CAIL2_FIT** puede considerarse favorable en ajuste aproximado, pero limitado en ajuste exacto, por lo que debe reportarse con cautela respecto de las pruebas exactas basadas en **Bootstrapping** (Henseler et al., 2015).

La conclusión debe evitar afirmaciones absolutas como “el modelo queda validado” o “**el ajuste es perfecto**”. Una formulación metodológicamente más adecuada sería señalar que el **modelo saturado CAITIZEN** presenta evidencia favorable de **ajuste aproximado** mediante **SRMR**, aunque las pruebas exactas basadas en **SRMR Bootstrap**, **d_ULS** y **d_G** sugieren cautela en la interpretación del ajuste global. Esta redacción reconoce la evidencia favorable sin ocultar las limitaciones detectadas por el procedimiento bootstrap. Además, conserva coherencia con la lógica del análisis **PLS-SEM**, donde las decisiones metodológicas no deben depender de un único indicador, sino de la convergencia entre criterios de medición, ajuste, significancia estadística y coherencia teórica (Hair et al., 2021).

Esta decisión no impide continuar con la evaluación del modelo estructural. El ajuste global del **modelo saturado** constituye una evidencia complementaria, no una condición única ni absoluta para interpretar el modelo completo. Por ello, aun cuando las pruebas exactas sugieren cautela, puede avanzarse hacia la evaluación estructural siempre que el investigador reporte los resultados con transparencia y continúe revisando la colinealidad, los coeficientes de trayectoria, **R²**, **f²**, la significancia estadística y la relevancia sustantiva de las relaciones. De esta manera, el **Capítulo 7** cierra la evaluación inferencial del ajuste global del **modelo saturado CAITIZEN** y abre paso al análisis estructural, bajo una postura metodológica equilibrada: continuar el análisis, pero sin afirmar ajuste global pleno ni formular conclusiones estructurales

Juan Mejía Trejo

definitivas antes de completar la evaluación posterior (Hair et al., 2019). Por tanto, el **modelo saturado CAITIZEN** no debe descartarse automáticamente; debe reportarse como un modelo con **ajuste aproximado favorable**, pero con **ajuste exacto limitado**, lo que exige continuar hacia la evaluación estructural con cautela metodológica. Ver **Tabla 7.2**.

Tabla 7.2. Valoración final del ajuste global del modelo saturado CAITIZEN

Criterio	Tipo de ajuste	Valor original	Límite / criterio de referencia	Resultado	Interpretación metodológica
SRMR	Ajuste aproximado convencional	0.047	< 0.08 criterio estricto; < 0.10 criterio aceptable	Cumple	El valor se ubica por debajo de los umbrales convencionales, por lo que sugiere ajuste aproximado favorable del modelo saturado.
SRMR bootstrap/ajuste exacto	Ajuste exacto basado en <i>Bootstrapping</i>	0.047	HI95 = 0.033; HI99 = 0.034	No cumple	Aunque el SRMR es aceptable por umbral convencional, supera los límites bootstrap; por ello, SmartPLS4 lo marca en rojo y debe reportarse con cautela.
d_ULS	Ajuste exacto basado en <i>Bootstrapping</i>	3.903	HI95 = 1.898; HI99 = 2.091	No cumple	El valor original supera los límites Bootstrap , por lo que no se cumple plenamente el ajuste exacto mediante discrepancia de mínimos cuadrados no ponderados.
d_G	Ajuste exacto basado en <i>Bootstrapping</i>	1.066	HI95 = 0.926; HI99 = 0.991	No cumple	El valor original supera los límites Bootstrap , por lo que no se cumple plenamente el ajuste exacto mediante discrepancia geodésica.
Decisión global	Integración de criterios	—	Convergencia entre ajuste aproximado y exacto	Ajuste aproximado o favorable; ajuste exacto limitado	El modelo presenta evidencia favorable de ajuste aproximado mediante SRMR, pero evidencia limitada en las pruebas exactas basadas en Bootstrapping . No se descarta automáticamente, pero debe reportarse con cautela metodológica.

Nota El **SRMR** se interpreta como índice de **ajuste aproximado** del modelo saturado. En este caso, el valor **SRMR = 0.047** cumple los umbrales convencionales de **< 0.08** y **< 0.10**, por lo que respalda un ajuste aproximado favorable. Los criterios **d_ULS** y **d_G** corresponden a diagnósticos más estrictos de **ajuste exacto** basados en **Bootstrapping**; al superar sus límites **HI95** y **HI99**, deben interpretarse como señales de cautela metodológica, no como criterios automáticos de rechazo del modelo. Por tanto, el

modelo **CAITIZEN** puede continuar hacia la evaluación estructural, siempre que se integren los resultados de cargas externas, confiabilidad, **AVE**, validez convergente, validez discriminante, colinealidad, coeficientes de trayectoria, significancia estadística y coherencia teórica (Henseler et al., 2015; Hair et al., 2019; Hair et al., 2021).

Fuente: Elaboración propia con base en resultados de **SmartPLS4**.

En síntesis, el modelo saturado **CAITIZEN** muestra un **SRMR = 0.047**, valor favorable bajo criterios convencionales de ajuste aproximado. Sin embargo, al comparar el valor original con los límites **Bootstrap**, el **SRMR**, **d_ULS** y **d_G** no cumplen plenamente los criterios de ajuste exacto. Por tanto, el modelo no debe rechazarse automáticamente, pero tampoco debe presentarse como un modelo con ajuste global pleno. La decisión metodológica adecuada es continuar hacia la evaluación estructural con cautela, considerando de manera integrada la colinealidad, los coeficientes de trayectoria, **R²**, **f²**, significancia estadística y relevancia sustantiva de las relaciones.

Criterio editorial para reportar el ajuste global en artículos científicos

En la presente obra, resulta pertinente presentar de manera amplia los criterios **SRMR**, **d_ULS** y **d_G**, porque permiten explicar al lector cómo se valora el ajuste global del **modelo saturado** en **SmartPLS4** y cómo se distinguen los índices de ajuste aproximado de las pruebas exactas basadas en **Bootstrapping**. Sin embargo, en un artículo científico empírico la decisión de reportar estos criterios debe responder al propósito del estudio, a las normas de la revista y a la economía argumentativa del manuscrito. No todo lo que se calcula en **SmartPLS4** debe incluirse necesariamente en el cuerpo principal del artículo, especialmente si su inclusión abre una discusión técnica secundaria que no es central para los objetivos de investigación.

En estudios aplicados con **PLS-SEM**, el reporte principal suele concentrarse en la calidad del modelo de medición y del modelo estructural. Por ello, si el artículo ya presenta adecuadamente **cargas externas**, **confiabilidad interna** mediante **Cronbach Alpha**, **rho_A** y **rho_c**, **validez convergente** mediante **AVE**, **validez discriminante** mediante **HTMT** o cargas cruzadas, así como **colinealidad**, **coeficientes de trayectoria**, **t-values**, **p-values**, **R²** y **f²**, entonces el manuscrito cuenta con los elementos centrales para sustentar la evaluación **PLS-SEM**. En ese contexto, los índices de ajuste global del **modelo saturado** pueden tratarse como **evidencia complementaria**, no necesariamente como una tabla principal obligatoria (Hair et al., 2019; Hair et al., 2021).

El **SRMR** puede reportarse de forma breve cuando ofrece una lectura favorable de ajuste aproximado. En el caso **CAITIZEN**, el valor **SRMR = 0.047** se encuentra por debajo del umbral convencional de **0.08**, lo que permite afirmar un ajuste aproximado aceptable del **modelo saturado** (Hu & Bentler, 1999). En un artículo científico, esta

Juan Mejía Trejo

información puede incluirse como una nota debajo de la tabla principal de resultados o como una frase breve dentro del apartado de resultados. Por ejemplo: “El modelo saturado mostró ajuste aproximado aceptable (**SRMR = 0.047 < 0.08**), lo que respalda la continuidad de la evaluación del modelo de medición y del modelo estructural”. Esta forma de reporte informa el resultado relevante sin sobrecargar el artículo con detalles técnicos adicionales.

Los criterios **d_ULS** y **d_G** corresponden a diagnósticos más estrictos de ajuste exacto basados en **Bootstrapping**. Su utilidad es mayor en capítulos metodológicos, tesis, libros, reportes técnicos o anexos, donde se busca documentar con detalle el proceso completo de evaluación del modelo. No obstante, en un artículo empírico breve, su presentación puede generar una discusión adicional que no siempre aporta valor al argumento central, especialmente cuando la revista no solicita explícitamente estos índices y cuando el modelo de medición y el modelo estructural ya han sido reportados con suficiencia. En estos casos, **d_ULS** y **d_G** pueden conservarse como evidencia de respaldo, material suplementario o información disponible para responder a evaluadores si fuese necesario (Henseler et al., 2015; Hair et al., 2021).

Por tanto, la decisión editorial no debe formularse como “*reportar u omitir*”, sino como “*reportar lo necesario según el propósito del artículo*”. Si el objetivo del manuscrito es contrastar hipótesis, explicar relaciones entre constructos o validar empíricamente un modelo aplicado, el reporte principal debe privilegiar los criterios directamente vinculados con la calidad de la medición y la evaluación estructural. En cambio, si el artículo tiene un propósito metodológico, comparativo o centrado en ajuste global, entonces sí resulta conveniente incluir una tabla específica con **SRMR**, **d_ULS** y **d_G**. Esta distinción permite mantener transparencia metodológica sin distraer al lector ni generar dudas innecesarias sobre la continuidad del estudio.

En el caso **CAITIZEN**, si el artículo no exige una tabla completa de ajuste global, una estrategia editorial adecuada consiste en reportar el **SRMR** como evidencia de ajuste aproximado aceptable y reservar **d_ULS** y **d_G** para libros, anexos, material suplementario o reportes metodológicos ampliados. Esta decisión es razonable porque el modelo ya cuenta con evidencia central de calidad de medición mediante cargas externas, confiabilidad, **AVE**, validez convergente y validez discriminante. Así, el artículo mantiene una narrativa clara y enfocada, mientras que el libro conserva la explicación metodológica completa para fines formativos.

Síntesis del procedimiento

El procedimiento de **Bootstrapping** inicia con el modelo **CAITIZEN_PLS_SIN CAIL2_FIT**, previamente configurado en **Modo A**. Desde **Calculate** → **Bootstrapping**, se selecciona el procedimiento con **5,000 submuestras**, resultados completos, prueba de **dos colas** y nivel de significancia de **0.05**. Esta configuración permite obtener no solo la significancia de rutas e indicadores, sino también los límites **Bootstrap** necesarios para valorar el ajuste global del **modelo saturado** mediante **SRMR**, **d_ULS**

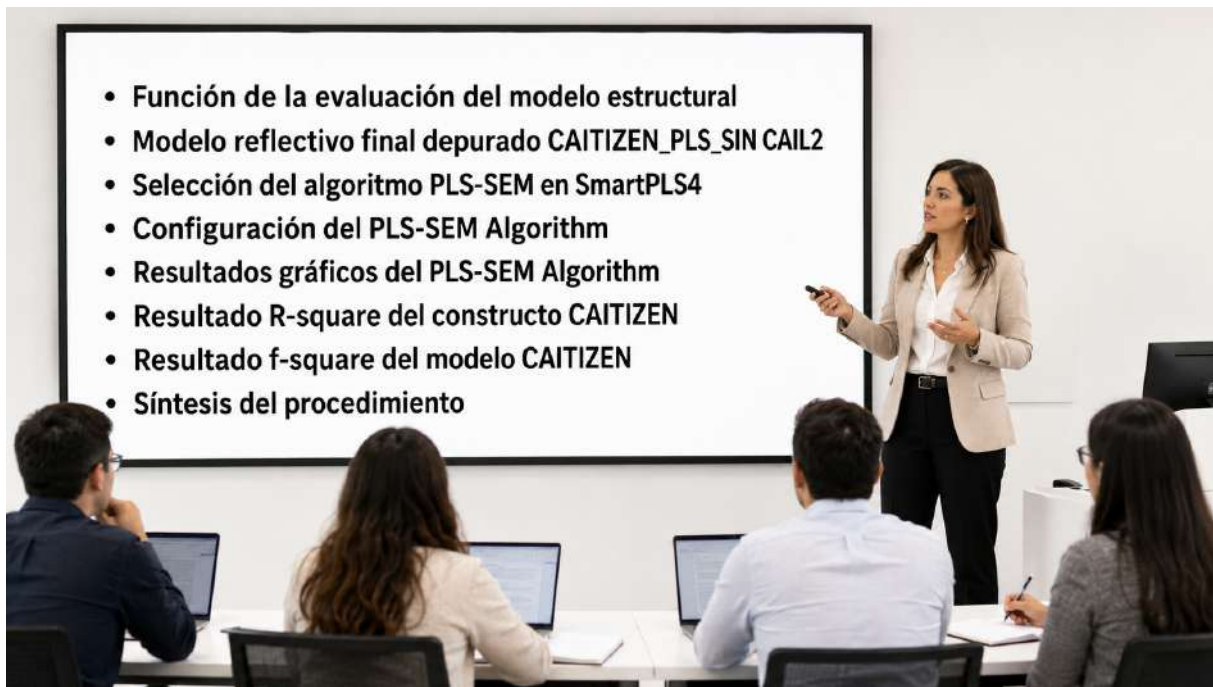
Juan Mejía Trejo

y **d_G** (Hair et al., 2021). Los resultados muestran que algunas rutas estructurales son estadísticamente significativas, especialmente **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN**, mientras que **CAIL** → **CAITIZEN** y **EAR** → **CAITIZEN** no presentan significancia directa. En cuanto al ajuste global, el **SRMR** indica un ajuste aproximado favorable, pero **d_ULS** y **d_G** no cumplen plenamente los criterios de ajuste exacto basados en límites bootstrap. Por tanto, la conclusión del capítulo debe ser equilibrada: el modelo cuenta con evidencia favorable parcial, pero requiere reportarse con cautela metodológica (Hair et al., 2019). En síntesis, el procedimiento puede resumirse en los siguientes puntos:

- El **Bootstrapping** no reemplaza al algoritmo **PLS-SEM**. El algoritmo estima los valores originales; el **Bootstrapping** evalúa su estabilidad inferencial.
- El análisis se ejecuta sobre la versión preparada en el Capítulo 6. La versión correcta es **CAITIZEN_PLS_SIN_CAIL2_FIT**, configurada como **composite en Modo A**.
- Se selecciona **Bootstrapping** desde **Calculate**. La ruta operativa es **Calculate** → **Bootstrapping**.
- Se utiliza **Complete Bootstrapping**. La opción **Complete (slower)** permite obtener resultados completos para rutas, cargas, pesos, confiabilidad, **HTMT**, **R²** y ajuste del modelo.
- Se configuran **5,000 submuestras**. Este número permite obtener estimaciones **Bootstrap** más estables.
- Se usa prueba de dos colas. La opción **Two tailed** con $\alpha = 0.05$ permite una inferencia más conservadora.
- El progreso del cálculo debe revisarse. La línea **Running saturated model fit [5000 done]** confirma que se calcularon los criterios **Bootstrap** del modelo saturado.
- La salida gráfica ofrece una lectura inicial. Los valores **p** mostrados en el modelo ayudan a identificar relaciones significativas, pero la decisión formal debe basarse en tablas.
- Las rutas significativas son **AFDJ**, **HAIC** y **MTPP** hacia **CAITIZEN**. Estas relaciones presentan **p = 0.000** en los resultados mostrados.
- Las rutas **CAIL** y **EAR** hacia **CAITIZEN** no son significativas. Sus valores **p = 0.658** y **p = 0.324** superan el umbral de **0.05**.
- El **SRMR** muestra ajuste aproximado favorable. El valor **0.047** se ubica por debajo de los umbrales **0.08** y **0.10**.
- **d_ULS** no cumple el ajuste exacto. El valor original **3.903** supera los límites **95% = 1.898** y **99% = 2.091**.
- **d_G** tampoco cumple plenamente el ajuste exacto. El valor original **1.066** supera los límites **95% = 0.926** y **99% = 0.991**.
- El modelo no debe descartarse automáticamente. La falta de ajuste exacto en **d_ULS** y **d_G** indica cautela, no inutilidad del modelo.

- **La conclusión debe ser prudente.** El modelo presenta ajuste aproximado favorable, pero evidencia limitada en las pruebas exactas basadas en ***Bootstrapping***.
- **El capítulo prepara la transición hacia la evaluación estructural del modelo.** Los resultados de rutas, significancia y ajuste global sirven como base para el siguiente capítulo sobre evaluación del modelo estructural.

CAPÍTULO 8. Evaluación del modelo estructural



El **Capítulo 8. Evaluación del modelo estructural** constituye la etapa en la que el análisis **PLS-SEM** deja de centrarse en la calidad de los indicadores y pasa a examinar la **capacidad explicativa de las relaciones entre constructos latentes**. Después de haber revisado la preparación de datos, la construcción del modelo, el modelo de medición reflectivo, el ajuste global del modelo saturado y la inferencia mediante Bootstrapping, este capítulo se enfoca en interpretar cómo los constructos antecedentes **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** contribuyen a explicar el constructo endógeno **CAITIZEN**. Esta fase permite responder no solo si el modelo está bien medido, sino qué tanto explica la ciudadanía sostenible asistida por inteligencia artificial (Hair et al., 2021).

Un punto metodológico central del capítulo es distinguir la versión correcta del modelo que debe utilizarse para la evaluación estructural. En esta etapa no debe emplearse la copia **CAITIZEN_PLS_SIN CAIL2_FIT**, configurada como composite en Modo A para el análisis de ajuste global, sino el modelo reflectivo final depurado **CAITIZEN_PLS_SIN CAIL2**. Esta precisión evita confundir los resultados del ajuste del modelo saturado con los criterios estructurales ordinarios, tales como **path coefficients**, **R^2** , **R^2 ajustado**, **f^2** , **Inner VIF** y **correlaciones entre variables latentes**. Por tanto, la trazabilidad entre versiones del modelo se convierte en una condición necesaria para reportar resultados válidos y reproducibles (Hair et al., 2019).

La evaluación estructural permite pasar de la pregunta “¿los constructos están correctamente medidos?” a la pregunta “¿qué tanto explican los constructos antecedentes al constructo dependiente?”. En el modelo **CAITIZEN**, esta transición implica analizar la dirección y magnitud de las rutas **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN**. Cada coeficiente de trayectoria expresa una relación directa estimada entre constructos latentes, pero su interpretación no debe limitarse al signo positivo o negativo. Deben considerarse de forma integrada la magnitud del efecto, la significancia estadística, la colinealidad, la varianza explicada y la coherencia teórica del modelo (Chin, 1998).

Antes de interpretar los coeficientes estructurales, el capítulo retoma la revisión de **colinealidad interna** mediante el **Inner VIF**. Esta revisión permite confirmar que los predictores del modelo no presentan redundancia crítica al explicar **CAITIZEN**. Si existiera colinealidad elevada, los coeficientes de trayectoria podrían distorsionarse, dificultando identificar el aporte individual de cada constructo antecedente. En este sentido, la evaluación estructural no comienza directamente con los valores beta, sino con la verificación de que las relaciones entre predictores permiten una interpretación estable y diferenciada. Así, el **Inner VIF** funciona como una condición previa para una lectura estructural confiable (Kock, 2015).

Uno de los criterios más relevantes del capítulo es el **coeficiente de determinación R^2** , porque indica la proporción de varianza del constructo endógeno explicada por sus predictores. En el modelo **CAITIZEN**, el valor **$R^2 = 0.682$** significa que **CAIL, EAR, AFDJ, HAIC y MTPP** explican conjuntamente el **68.2% de la varianza** de la ciudadanía sostenible asistida por inteligencia artificial. Además, el **R^2 ajustado = 0.679** muestra que la capacidad explicativa se mantiene estable al considerar el número de predictores incluidos. Este resultado sugiere una capacidad explicativa relevante para un modelo aplicado en ciencias sociales, administración, educación e inteligencia artificial (Hair et al., 2021).

El capítulo también incorpora la **descomposición de la varianza explicada**, recurso útil para identificar qué proporción del R^2 corresponde aproximadamente a cada predictor. Esta lectura complementa el R^2 global, porque permite observar que la explicación de **CAITIZEN** se concentra principalmente en **HAIC y AFDJ**, seguidos por **MTPP**, mientras que **EAR y CAIL** muestran aportes estructurales reducidos o prácticamente nulos. Esta descomposición no sustituye la significancia estadística ni el tamaño del efecto, pero ayuda a traducir el resultado global del modelo en una interpretación sustantiva sobre la contribución relativa de cada dimensión explicativa (Hair et al., 2019).

Finalmente, el capítulo revisa el **tamaño del efecto f^2** , criterio que permite valorar cuánto cambia el R^2 cuando un predictor específico se incluye o excluye del modelo. En **CAITIZEN**, **HAIC** presenta un efecto mediano, **AFDJ** un efecto pequeño cercano a moderado, **MTPP** un efecto pequeño, mientras que **CAIL y EAR** muestran efectos prácticamente nulos. Esta lectura debe ser prudente: un bajo f^2 no significa necesariamente irrelevancia conceptual, sino baja contribución directa dentro de esta

Juan Mejía Trejo

especificación estructural. Por ello, el capítulo aporta una interpretación integrada de magnitud, explicación, contribución relativa y relevancia sustantiva del modelo **CAITIZEN** (Cohen, 1992).

Función de la evaluación del modelo estructural

La evaluación del **modelo estructural** constituye la etapa en la que el investigador analiza las relaciones entre los constructos latentes del modelo, una vez que el modelo de medición ha mostrado evidencia suficiente de calidad. En los capítulos anteriores se revisaron la preparación de datos, la construcción del modelo, la ejecución inicial del algoritmo **PLS-SEM**, la colinealidad preliminar, las cargas externas, la confiabilidad interna, la validez convergente, la validez discriminante, el ajuste global del modelo saturado y la inferencia mediante **Bootstrapping**. Con esa base, el presente capítulo se concentra en interpretar la estructura explicativa del modelo **CAITIZEN**, es decir, en valorar cómo los constructos antecedentes **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** contribuyen a explicar el constructo endógeno **CAITIZEN** (Hair et al., 2021).

En esta etapa es indispensable precisar qué versión del modelo debe utilizarse en **SmartPLS4**. La evaluación del modelo estructural no se realiza sobre la copia configurada para el análisis de ajuste global, sino sobre el **modelo reflectivo final depurado**. En el caso **CAITIZEN**, esto significa trabajar con el modelo **CAITIZEN_PLS_SIN_CAIL2**, siempre que la eliminación de **CAIL2** haya sido adoptada como decisión metodológica final. Este modelo conserva la especificación reflectiva de los constructos y permite obtener los criterios estructurales mediante **PLS-SEM Algorithm**, tales como *Inner VIF*, *path coefficients*, R^2 , R^2 ajustado, f^2 y correlaciones entre variables latentes. En cambio, la copia **CAITIZEN_PLS_SIN_CAIL2_FIT**, configurada como **composite en Modo A**, debe reservarse exclusivamente para el análisis de ajuste global del modelo saturado, mediante criterios como **SRMR**, **d_ULS** y **d_G** (Hair et al., 2019). Ver **Tabla 8.1**.

Tabla 8.1. Modelos utilizados según la etapa de análisis en SmartPLS4

Modelo en SmartPLS4	Configuración	Uso metodológico
CAITIZEN_PLS_SIN_CAIL2	Reflectivo	Se utiliza para la evaluación del modelo estructural mediante PLS-SEM Algorithm: Inner VIF, path coefficients, R^2, R^2 ajustado, f^2 , correlaciones entre variables latentes, estructura e hipótesis.
CAITIZEN_PLS_SIN_CAIL2_FIT	Composite en Modo A	Se utiliza exclusivamente para el análisis de ajuste global del modelo saturado: SRMR, d_ULS y d_G .

Fuente: Elaboración propia.

Por tanto, cuando en este capítulo se haga referencia a la evaluación estructural del modelo **CAITIZEN**, debe entenderse que el análisis se realiza sobre el modelo **reflectivo final depurado**, no sobre la copia **FIT** configurada en **Modo A** para el ajuste global. Esta aclaración es necesaria porque ambos modelos provienen de la misma

estructura conceptual, pero responden a finalidades metodológicas distintas dentro de la ruta **PLS-SEM** (Hair et al., 2021).

La función principal de esta etapa no es volver a evaluar si los indicadores representan adecuadamente a sus constructos, porque esa tarea corresponde al **modelo de medición reflectivo**. Tampoco consiste únicamente en revisar si una ruta tiene un **p-value** significativo, porque la significancia estadística ya fue abordada mediante **Bootstrapping**. La evaluación estructural tiene una finalidad más amplia: determinar si las relaciones propuestas entre constructos presentan dirección esperada, magnitud interpretable, ausencia de colinealidad crítica, capacidad explicativa suficiente y relevancia sustantiva. Por ello, esta etapa permite pasar de la pregunta “¿los constructos están bien medidos?” a la pregunta “¿qué tanto explican los constructos antecedentes al constructo dependiente?” (Chin, 1998).

En el modelo **CAITIZEN**, el constructo dependiente es **CAITIZEN**, entendido como ciudadanía sostenible asistida por inteligencia artificial. Los constructos antecedentes son **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP**. Cada uno representa una dimensión explicativa distinta: alfabetización crítica en inteligencia artificial, conciencia ética y responsabilidad, conciencia de equidad y justicia de datos, colaboración creativa humano-IA y transparencia metacognitiva en prácticas de prompting. La evaluación estructural permite identificar cuáles de estas dimensiones tienen mayor peso explicativo, cuáles muestran efectos más moderados y cuáles pueden conservar valor teórico aunque su efecto directo no sea estadísticamente significativo (Hair et al., 2021).

El primer criterio es el antecedente de la evaluación estructural, debe retomarse la **revisión de colinealidad** entre predictores realizada mediante el **Inner VIF**. Esta revisión ya fue abordada en el Capítulo 4 como parte de la ejecución inicial del **PLS-SEM Algorithm** y permitió verificar que los constructos antecedentes del modelo **CAITIZEN** no presentan colinealidad crítica al explicar el constructo endógeno. En esta etapa, el **Inner VIF** no se reinterpreta como un resultado nuevo, sino como una condición metodológica previamente satisfecha que permite avanzar hacia la lectura de los **path coefficients**, el R^2 , la descomposición de la varianza explicada y el tamaño del efecto f^2 . No obstante, si el investigador ejecuta nuevamente el **PLS-SEM Algorithm** sobre el modelo reflectivo final depurado, debe confirmar que los valores **VIF** se mantienen dentro de rangos aceptables antes de interpretar los resultados estructurales (Kock, 2015).

Después de revisar la colinealidad, se interpretan los coeficientes de trayectoria o **path coefficients**. Estos coeficientes indican la dirección y magnitud de las relaciones entre constructos. En el modelo **CAITIZEN**, las rutas estructurales son **CAIL** → **CAITIZEN**, **EAR** → **CAITIZEN**, **AFDJ** → **CAITIZEN**, **HAIC** → **CAITIZEN** y **MTPP** → **CAITIZEN**. Un coeficiente positivo indica que, al aumentar el predictor, también tiende a aumentar el constructo endógeno. Sin embargo, la interpretación no debe limitarse al signo del coeficiente. También deben considerarse su magnitud, significancia estadística, tamaño del efecto y coherencia con el marco conceptual del modelo (Hair et al., 2021).

Juan Mejía Trejo

Un criterio central de la evaluación estructural es el **coeficiente de determinación R^2** . Este indicador muestra la proporción de varianza del constructo endógeno que es explicada por sus predictores. En el caso **CAITIZEN**, el R^2 permite estimar qué tanto explican conjuntamente **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** la ciudadanía sostenible asistida por inteligencia artificial. Un valor elevado de R^2 sugiere mayor capacidad explicativa del modelo, aunque su interpretación siempre debe considerar el área de estudio, la complejidad del fenómeno, la naturaleza de los constructos y el propósito de la investigación (Hair et al., 2019).

La evaluación estructural también puede complementarse con la **descomposición de la varianza explicada**. Mientras el R^2 indica cuánto explica el conjunto de predictores, la descomposición permite estimar qué proporción de esa varianza explicada corresponde aproximadamente a cada constructo antecedente. Para ello, se multiplica el **coeficiente path** de cada predictor por su correlación con el constructo endógeno. Esta operación permite distribuir el R^2 entre los predictores y facilita una interpretación más sustantiva del modelo. Así, el investigador no solo reporta que el modelo explica una determinada proporción de **CAITIZEN**, sino que puede identificar qué factores aportan mayor o menor proporción a esa explicación (Hair et al., 2021).

Otro criterio relevante es el **tamaño del efecto f^2** . Este indicador permite valorar cuánto cambia el R^2 cuando un predictor específico se incluye o se excluye del modelo. Por tanto, el f^2 no responde exactamente la misma pregunta que la descomposición de la varianza. La descomposición permite observar cómo se distribuye el R^2 entre los predictores incluidos; el f^2 permite evaluar la relevancia individual de un predictor al comparar el modelo con y sin dicho constructo. Como criterio general, valores de $f^2 \geq 0.02$ indican efecto pequeño, $f^2 \geq 0.15$ efecto mediano y $f^2 \geq 0.35$ efecto grande (Cohen, 1992).

En síntesis, la evaluación del modelo estructural permite valorar la capacidad explicativa del modelo **CAITIZEN** después de haber establecido una base suficiente de medición, ajuste e inferencia. Esta etapa no sustituye la evaluación del modelo de medición ni el **Bootstrapping**, sino que los complementa. Su función es convertir los resultados estadísticos en interpretación estructural: qué predictores explican **CAITIZEN**, cuánto explican en conjunto, qué proporción aporta cada uno, cuál es su efecto individual y cómo se articulan estos resultados con la teoría del modelo (Hair et al., 2021).

Modelo reflectivo final depurado CAITIZEN_PLS_SIN CAIL2

La muestra el modelo **CAITIZEN_PLS_SIN CAIL2** en **SmartPLS4**, es decir, la versión reflectiva final depurada que se utiliza para la evaluación del modelo estructural. En esta etapa ya no se trabaja con la copia configurada para ajuste global en **Modo A**, sino con el modelo reflectivo donde los constructos latentes se relacionan con sus indicadores mediante flechas que salen del constructo hacia los ítems. Esta

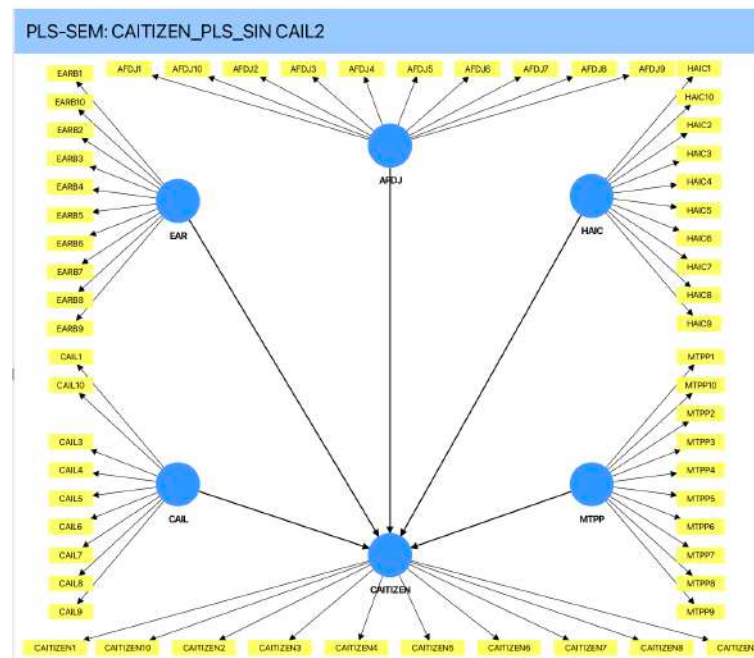
Juan Mejía Trejo

distinción es relevante porque el análisis estructural requiere conservar la lógica original de medición reflectiva para interpretar correctamente los **path coefficients**, el **R^2** , el **R^2 ajustado**, el **f^2** , las correlaciones entre variables latentes y los resultados asociados a las hipótesis del modelo (Hair et al., 2021).

En la **Figura 8.1** se observa que el constructo endógeno **CAITIZEN** recibe cinco relaciones estructurales provenientes de **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP**. Estos constructos representan las dimensiones antecedentes propuestas para explicar la ciudadanía sostenible asistida por inteligencia artificial. La estructura muestra que **CAITIZEN** no se interpreta como un constructo aislado, sino como una variable dependiente explicada por factores asociados con alfabetización crítica en **IA**, conciencia ética, justicia de datos, colaboración humano-**IA** y transparencia metacognitiva en prácticas de prompting. Por ello, esta captura funciona como punto de partida visual para justificar la evaluación estructural del modelo (Chin, 1998).

La eliminación de **CAIL2** indica que el modelo ya fue depurado previamente en la etapa de evaluación del modelo de medición. Por tanto, el análisis que se realiza en este capítulo no busca volver a justificar la exclusión de indicadores, sino interpretar la capacidad explicativa del modelo final. En términos metodológicos, esto significa que el investigador parte de un modelo de medición ya revisado y se concentra en explicar qué constructos tienen mayor peso sobre **CAITIZEN**, qué relaciones son más fuertes, qué proporción de varianza se explica y qué tamaño de efecto presenta cada predictor. Esta transición permite pasar del análisis de calidad de medición al análisis de relaciones estructurales (Hair et al., 2019).

Figura 8.1. Modelo reflectivo final depurado CAITIZEN sin CAIL2



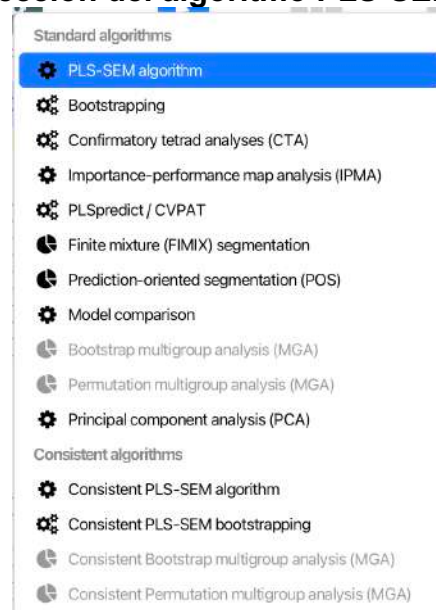
Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

Selección del algoritmo PLS-SEM en SmartPLS4

La **Figura 8.2** muestra el menú de algoritmos disponibles en **SmartPLS4**. En esta etapa se selecciona la opción **PLS-SEM algorithm**, ubicada dentro de los algoritmos estándar. Esta selección es la adecuada para obtener los resultados básicos de estimación del modelo estructural, incluyendo **path coefficients**, R^2 , R^2 ajustado, f^2 , cargas externas, criterios de calidad y correlaciones entre variables latentes. Aunque el **Bootstrapping** ya fue utilizado en el capítulo anterior para evaluar la significancia estadística mediante **t-values** y **p-values**, el **PLS-SEM algorithm** sigue siendo necesario para recuperar los valores estructurales originales del modelo (Hair et al., 2021).

Figura 8.2. Selección del algoritmo PLS-SEM en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Es importante aclarar que volver a ejecutar el **PLS-SEM algorithm** no implica repetir innecesariamente el análisis. En el Capítulo 4, el algoritmo se utilizó como ejecución inicial para estimar el modelo, revisar cargas preliminares y diagnosticar colinealidad. En el Capítulo 8, en cambio, se vuelve a consultar el mismo procedimiento sobre el modelo reflectivo final depurado para desplegar los criterios estructurales que permiten interpretar la capacidad explicativa del modelo. Por ello, la finalidad ahora no es revisar nuevamente todo el modelo de medición, sino recuperar los resultados necesarios para explicar la relación entre los constructos antecedentes y **CAITIZEN** (Hair et al., 2019).

La diferencia entre **PLS-SEM algorithm** y **Bootstrapping** debe mantenerse clara. El algoritmo **PLS-SEM** estima los valores originales del modelo con la muestra disponible; por ejemplo, los coeficientes de trayectoria, el R^2 y el f^2 . En cambio, el **Bootstrapping** genera múltiples submuestras para estimar la estabilidad y significancia de esos resultados. Por tanto, ambos procedimientos son

Juan Mejía Tréjo

complementarios. El primero responde cuánto explica el modelo y qué tan fuertes son las relaciones; el segundo responde si esas relaciones son estadísticamente significativas bajo un procedimiento de remuestreo (Hair et al., 2021).

Configuración del PLS-SEM Algorithm

La **Figura 8.3** presenta la ventana de configuración del **PLS-SEM algorithm**. Para la evaluación estructural del modelo **CAITIZEN**, la configuración mostrada es adecuada: **Weighting scheme = Path**, **Type of results = Standardized** e **Initial weight = Default**. El esquema **Path** es el más utilizado en modelos **PLS-SEM** orientados a relaciones estructurales, porque estima los constructos considerando la dirección de las relaciones especificadas en el modelo. Esta configuración permite obtener coeficientes estandarizados, adecuados para comparar la magnitud relativa de los efectos entre los distintos predictores de **CAITIZEN** (Hair et al., 2019).

Figura 8.3. Configuración del PLS-SEM Algorithm

PLS-SEM algorithm

The partial least squares structural equation modeling (PLS-SEM) is a composite-based estimator of latent variable structural equation models. Essentially, the PLS-SEM algorithm is a sequence of partial regressions.

[More info](#)

PLS setup **Data**

[Weighting scheme](#) ☐ Factor ☒ Path ☐ PCA

[Type of results](#) Standardized

[Initial weight](#) Default

[Default settings](#) ☐ Open report

Close Start calculation

Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

La opción **Standardized** es especialmente importante porque permite interpretar los coeficientes de trayectoria en una escala común. Esto significa que los efectos de **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** sobre **CAITIZEN** pueden compararse directamente. Por ejemplo, un coeficiente estandarizado de **0.374** puede interpretarse como un efecto directo más fuerte que uno de **0.164**, siempre que ambos se refieran al mismo constructo endógeno. En modelos aplicados a ciencias sociales, administración e innovación, esta estandarización facilita la lectura sustantiva de las relaciones entre variables latentes (Chin, 1998).

La configuración **Initial weight = Default** es apropiada cuando no existe una razón metodológica específica para modificar los pesos iniciales. En una aplicación didáctica y replicable, como la desarrollada en este libro, mantener los valores por defecto ayuda a que el lector reproduzca el procedimiento de manera sencilla. Una vez configurado el algoritmo, se selecciona **Start calculation** para ejecutar el modelo. A partir de esta ejecución se obtienen los resultados que serán consultados en las secciones siguientes: coeficientes de trayectoria, **R²**, **R² ajustado**, **f²**, cargas externas y correlaciones entre variables latentes (Hair et al., 2021).

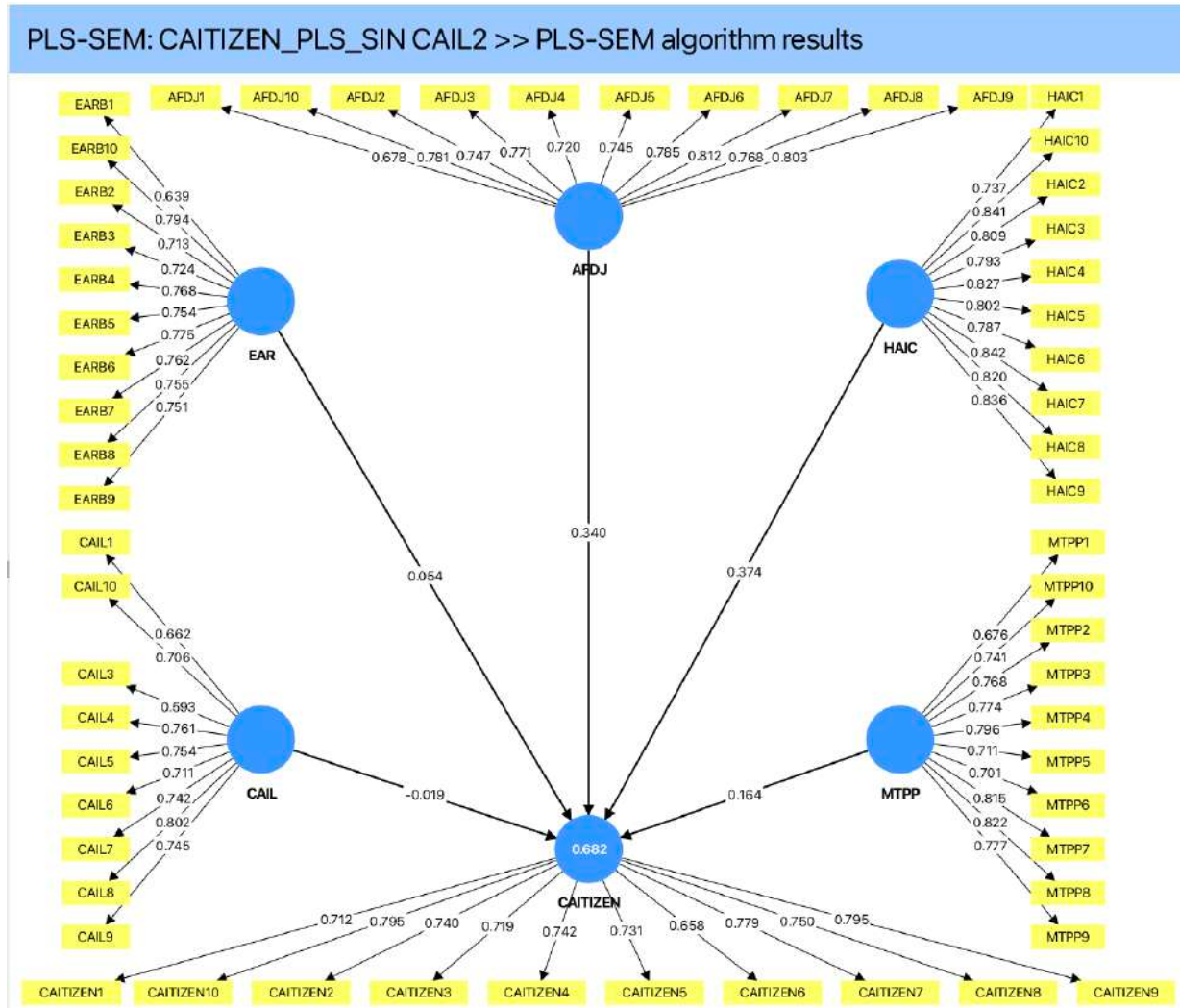
Resultados gráficos del PLS-SEM Algorithm

La **Figura 8.4** muestra los resultados gráficos generados después de ejecutar el **PLS-SEM algorithm** sobre el modelo **CAITIZEN_PLS_SIN CAIL2**. En el diagrama se observan los coeficientes de trayectoria entre los constructos antecedentes y el constructo endógeno **CAITIZEN**. Los valores visibles son **AFDJ → CAITIZEN = 0.340**, **HAIC → CAITIZEN = 0.374**, **MTPP → CAITIZEN = 0.164**, **EAR → CAITIZEN = 0.054** y **CAIL → CAITIZEN = -0.019**. Estos coeficientes indican la dirección y magnitud de las relaciones estructurales estimadas con la muestra original (Hair et al., 2021).

La lectura inicial del diagrama sugiere que **HAIC** presenta el coeficiente directo más alto sobre **CAITIZEN**, seguido por **AFDJ** y **MTPP**. En contraste, **EAR** muestra un coeficiente positivo muy bajo y **CAIL** presenta un coeficiente negativo cercano a cero. Esta lectura no debe confundirse todavía con significancia estadística, porque los **t-values** y **p-values** proceden del **Bootstrapping**. El diagrama del algoritmo **PLS-SEM** permite observar la magnitud original de las rutas, pero la inferencia estadística debe integrarse posteriormente con los resultados de remuestreo (Hair et al., 2019).

En el centro del constructo **CAITIZEN** aparece el valor **0.682**, que corresponde al **R²** del constructo endógeno. Esto significa que, en el modelo reflectivo final depurado mostrado en la captura, los predictores **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** explican conjuntamente el **68.2%** de la varianza de **CAITIZEN**. Este valor representa la capacidad explicativa global del modelo estructural. Sin embargo, el **R²** no indica por sí solo cuánto aporta cada predictor; para ello es necesario complementar el análisis con la descomposición de la varianza explicada y el tamaño del efecto **f²** (Chin, 1998).

Figura 8.4. Resultados gráficos del PLS-SEM Algorithm



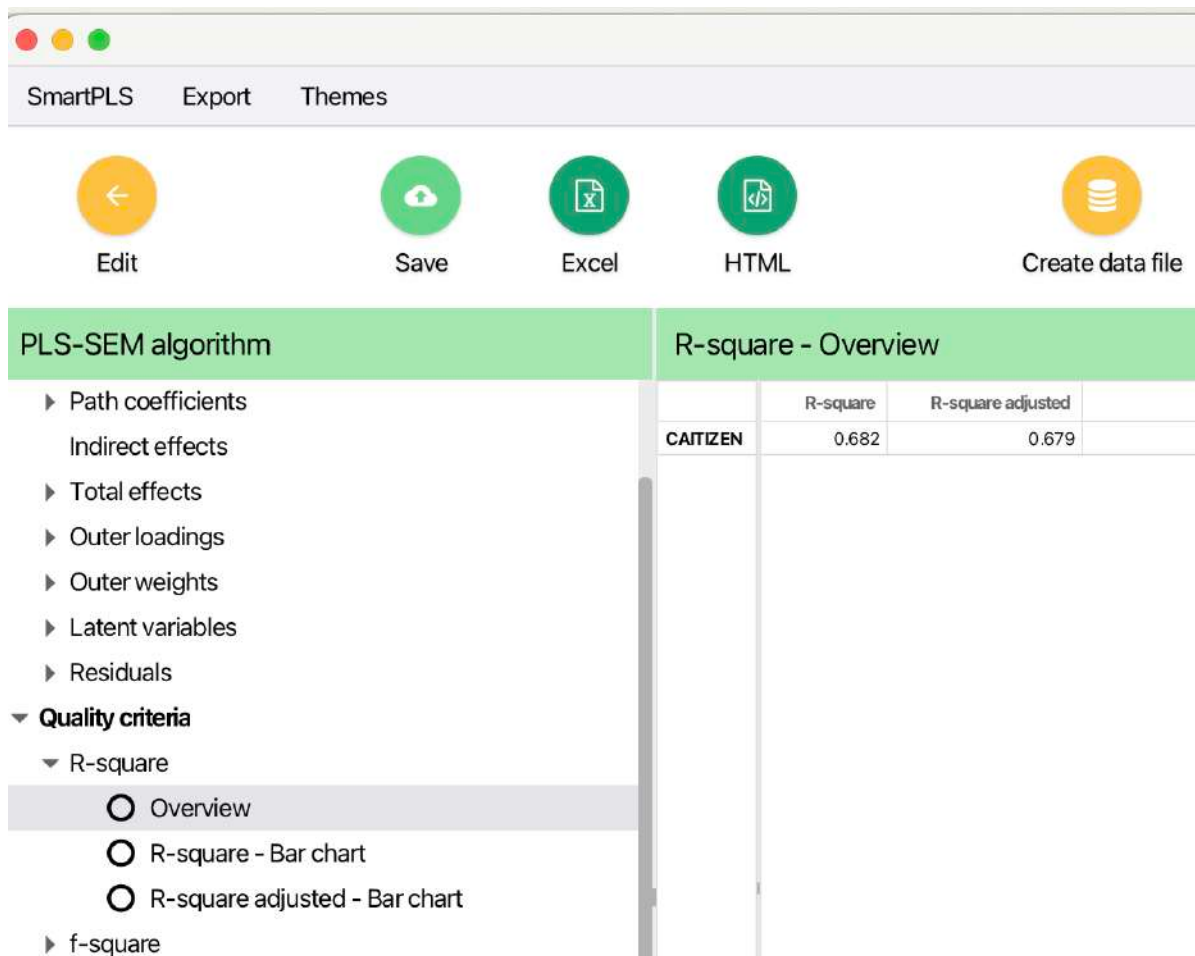
Fuente: Elaboración propia con **SmartPLS4**.

El diagrama también muestra las cargas externas de los indicadores reflectivos. Aunque estas cargas ya fueron evaluadas en el capítulo correspondiente al modelo de medición, su presencia en el gráfico confirma que el modelo sigue conservando la estructura reflectiva. En esta etapa, las cargas no son el foco principal de interpretación, pero ayudan a verificar que el análisis estructural se está realizando sobre el modelo correcto. Por ello, la figura debe explicarse como una salida integral del algoritmo, pero el énfasis del **Capítulo 8** debe ponerse en rutas estructurales, R^2 , f^2 y contribución explicativa de los predictores (Hair et al., 2021).

Resultado R-square del constructo CAITIZEN

La **Figura 8.5** muestra el resultado **R-square - Overview** dentro de los criterios de calidad del **PLS-SEM algorithm**. En esta salida se observa que el constructo **CAITIZEN** presenta un $R^2 = 0.682$ y un R^2 ajustado = **0.679**. El R^2 indica la proporción de varianza del constructo endógeno explicada por sus predictores. En este caso, el modelo explica el **68.2%** de la varianza de **CAITIZEN**, lo que representa una capacidad explicativa relevante para un modelo aplicado en ciencias sociales, administración, educación e inteligencia artificial (Hair et al., 2019).

Figura 8.5. Resultado R-square del constructo CAITIZEN



Fuente: Elaboración propia con **SmartPLS4**.

El R^2 ajustado = **0.679** corrige el valor de R^2 considerando el número de predictores incluidos en el modelo. Esta diferencia mínima entre R^2 y R^2 ajustado indica que la capacidad explicativa del modelo no se reduce de manera importante al considerar la cantidad de constructos antecedentes. En otras palabras, el modelo mantiene una explicación robusta de **CAITIZEN** aun después de ajustar por el número de predictores.

Juan Mejía Trejo

Este resultado es útil porque evita interpretar el R^2 únicamente como una cifra alta, y permite valorar si el modelo mantiene estabilidad explicativa (Hair et al., 2021).

Ejemplo didáctico de descomposición de la varianza explicada

La descomposición de la varianza explicada permite responder una pregunta que el R^2 por sí solo no contesta: ¿qué proporción del total explicado corresponde a cada predictor? Por ejemplo, si un modelo reportara $R^2 = 0.682$, ello significaría que los predictores explican conjuntamente el **68.2%** de la varianza del constructo endógeno. Sin embargo, ese valor global no muestra cuánto aporta **CAIL**, **EAR**, **AFDJ**, **HAIC** o **MTPP** por separado. Para aproximar esa distribución, se multiplica el **coeficiente path** de cada predictor por su correlación con el constructo endógeno (Hair et al., 2021).

La fórmula operativa es la siguiente: **varianza explicada = coeficiente path × correlación con el constructo endógeno**. Por ejemplo, para **HAIC** → **CAITIZEN** el coeficiente path de **0.374** (Ver **Figura 8.6**) y la correlación entre **HAIC** y **CAITIZEN** (Ver **Figura 8.7**) es **0.765**.

Figura 8.7. Matriz de correlaciones entre las variables latentes del modelo CAITIZEN

SmartPLS Export Themes

 Edit

 Save

 Excel

 HTML

 Create data file

 Compare

PLS-SEM algorithm

Path coefficients - Matrix

▼ Graphical

Graphical output

▼ Final results

▼ Path coefficients

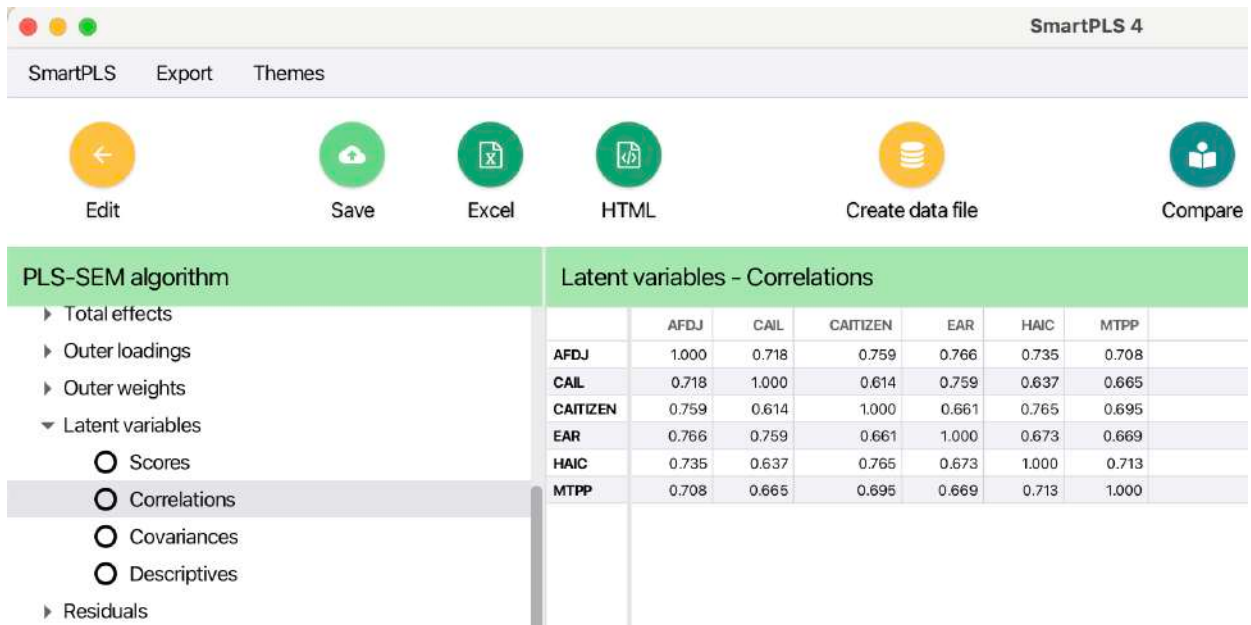
☒ Matrix

☐ List

☐ Bar chart

	AFDJ	CAIL	CAITIZEN	EAR	HAIC	MTPP
AFDJ			0.340			
CAIL			-0.019			
CAITIZEN						
EAR			0.054			
HAIC			0.374			
MTPP			0.164			

Fuente: Elaboración propia con **SmartPLS4**.

Figura 8.7. Matriz de coeficientes *path* del modelo CAITIZEN

Fuente: Elaboración propia con **SmartPLS4**.

El procedimiento práctico en **SmartPLS4** consiste en ejecutar *primero PLS-SEM Algorithm* y abrir los resultados de *Path coefficients*. De ahí se copian a Excel los coeficientes de las rutas que llegan al constructo endógeno que se desea explicar; en este caso, **CAITIZEN**. Después se entra a *Latent variables* → *Correlations* y se copian las correlaciones entre cada predictor y **CAITIZEN**. En Excel se multiplica cada coeficiente path por su correlación correspondiente. Finalmente, se suman las contribuciones obtenidas; la suma debe aproximarse al valor de R^2 del constructo endógeno (Hair et al., 2021). La lógica puede expresarse con la siguiente **Tabla 8.2**.

Tabla 8.2. Descomposición de la varianza explicada de CAITIZEN por predictor

Predictor	Coeficiente path (β) hacia CAITIZEN	Correlación (r) hacia CAITIZEN	Varianza explicada aproximada	
			$(\beta \times r)$	%
CAIL	-0.019	0.614	-0.0117	-1.17%
EAR	0.054	0.661	0.0357	3.57%
AFDJ	0.340	0.759	0.2581	25.81%
HAIC	0.374	0.765	0.2861	28.61%
MTPP	0.164	0.695	0.1140	11.4%
Total aproximado	----	----	0.6822	68.22%

Fuente: Elaboración propia

La descomposición de la varianza debe entenderse como una herramienta interpretativa complementaria. Su utilidad consiste en mostrar la contribución relativa de cada predictor al total explicado por el modelo. No sustituye la evaluación de significancia mediante *Bootstrapping*, ni reemplaza el tamaño del efecto f^2 . La

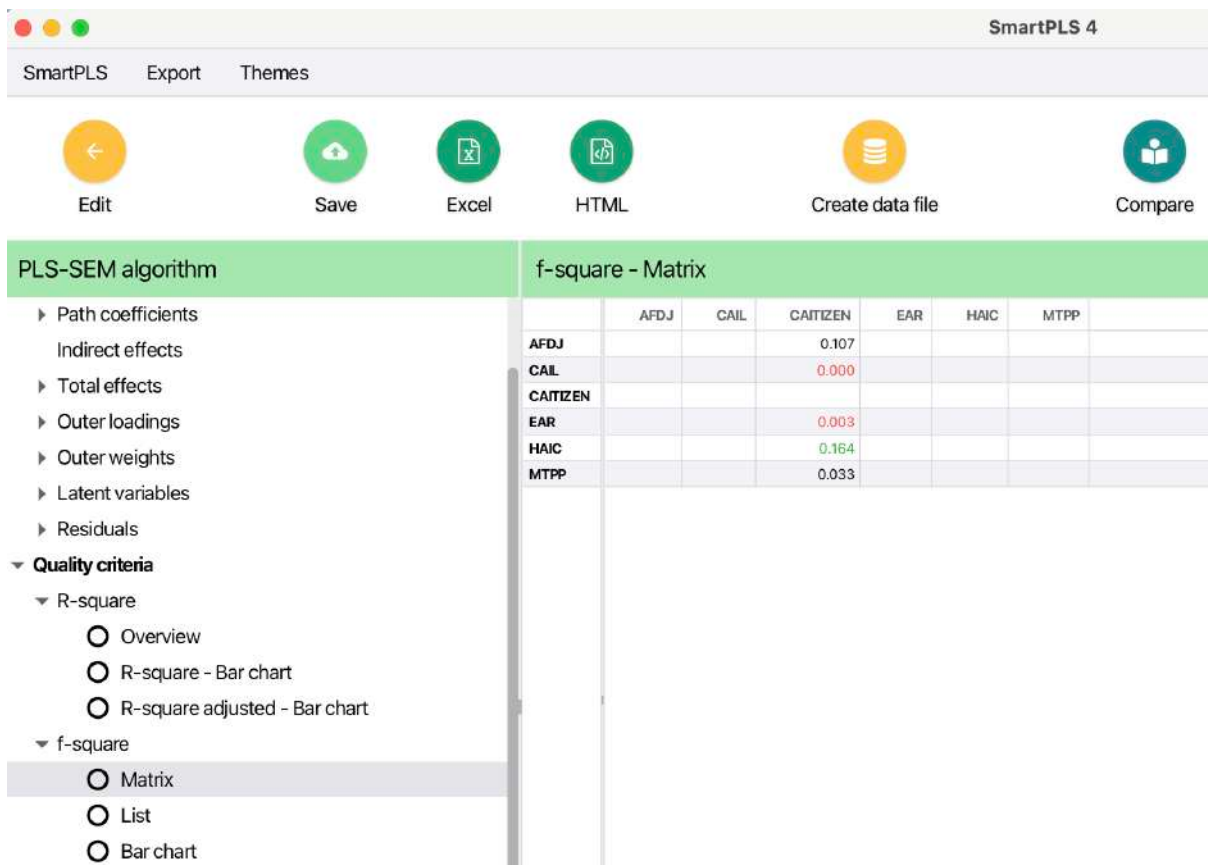
Juan Mejía Trejo

significancia indica si una relación es estadísticamente distinta de cero; el f^2 indica cuánto cambia el R^2 cuando se elimina un predictor; y la descomposición de varianza permite observar cómo se distribuye el R^2 entre los predictores incluidos. Por ello, los tres criterios deben leerse de manera integrada (Cohen, 1992).

Resultado f-square del modelo CAITIZEN

La **Figura 8.8** muestra la matriz **f-square** dentro de los criterios de calidad del **PLS-SEM algorithm**. El f^2 permite valorar el tamaño del efecto de cada predictor sobre el constructo endógeno. A diferencia del R^2 , que muestra la capacidad explicativa global del conjunto de predictores, el f^2 evalúa la contribución individual de cada constructo antecedente. En términos prácticos, indica cuánto se reduciría la capacidad explicativa del modelo si un predictor específico se excluyera de la explicación de **CAITIZEN** (Cohen, 1992).

Figura 8.8. Resultado f-square del modelo CAITIZEN



Fuente: Elaboración propia con **SmartPLS4**.

En la matriz se observan los siguientes valores: **AFDJ** → **CAITIZEN** = 0.107, **CAIL** → **CAITIZEN** = 0.000, **EAR** → **CAITIZEN** = 0.003, **HAIC** → **CAITIZEN** = 0.164 y **MTPP** → **CAITIZEN** = 0.033. De acuerdo con los criterios usuales, valores de $f^2 \geq 0.02$ indican

Juan Mejía Trejo

efecto pequeño, $f^2 \geq 0.15$ efecto mediano y $f^2 \geq 0.35$ efecto grande. Bajo esta referencia, **HAIC** presenta un efecto mediano, **AFDJ** un efecto pequeño cercano a moderado, **MTPP** un efecto pequeño, mientras que **CAIL** y **EAR** muestran efectos prácticamente nulos sobre **CAITIZEN** (Hair et al., 2021).

Estos resultados son congruentes con los coeficientes de trayectoria observados en el diagrama del algoritmo. **HAIC** muestra el coeficiente path más alto y también el mayor f^2 , lo que refuerza su importancia estructural dentro del modelo. **AFDJ** también presenta una contribución relevante, aunque menor que **HAIC**. **MTPP** aporta de manera más limitada, pero todavía con un efecto pequeño. En cambio, **CAIL** y **EAR** tienen valores de f^2 muy bajos, lo que indica que su exclusión tendría poco impacto sobre el R^2 de **CAITIZEN** en esta especificación particular del modelo (Hair et al., 2019).

La interpretación del f^2 debe hacerse con cautela. Un valor bajo no significa necesariamente que el constructo carezca de importancia conceptual. Puede indicar que su efecto directo sobre **CAITIZEN** queda absorbido por otros predictores, que su influencia es indirecta, que comparte varianza con otros constructos o que su papel teórico requiere una especificación distinta. Por ello, los valores de f^2 deben interpretarse junto con los **path coefficients**, los **t-values**, los **p-values**, la descomposición de varianza y la fundamentación teórica del modelo. En estudios sociales y administrativos, la relevancia metodológica y la relevancia conceptual no siempre coinciden plenamente (Hair et al., 2021).

Síntesis del procedimiento

Las capturas muestran una secuencia metodológica coherente para la evaluación del modelo estructural. Primero, se abre el modelo reflectivo final depurado **CAITIZEN_PLS_SIN CAIL2**. Segundo, se selecciona **PLS-SEM algorithm** dentro de los algoritmos estándar de **SmartPLS4**. Tercero, se ejecuta el algoritmo con esquema **Path**, resultados **Standardized** y peso inicial **Default**. Cuarto, se interpretan los coeficientes estructurales y el R^2 del diagrama. Quinto, se consulta **R-square - Overview** para reportar R^2 y R^2 ajustado. Sexto, se revisa la matriz **f-square** para valorar el tamaño del efecto de cada predictor (Hair et al., 2021).

Esta secuencia permite distinguir entre estimación, explicación e inferencia. El **PLS-SEM algorithm** proporciona los valores originales del modelo: **path coefficients**, R^2 , R^2 ajustado, f^2 y correlaciones entre variables latentes. El **Bootstrapping**, revisado previamente, permite evaluar la significancia estadística de las relaciones mediante **t-values** y **p-values**. La descomposición de la varianza explicada, por su parte, permite distribuir el R^2 entre los predictores mediante el producto entre coeficiente path y correlación con el constructo endógeno. Así, el Capítulo 8 integra magnitud, explicación, contribución relativa y tamaño del efecto dentro de una lectura estructural del modelo **CAITIZEN** (Hair et al., 2019), del que presentamos una síntesis de procedimiento:

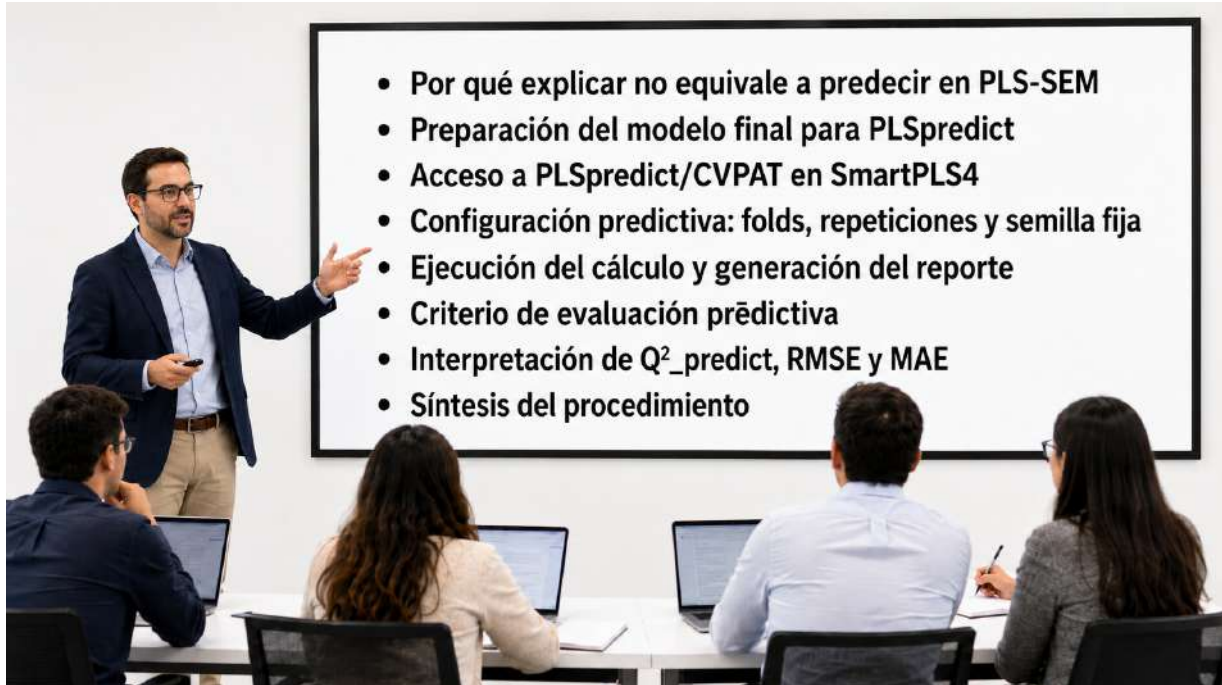
Juan Mejía Trejo

- Abrir en **SmartPLS4** el modelo reflectivo final depurado **CAITIZEN_PLS_SIN CAIL2**. Esta es la versión correcta para evaluar el modelo estructural, porque conserva la especificación reflectiva de los constructos y permite interpretar las relaciones entre **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**.
- No utilizar para esta etapa la copia **CAITIZEN_PLS_SIN CAIL2_FIT**, ya que dicha versión fue configurada como **composite en Modo A** únicamente para revisar el ajuste global del modelo saturado mediante **SRMR**, **d_ULS** y **d_G**.
- Seleccionar la opción **Calculate** → **PLS-SEM algorithm**. Aunque el algoritmo ya se utilizó previamente en el **Capítulo 4**, en esta etapa se vuelve a ejecutar con una finalidad distinta: desplegar los resultados estructurales finales del modelo depurado.
- Configurar el algoritmo con los parámetros: **Weighting scheme = Path**, **Type of results = Standardized** e **Initial weight = Default**. Esta configuración permite obtener coeficientes estandarizados y comparables entre los distintos predictores del constructo endógeno **CAITIZEN**.
- Ejecutar el cálculo mediante **Start calculation**. Una vez procesado el modelo, revisar primero el diagrama gráfico de resultados para identificar los **path coefficients** de las rutas estructurales.
- Registrar los coeficientes de trayectoria del modelo: **AFDJ** → **CAITIZEN** = 0.340, **HAIC** → **CAITIZEN** = 0.374, **MTPP** → **CAITIZEN** = 0.164, **EAR** → **CAITIZEN** = 0.054 y **CAIL** → **CAITIZEN** = -0.019.
- Interpretar los **path coefficients** en términos de dirección y magnitud. En el modelo **CAITIZEN**, **HAIC** presenta el coeficiente directo más alto, seguido por **AFDJ** y **MTPP**. En contraste, **EAR** muestra un efecto positivo muy bajo y **CAIL** un efecto negativo cercano a cero.
- Recordar que los **path coefficients** del **PLS-SEM algorithm** muestran la magnitud original de las relaciones, pero no sustituyen la inferencia estadística. Los **t-values** y **p-values** deben integrarse desde los resultados de **Bootstrapping**.
- Consultar **Quality criteria** → **R-square** → **Overview** para obtener el coeficiente de determinación del constructo endógeno. En el modelo actual, **CAITIZEN** presenta **R² = 0.682** y **R² ajustado = 0.679**.
- Interpretar el **R² = 0.682** como evidencia de que los constructos antecedentes **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTPP** explican conjuntamente el **68.2%** de la varianza de **CAITIZEN**.
- Interpretar el **R² ajustado = 0.679** como una corrección del **R²** que considera el número de predictores incluidos. La diferencia mínima entre **R²** y **R² ajustado** indica que la capacidad explicativa del modelo se mantiene estable.
- Consultar **Quality criteria** → **f-square** → **Matrix** para revisar el tamaño del efecto de cada predictor sobre **CAITIZEN**.
- Registrar los valores **f²** del modelo: **AFDJ** → **CAITIZEN** = 0.107, **CAIL** → **CAITIZEN** = 0.000, **EAR** → **CAITIZEN** = 0.003, **HAIC** → **CAITIZEN** = 0.164 y **MTPP** → **CAITIZEN** = 0.033.
- Interpretar los valores **f²** con los criterios convencionales: **f² ≥ 0.02** indica efecto pequeño, **f² ≥ 0.15** indica efecto mediano y **f² ≥ 0.35** indica efecto grande.
- Con base en esos criterios, interpretar **HAIC** como predictor con efecto mediano, **AFDJ** como predictor con efecto pequeño cercano a moderado, **MTPP** como

predictor con efecto pequeño, y **CAIL** y **EAR** como predictores con efecto prácticamente nulo sobre **CAITIZEN**.

- Consultar **Latent variables** → **Correlations** para obtener las correlaciones entre cada predictor y el constructo endógeno **CAITIZEN**. Estas correlaciones se requieren para calcular la descomposición de la varianza explicada.
- Copiar a Excel los **path coefficients** de las rutas hacia **CAITIZEN** y las correlaciones de cada predictor con **CAITIZEN**.
- Calcular en Excel la descomposición de la varianza explicada mediante la fórmula: **varianza explicada = coeficiente path × correlación con CAITIZEN**.
- Expresar cada contribución como proporción o porcentaje. Por ejemplo, si un predictor tiene un coeficiente path de **0.374** y una correlación con **CAITIZEN** de **0.800**, su contribución aproximada sería $0.374 \times 0.800 = 0.299$, equivalente a **29.9%** de varianza explicada.
- Sumar las contribuciones aproximadas de todos los predictores. La suma debe aproximarse al valor de $R^2 = 0.682$, es decir, al **68.2%** de varianza explicada de **CAITIZEN**.
- Distinguir claramente la descomposición de varianza del tamaño del efecto f^2 . La descomposición muestra cómo se distribuye el R^2 entre los predictores; el f^2 muestra cuánto cambia el R^2 cuando un predictor específico se incluye o se excluye del modelo.
- Integrar los resultados del **PLS-SEM algorithm** con los resultados del **Bootstrapping**. El algoritmo proporciona **path coefficients**, R^2 , R^2 ajustado, f^2 y correlaciones; el **Bootstrapping** proporciona **t-values**, **p-values** e intervalos de confianza.
- Elaborar la interpretación estructural final del modelo **CAITIZEN** considerando conjuntamente magnitud de las rutas, significancia estadística, varianza explicada, contribución relativa y tamaño del efecto.
- Concluir que el **Capítulo 8** no repite la evaluación del modelo de medición ni sustituye el **Bootstrapping**, sino que integra los resultados necesarios para explicar qué constructos contribuyen más a la ciudadanía sostenible asistida por inteligencia artificial.
- Reportar con cautela los constructos que presentan baja contribución estructural. Un valor bajo de **path coefficient** o f^2 no significa necesariamente que el constructo carezca de importancia teórica; puede indicar que su efecto directo es débil, que comparte varianza con otros predictores o que su papel podría ser indirecto.
- Cerrar el análisis estructural destacando que, en el modelo **CAITIZEN**, la explicación de **CAITIZEN** se concentra principalmente en **HAIC**, **AFDJ** y, en menor medida, **MTPP**, mientras que **CAIL** y **EAR** requieren una interpretación conceptual más cuidadosa.

CAPÍTULO 9. De la explicación estructural a la predicción del modelo



El Capítulo 9. De la explicación estructural a la predicción del modelo mediante *PLSpredict/CVPAT* introduce una etapa metodológica necesaria para distinguir entre **capacidad explicativa** y **capacidad predictiva** en **PLS-SEM**. En los capítulos previos se evaluaron la calidad del modelo de medición, el ajuste global, la inferencia mediante Bootstrapping y la estructura explicativa del modelo. Sin embargo, **explicar no equivale automáticamente a predecir**. Un modelo puede mostrar coeficientes significativos, R^2 aceptable y tamaños de efecto relevantes, pero aun así presentar errores elevados al anticipar indicadores observables. Por ello, este capítulo permite valorar si **CAITIZEN** no solo explica relaciones entre constructos, sino si también ofrece **desempeño predictivo** sobre sus indicadores endógenos (Shmueli, 2010).

La importancia del análisis predictivo radica en evitar una **interpretación excesiva de los resultados estructurales**. En **PLS-SEM**, los **path coefficients**, R^2 y f^2 permiten comprender la dirección, magnitud y relevancia de las relaciones entre constructos latentes; no obstante, estos criterios se orientan principalmente a la **explicación interna del modelo**. La predicción, en cambio, exige valorar qué tan bien el modelo anticipa valores observables mediante procedimientos de **validación cruzada**. Así, *PLSpredict* no repite la evaluación estructural del capítulo anterior, sino que la complementa al responder una pregunta distinta: si el modelo tiene **utilidad anticipatoria** sobre los indicadores de **CAITIZEN** (Hair et al., 2019).

En el caso del modelo **CAITIZEN**, la evaluación predictiva debe realizarse sobre el **modelo final depurado CAITIZEN_PLS_SIN CAIL2**, no sobre versiones provisionales ni sobre la copia FIT utilizada para el ajuste global del modelo saturado. Esta precisión garantiza **trazabilidad metodológica**, porque cada versión del modelo cumple una función distinta dentro de la ruta **PLS-SEM**. La predicción debe ejecutarse sobre el modelo que conserva la **especificación teórica final**, sus constructos reflectivos, sus indicadores depurados y sus relaciones estructurales. De esta manera, la evaluación predictiva se conecta coherentemente con la **teoría, la medición y la estructura explicativa** ya revisadas (Chin, 1998).

El acceso a **PLSpredict/CVPAT** desde el menú **Calculate** de **SmartPLS4** marca el tránsito del **análisis explicativo** hacia el **análisis predictivo**. Mientras el algoritmo **PLS-SEM** estima relaciones estructurales y el **Bootstrapping** evalúa su estabilidad inferencial, **PLSpredict** se centra en los **errores de predicción de los indicadores endógenos**. Esta diferencia es relevante porque un modelo puede explicar una proporción importante de varianza y, aun así, no superar el desempeño de un **modelo lineal de referencia**. Por tanto, **PLSpredict** permite someter el modelo **CAITIZEN** a una evaluación más exigente de **utilidad empírica** y no solo de coherencia estructural (Shmueli et al., 2019).

La configuración del análisis predictivo constituye un paso fundamental del capítulo. Para el caso **CAITIZEN**, se propone utilizar **17 folds**, **10 repeticiones**, tipo de predicción **Early antecedents** y generador aleatorio **Fixed seed**. Esta configuración responde al tamaño muestral disponible de **511 casos** y permite dividir la muestra en secciones de validación con aproximadamente **30 observaciones**. Con ello, se busca reducir la sensibilidad a particiones demasiado pequeñas y fortalecer la estabilidad de los **errores predictivos**. Además, el uso de semilla fija favorece la **reproducibilidad del análisis**, condición indispensable para documentar decisiones metodológicas en investigaciones aplicadas (Green, 1991).

La interpretación de los resultados debe iniciar con **Q²_predict**, porque este criterio indica si el modelo tiene **relevancia predictiva** para los indicadores endógenos. Un valor positivo sugiere que el modelo predice mejor que una referencia ingenua basada en promedios. En **CAITIZEN**, el interés se concentra en los indicadores del constructo dependiente, ya que representan las manifestaciones observables de la **ciudadanía sostenible asistida por inteligencia artificial**. Posteriormente, se revisan los errores **RMSE** y **MAE**, los cuales permiten valorar la magnitud del **error predictivo**. Esta lectura evita sostener afirmaciones predictivas solo con base en **R²** o significancia estructural (Hair et al., 2021).

Finalmente, el capítulo enfatiza la comparación entre los errores del modelo **PLS-SEM** y los errores del modelo lineal de referencia **LM**. Si el **PLS-SEM** presenta valores menores de **RMSE** y **MAE** que el **LM**, puede sostenerse que el modelo tiene **desempeño predictivo favorable**. Esta comparación es central porque **Q²_predict** positivo solo demuestra relevancia predictiva básica, mientras que la comparación contra **LM** permite valorar si el modelo supera una alternativa simple. En síntesis, el

capítulo completa la evaluación metodológica del modelo **CAITIZEN** al integrar **explicación, predicción, validación cruzada y errores de desempeño** dentro de una lectura aplicada y académicamente defendible (Ringle et al., 2012).

Por qué explicar no equivale a predecir en PLS-SEM

La distinción entre **explicación y predicción** es fundamental porque evita una interpretación excesiva de los resultados **PLS-SEM**. En muchos estudios aplicados, el investigador observa coeficientes path significativos, un R^2 aceptable o tamaños de efecto f^2 relevantes y concluye que el modelo “predice”. Sin embargo, esos indicadores muestran principalmente **capacidad explicativa interna**, no necesariamente desempeño predictivo. Explicar significa comprender la relación entre constructos dentro de una estructura teórica; predecir implica evaluar qué tan bien el modelo anticipa valores observables. Por ello, este capítulo es necesario: permite pasar de “*qué constructos explican el fenómeno*” a “*qué tan bien el modelo predice los indicadores endógenos*” (Shmueli, 2010).

El capítulo anterior evaluó el modelo estructural mediante **coeficientes path, R^2 y f^2** . Esa evaluación permite identificar qué constructos tienen mayor peso explicativo sobre el constructo endógeno y qué tan relevante es cada predictor dentro de la estructura propuesta. No obstante, esa información no responde por completo si el modelo tiene **utilidad predictiva**. Un modelo puede explicar una proporción importante de varianza y, aun así, producir errores predictivos altos cuando se evalúan sus indicadores. Por esta razón, el análisis predictivo no repite la evaluación estructural, sino que la complementa. Su función es verificar si el modelo también ofrece desempeño anticipatorio aceptable (Hair et al., 2019).

La explicación se concentra en la **lógica teórica del modelo**. En **PLS-SEM**, explicar implica justificar conceptualmente las relaciones entre constructos, estimar la dirección e intensidad de esas relaciones y valorar la varianza explicada de los constructos dependientes. En el modelo **CAITIZEN**, esta dimensión permite analizar cómo las capacidades críticas, éticas, colaborativas, metacognitivas y de justicia de datos se relacionan con la ciudadanía sostenible asistida por inteligencia artificial. Por tanto, la explicación ayuda a comprender **por qué** ciertos constructos son relevantes dentro del modelo. Sin embargo, esta comprensión todavía no demuestra si el modelo puede anticipar adecuadamente los indicadores de **CAITIZEN** (Chin, 1998).

La predicción tiene otra finalidad metodológica. Su interés principal no es confirmar si una hipótesis estructural resulta significativa, sino evaluar si el modelo produce valores predictivos con **errores aceptables**. Para ello, PLSpredict incorpora criterios como **$Q^2_{predict}$, RMSE, MAE** y la comparación entre el error del modelo **PLS-SEM** y el error de un modelo lineal de referencia. Esta evaluación permite saber si el modelo tiene capacidad predictiva real sobre los indicadores endógenos. En consecuencia,

PLSpredict no sustituye el análisis de R^2 , f^2 o **Bootstrapping**; lo amplía mediante una prueba específica de **desempeño predictivo** (Shmueli et al., 2019).

La importancia del tema radica en que un modelo puede tener buena apariencia estadística y aun así no justificar afirmaciones predictivas. Si el investigador reporta únicamente **R^2 y coeficientes significativos**, puede sostener que el modelo explica el fenómeno, pero no que predice adecuadamente sus indicadores. Para afirmar capacidad predictiva se necesita evidencia adicional. Esta distinción fortalece el reporte académico porque obliga a separar tres niveles: primero, la calidad del modelo de medición; segundo, la significancia y relevancia del modelo estructural; tercero, el desempeño predictivo de los indicadores endógenos. Así, el análisis evita confundir **explicación, significancia y predicción** como si fueran equivalentes (Ringle et al., 2012).

En modelos aplicados y emergentes, la diferencia entre explicación y predicción adquiere mayor relevancia. En administración, innovación, educación superior e inteligencia artificial, los constructos suelen representar capacidades, percepciones o disposiciones difíciles de observar directamente. Por ello, no basta con comprobar que las relaciones estructurales sean significativas; también es importante valorar si el modelo produce **predicciones empíricamente aceptables**. Esta lógica permite que **PLS-SEM** no se limite a explicar relaciones internas, sino que avance hacia una evaluación más completa del desempeño del modelo. En este sentido, la predicción fortalece la utilidad aplicada del análisis estructural (Reinartz et al., 2009).

La diferencia entre explicación y predicción también tiene implicaciones para la **contribución teórica**. Un modelo explicativo contribuye cuando organiza relaciones conceptuales, aclara mecanismos y permite comprender mejor un fenómeno. Un modelo predictivo contribuye cuando, además, demuestra utilidad para anticipar resultados observables. Por ello, un modelo explicativo-predictivo debe sostener ambas dimensiones con evidencia diferenciada: primero, mediante coherencia teórica, calidad de medición y evaluación estructural; segundo, mediante comprobación del desempeño predictivo. Esta doble lectura permite que el investigador no solo afirme que el modelo tiene sentido conceptual, sino que también muestre su posible utilidad diagnóstica o aplicada (Gregor, 2006).

En el caso **CAITIZEN**, esta diferencia es especialmente importante porque el modelo se presenta como una estructura **explicativa-predictiva**. La parte explicativa permite interpretar qué dimensiones contribuyen a la ciudadanía sostenible asistida por inteligencia artificial; la parte predictiva permite valorar si los indicadores de **CAITIZEN** pueden estimarse con suficiente precisión. De esta forma, **PLSpredict** completa la ruta iniciada en los capítulos anteriores: preparación de datos, construcción del modelo, evaluación de medición, evaluación estructural y análisis de capacidad predictiva. Por ello, **PLSpredict** no debe verse como un agregado opcional, sino como una etapa necesaria cuando el modelo declara alcance predictivo (Hair et al., 2021).

En síntesis, este apartado es importante porque evita una confusión frecuente: asumir que **R^2 alto, paths significativos o f^2 relevantes equivalen automáticamente a capacidad predictiva**. La evaluación estructural muestra qué tanto el modelo explica; **PLSpredict** muestra qué tan bien el modelo anticipa indicadores endógenos. Esta diferencia permite reportar resultados con mayor rigor, especialmente cuando el estudio declara un propósito explicativo-predictivo. En el modelo **CAITIZEN**, la incorporación de PLSpredict permite cerrar la evaluación metodológica con una pregunta adicional: no solo qué constructos explican la ciudadanía sostenible asistida por inteligencia artificial, sino qué tan bien el modelo puede predecir sus indicadores observables (SmartPLS GmbH, s. f.-a).

Preparación del modelo final para PLSpredict

La evaluación con **PLSpredict** debe iniciar a partir del **modelo final depurado**, no desde una versión provisional, exploratoria o usada únicamente para ajuste global. En el caso **CAITIZEN**, el modelo adecuado es **CAITIZEN_PLS_SIN CAIL2**, porque ya incorpora la decisión metodológica de retirar el indicador **CAIL2** y conserva la **estructura reflectiva validada** en los capítulos anteriores. Esta preparación es importante porque la predicción solo tiene sentido cuando el **modelo de medición** y el **modelo estructural** han sido revisados previamente. De lo contrario, se correría el riesgo de evaluar capacidad predictiva sobre una estructura todavía inestable o metodológicamente no defendible (Hair et al., 2019).

El modelo final depurado debe conservar la misma **especificación teórica** utilizada para la evaluación estructural. Esto significa que los **constructos, indicadores y relaciones entre variables latentes** deben corresponder al modelo que se desea reportar académicamente. En **PLS-SEM**, la predicción no debe ejecutarse sobre un modelo alternativo si este no representa la **estructura conceptual final**. Por ello, antes de abrir **PLSpredict**, el investigador debe verificar que los constructos estén correctamente conectados, que los indicadores correspondan a sus variables latentes y que las relaciones estructurales reflejen la **lógica teórica del estudio**. Esta coherencia es necesaria para conectar **teoría, medición, estructura y predicción** (Chin, 1998).

En el caso **CAITIZEN**, la preparación del modelo también implica reconocer que el **constructo endógeno principal** es **CAITIZEN**, por lo que la evaluación predictiva se concentrará en sus **indicadores observables**. PLSpredict no busca simplemente repetir los valores de **R^2 o f^2** obtenidos en el capítulo anterior, sino estimar qué tan bien el modelo anticipa indicadores como **CAITIZEN1, CAITIZEN2** y los demás ítems del constructo dependiente. Esta distinción permite que el lector comprenda que la predicción se evalúa a nivel de **desempeño empírico sobre indicadores endógenos**, no solo a partir de la fuerza interna de las relaciones estructurales (Shmueli et al., 2019).

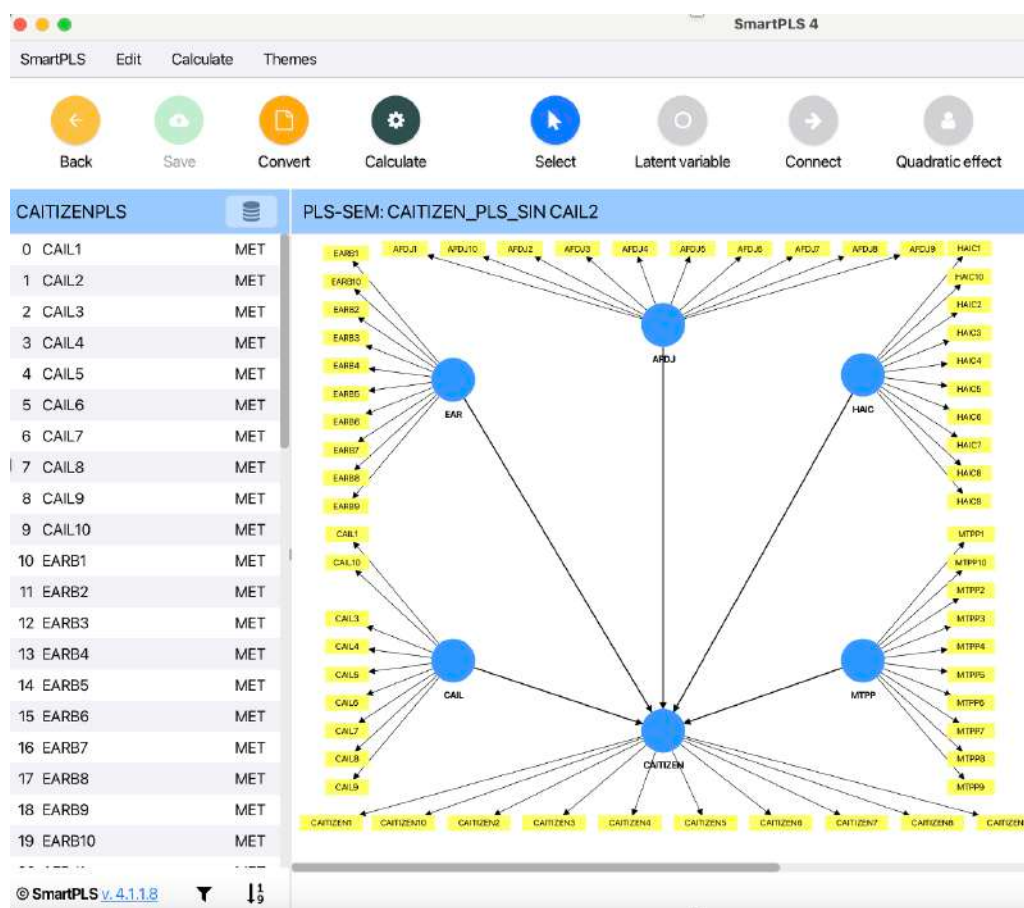
También debe evitarse ejecutar **PLSpredict** en la copia del modelo utilizada para el **análisis de ajuste global** o en configuraciones especiales como el **modelo FIT**.

Juan Mejía Tréjo

Dichas copias cumplen una función metodológica diferente y no necesariamente representan la estructura final que será sometida a predicción. El análisis predictivo debe aplicarse sobre el **modelo normal de PLS-SEM**, con sus **constructos reflectivos**, **relaciones estructurales** e **indicadores depurados**. Esta decisión permite mantener trazabilidad entre los capítulos: primero se prepara la base, después se construye el modelo, luego se evalúa la medición y estructura, y finalmente se analiza la **capacidad predictiva** (Hair et al., 2021).

La **Figura 9.1** debe colocarse al final de esta sección porque muestra visualmente el **modelo CAITIZEN final depurado** antes de ejecutar **PLSpredict**. Su función no es presentar todavía resultados predictivos, sino documentar el **punto de partida del procedimiento**. La imagen permite verificar que el análisis se realizará sobre el modelo **CAITIZEN_PLS_SIN CAIL2**, con los constructos y relaciones preparados para la evaluación predictiva. En términos didácticos, esta figura ayuda al lector a identificar qué modelo debe abrirse en **SmartPLS4** antes de seleccionar la herramienta **PLSpredict/CVPAT** en el menú de cálculo (SmartPLS GmbH, s. f.-a).

Figura 9.1. Modelo CAITIZEN final depurado para ejecutar PLSpredict.



Nota. El análisis predictivo debe ejecutarse sobre el **modelo PLS-SEM final depurado**, no sobre una copia FIT ni sobre modelos provisionales.

Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

Acceso a PLSpredict/CVPAT en SmartPLS4

Una vez abierto el **modelo final depurado CAITIZEN_PLS_SIN CAIL2**, el siguiente paso consiste en acceder al módulo de predicción de **SmartPLS4**. Para ello, se debe ubicar el menú superior **Calculate** y seleccionar la opción **PLSpredict/CVPAT**. Esta ruta operativa es importante porque diferencia el análisis predictivo de otros procedimientos disponibles en el software, como **PLS-SEM algorithm**, **Bootstrapping**, **IPMA**, **Model comparison** o **FIMIX segmentation**. En esta fase, el objetivo ya no es estimar coeficientes de trayectoria ni contrastar hipótesis, sino evaluar el **desempeño predictivo** del modelo sobre los indicadores endógenos (SmartPLS GmbH, s. f.-a).

La selección de **PLSpredict/CVPAT** debe realizarse después de haber completado la evaluación del modelo de medición y del modelo estructural. Esto significa que el investigador ya debió revisar previamente las **cargas externas**, la **confiabilidad**, la **validez convergente**, la **validez discriminante**, la **colinealidad**, los **coeficientes path**, el R^2 y el f^2 . PLSpredict no sustituye esas etapas; las complementa. Su función es responder una pregunta posterior: si el modelo ya muestra calidad de medición y capacidad explicativa, ¿también puede anticipar adecuadamente los indicadores del constructo endógeno? Esta secuencia preserva el rigor metodológico del análisis **PLS-SEM** (Hair et al., 2019).

En términos didácticos, la captura del menú **Calculate** debe colocarse en esta sección porque muestra al lector la ruta exacta para activar el análisis predictivo. La imagen evita confusiones con otras herramientas de **SmartPLS4** y permite reconocer visualmente que **PLSpredict/CVPAT** es una función específica. Esta precisión es relevante para un libro aplicado, ya que el aprendizaje no debe limitarse a explicar conceptos, sino también a mostrar cómo se ejecutan en el software. La operación correcta del programa ayuda a mantener correspondencia entre la intención metodológica y el procedimiento técnico realizado por el investigador (Hair et al., 2021).

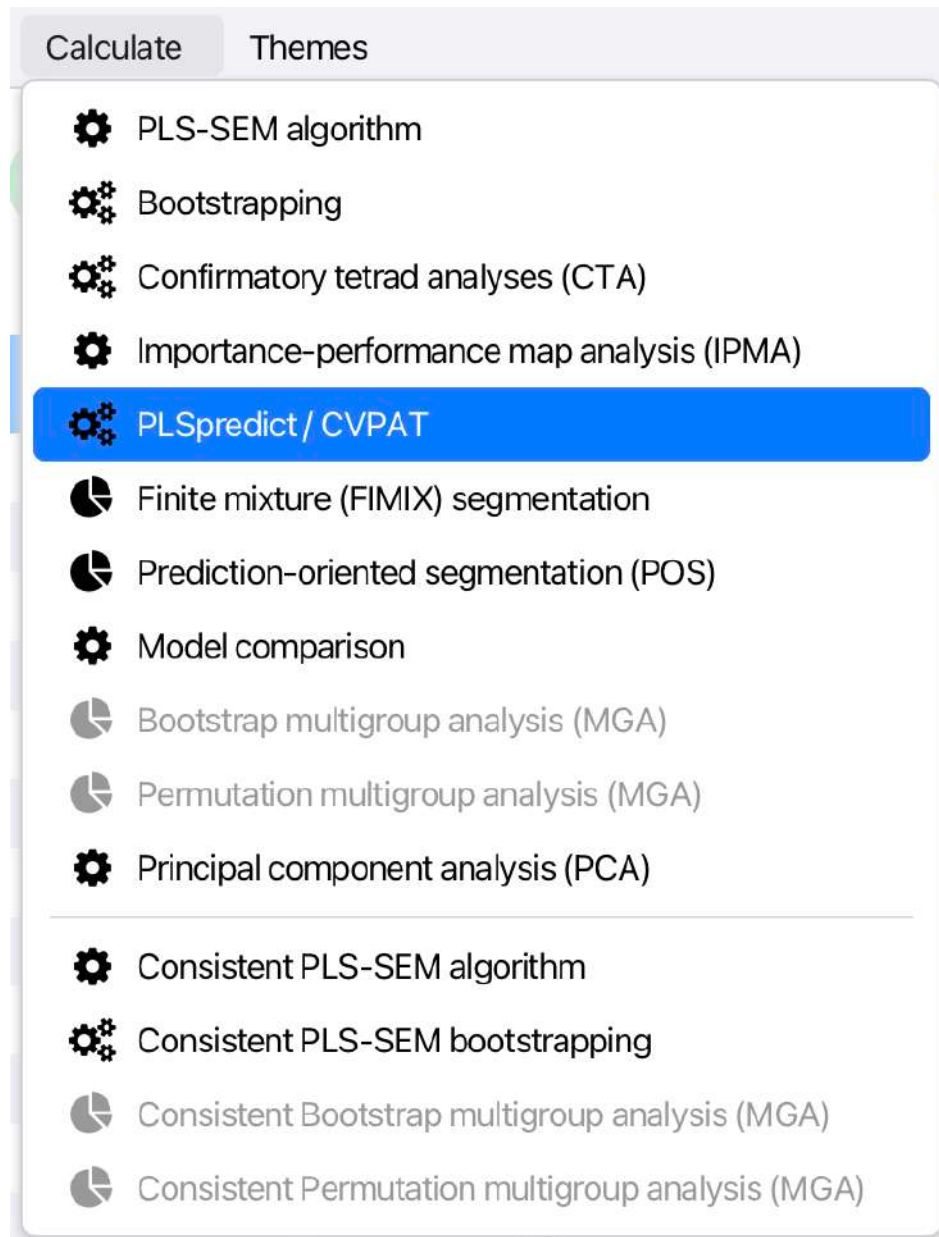
Desde el punto de vista metodológico, el acceso a **PLSpredict/CVPAT** marca el tránsito del análisis explicativo al análisis predictivo. El algoritmo **PLS-SEM** permite estimar relaciones entre constructos y maximizar la varianza explicada de las variables dependientes; en cambio, **PLSpredict** permite valorar si el modelo genera predicciones con errores aceptables. Esta distinción es central porque un modelo puede presentar relaciones significativas y, aun así, no mostrar desempeño predictivo suficiente. Por ello, seleccionar esta opción en **SmartPLS4** representa una decisión analítica orientada a comprobar la **utilidad anticipatoria** del modelo **CAITIZEN** (Shmueli et al., 2019).

La **Figura 9.2** debe insertarse al final de esta sección, inmediatamente después de explicar la ruta **Calculate** → **PLSpredict/CVPAT**. Su función es documentar visualmente el punto en que el investigador activa el procedimiento predictivo. En el pie de figura conviene señalar que esta opción se ejecuta sobre el **modelo final depurado**, no sobre modelos alternativos ni copias **FIT**. Con ello, el lector

Juan Mejía Trejo

comprenderá que **PLSpredict** forma parte de una secuencia ordenada: primero se valida el modelo, después se evalúa su estructura y finalmente se analiza su **capacidad predictiva** mediante el módulo correspondiente (Ringle et al., 2012).

Figura 9.2. Selección de PLSpredict/CVPAT desde el menú Calculate.



Nota. La opción **PLSpredict/CVPAT** debe seleccionarse desde el menú **Calculate** una vez abierto el **modelo PLS-SEM final depurado**. Este procedimiento activa la evaluación predictiva del modelo y no sustituye el algoritmo **PLS-SEM** ni el **Bootstrapping**.

Fuente: Elaboración propia con **SmartPLS4**.

Configuración predictiva: folds, repeticiones y semilla fija

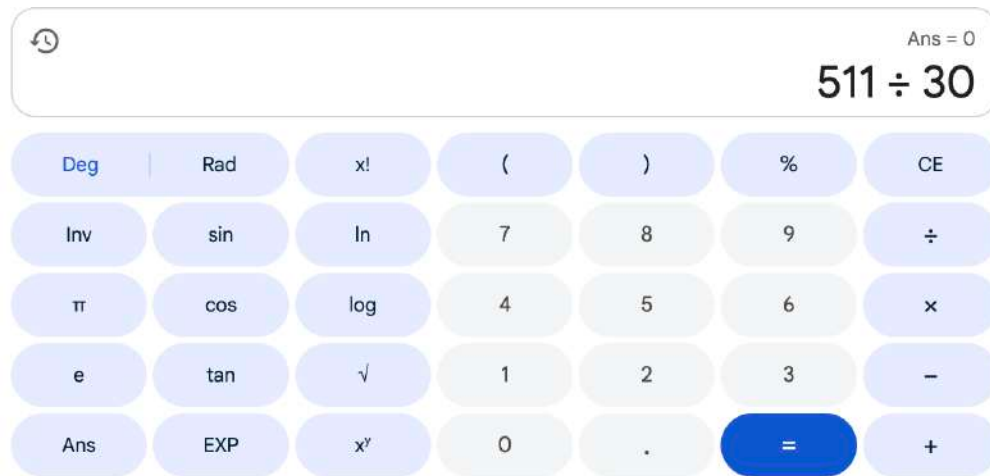
Una vez seleccionada la opción **PLSpredict/CVPAT**, **SmartPLS4** abre la ventana de configuración del análisis predictivo. En esta etapa, el investigador debe definir los parámetros que determinarán cómo se realizará la **validación cruzada predictiva**. Los elementos principales son el **número de folds**, el **número de repeticiones**, el **tipo de predicción** y el **generador aleatorio**. Estos parámetros no deben modificarse de manera arbitraria, porque influyen en la forma en que el software divide los datos para estimar y evaluar predicciones. Por ello, esta sección debe explicar no solo qué valores se seleccionan, sino también por qué se utilizan en el modelo **CAITIZEN** (Shmueli et al., 2019).

El parámetro **folds** indica en cuántos subconjuntos se dividirá la muestra para realizar la **validación cruzada predictiva**. En términos prácticos, **SmartPLS4** utiliza una parte de los datos para entrenar el modelo y otra parte para evaluar su capacidad predictiva. En el caso **CAITIZEN**, la muestra final fue de **511 casos válidos**, por lo que se propone una estimación didáctica dividiendo **511 entre 30**, lo que produce aproximadamente **17 folds** (ver Figura 9.3). El valor **30** se utiliza como tamaño mínimo práctico por subconjunto porque permite que cada sección de retención conserve un número suficiente de observaciones para calcular errores predictivos con mayor estabilidad. No debe interpretarse como una regla universal, sino como una decisión razonada para evitar subconjuntos demasiado pequeños (Hair et al., 2021).

El uso de **30 casos por sección** tiene una justificación metodológica básica: cuando los subconjuntos son demasiado reducidos, los errores predictivos pueden volverse más sensibles a casos extremos, respuestas atípicas o variaciones aleatorias de la muestra. En cambio, trabajar con secciones cercanas a **30 observaciones permite contar con una base mínima para estimar indicadores como RMSE y MAE con mayor consistencia**. Este criterio se relaciona con la idea general de mantener tamaños muestrales suficientes para análisis estadísticos aplicados, aunque no constituye un umbral exclusivo de **PLSpredict**. Por ello, se usa aquí como referencia didáctica y no como exigencia normativa (Green, 1991).

Así, la operación $511 \div 30 = 17$ permite traducir el tamaño muestral total del estudio **CAITIZEN** en una configuración operativa defendible para **PLSpredict**. Con **17 folds**, cada partición mantiene aproximadamente **30** casos en la sección de evaluación, mientras el resto de la muestra se utiliza para entrenar el modelo en cada iteración. Esta distribución conserva un equilibrio razonable entre entrenamiento y retención, evitando tanto folds excesivamente numerosos con subconjuntos pequeños como folds demasiado pocos con secciones de validación amplias. La decisión fortalece la trazabilidad del análisis predictivo y facilita que el lector comprenda por qué se ajustó la configuración predeterminada del software (Shmueli et al., 2019).

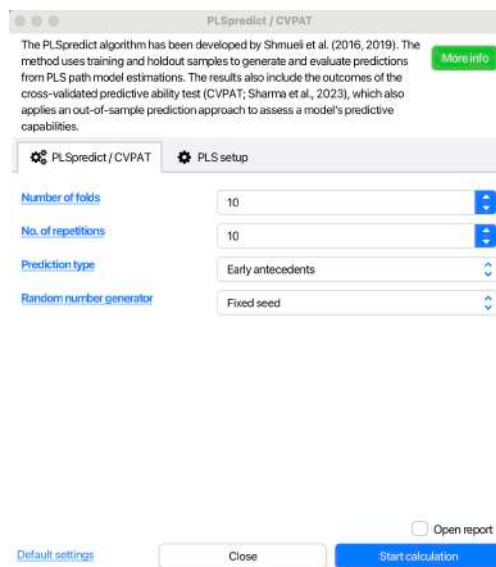
Figura 9.3. Estimación didáctica del número de folds mediante $511 \div 30 = 17$.



Fuente: Elaboración propia

La configuración inicial de **SmartPLS4** suele presentar valores predeterminados, como **10 folds** y **10 repeticiones**. Esta configuración puede ser suficiente en muchos análisis aplicados, pero el investigador puede ajustarla cuando desea documentar una decisión más alineada con el tamaño de la muestra. En el caso **CAITIZEN**, mostrar la configuración inicial tiene valor didáctico porque permite distinguir entre el valor predeterminado del software y la configuración finalmente adoptada. Por ello, la **Figura 9.4** debe insertarse después de explicar los valores iniciales, mostrando cómo aparece la ventana antes del ajuste del número de folds (SmartPLS GmbH, s. f.-d).

Figura 9.4. Configuración inicial predeterminada de PLSpredict/CVPAT con 10 folds.



Fuente: Elaboración propia con SmartPLS.

Juan Mejía Tréjo

Posteriormente, la configuración se ajusta a **17 folds** y se mantienen **10 repeticiones**. Las repeticiones permiten aumentar la estabilidad del análisis predictivo, ya que el proceso de validación cruzada se ejecuta varias veces. También se conserva el tipo de predicción **Early antecedents**, adecuado cuando se desea que el modelo utilice los antecedentes disponibles para generar predicciones de los indicadores endógenos. Asimismo, se selecciona **Fixed seed** como generador aleatorio, porque permite que el procedimiento sea reproducible. Esta decisión es relevante en investigación aplicada, donde la trazabilidad del análisis fortalece la confiabilidad del reporte metodológico (Hair et al., 2019). Ver **Figura 9.5**.

Figura 9.5. Configuración ajustada de PLSpredict con 17 folds, 10 repeticiones, Early antecedents y Fixed seed.

The screenshot displays the 'PLSpredict / CVPAT' window. At the top, a text box explains the algorithm: 'The PLSpredict algorithm has been developed by Shmueli et al. (2016, 2019). The method uses training and holdout samples to generate and evaluate predictions from PLS path model estimations. The results also include the outcomes of the cross-validated predictive ability test (CVPAT; Sharma et al., 2023), which also applies an out-of-sample prediction approach to assess a model's predictive capabilities.' A green 'More info' button is to the right. Below this, two tabs are visible: 'PLSpredict / CVPAT' (active) and 'PLS setup'. The configuration settings are as follows:

Parameter	Value
Number of folds	17
No. of repetitions	10
Prediction type	Early antecedents
Random number generator	Fixed seed

At the bottom, there is a 'Default settings' link, a 'Close' button, an 'Open report' checkbox (unchecked), and a blue 'Start calculation' button.

Fuente: Elaboración propia con **SmartPLS4**

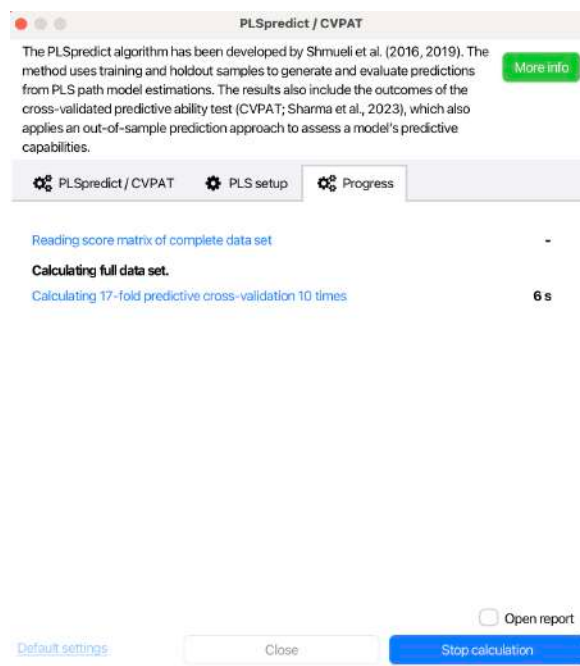
Juan Mejía Tréjo

Ejecución del cálculo y generación del reporte

Una vez definida la configuración predictiva, el siguiente paso consiste en presionar **Start calculation** para ejecutar **PLSpredict/CVPAT**. En este momento, **SmartPLS4** inicia el proceso de **validación cruzada predictiva** con los parámetros previamente establecidos: **17 folds**, **10 repeticiones**, tipo de predicción **Early antecedents** y generador aleatorio **Fixed seed**. Esta fase confirma que el análisis ya no se encuentra en la etapa de configuración, sino en la generación efectiva de predicciones. Por ello, debe entenderse como el punto en que el modelo **CAITIZEN** comienza a ser evaluado en términos de **desempeño predictivo** (Shmueli et al., 2019).

Durante la ejecución, **SmartPLS4** muestra una ventana de progreso donde se observa la lectura de datos, la construcción de matrices y el cálculo iterativo de la validación cruzada. En la captura se aprecia el mensaje **“Calculating 17-fold predictive cross-validation 10 times”**, lo cual documenta que el software está aplicando exactamente la configuración establecida en la sección anterior. Esta evidencia visual es importante porque permite comprobar que el análisis no se ejecutó con valores predeterminados, sino con los parámetros definidos para el caso **CAITIZEN**. Así, la figura cumple una función de **trazabilidad metodológica** (Hair et al., 2021). Ver **Figura 9.6**.

Figura 9.6. Ejecución de la validación cruzada predictiva en SmartPLS4.



Nota. **SmartPLS4** ejecuta **PLSpredict/CVPAT** con **17 folds** y **10 repeticiones**, conforme a la configuración definida para el modelo **CAITIZEN**.

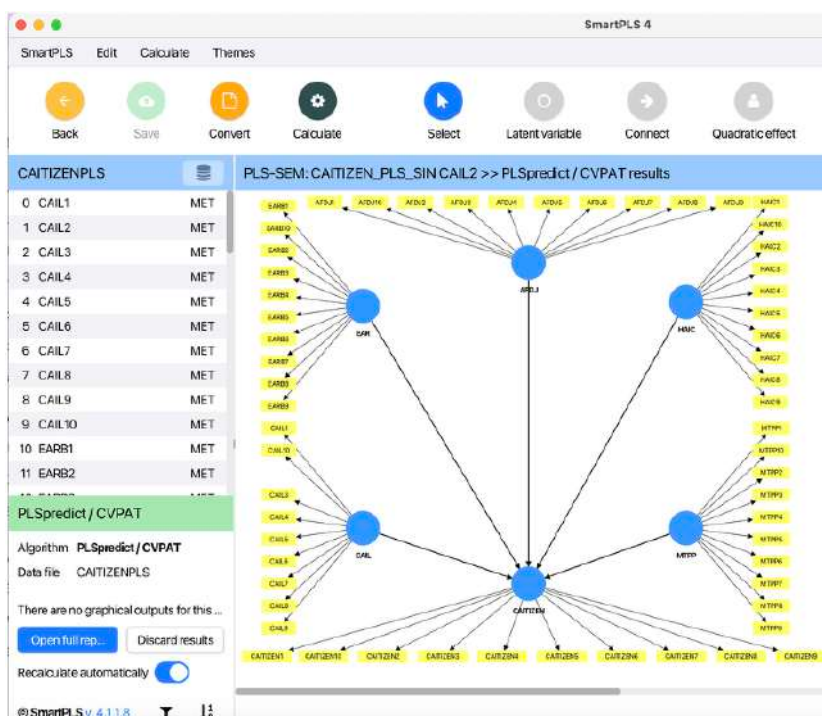
Fuente: Elaboración propia con SmartPLS.

Juan Mejía Trejo

La ejecución de **PLSpredict** debe diferenciarse claramente de otros cálculos realizados en capítulos anteriores. El **PLS-SEM algorithm** estima relaciones estructurales, el **Bootstrapping** evalúa la significancia de los coeficientes y **PLSpredict** analiza errores de predicción sobre indicadores o constructos endógenos. Por tanto, esta etapa no busca aceptar o rechazar hipótesis, sino generar evidencia sobre la capacidad del modelo para anticipar valores observables. Esta diferencia es relevante porque un modelo puede tener **R^2 aceptable** y **paths significativos**, pero aun así requerir una prueba adicional para sostener afirmaciones predictivas (Hair et al., 2019).

Cuando el cálculo termina, **SmartPLS4** habilita el acceso a los resultados mediante la opción **Open full report**. Esta pantalla debe mostrarse porque marca el tránsito entre la ejecución del análisis y la consulta de resultados. A diferencia de la evaluación estructural, donde el investigador suele interpretar coeficientes visibles en el diagrama, PLSpredict genera un reporte principalmente tabular. Por ello, el lector debe aprender a identificar dónde se abre el reporte completo y cuáles secciones deben revisarse. La opción **Open full report** permite acceder a los resúmenes de predicción, errores y comparaciones necesarias para la interpretación (SmartPLS GmbH, s. f.-d). Ver **Figura 9.7**.

Figura 9.7. Generación de resultados y acceso al reporte completo mediante *Open full report*.



Nota. Una vez finalizado el cálculo, **SmartPLS4** habilita la opción **Open full report**, desde donde se consultan los resultados predictivos del modelo.

Fuente: Elaboración propia con **SmartPLS4**.

Juan Mejía Tréjo

En síntesis, la ejecución del cálculo y la generación del reporte constituyen el puente entre la configuración técnica y la interpretación metodológica. Primero se confirma que **SmartPLS4** aplicó la validación cruzada con **17 folds** y **10 repeticiones**; después se abre el reporte completo para revisar los indicadores predictivos. Esta secuencia evita interpretar resultados sin documentar cómo fueron generados. En el caso **CAITIZEN**, la correcta ejecución de **PLSpredict** permite pasar a la siguiente sección, donde se analizarán **$Q^2_{predict}$** , **RMSE**, **MAE** y la comparación entre el modelo **PLS-SEM** y el modelo lineal de referencia **LM** (Ringle et al., 2012).

La configuración final debe reportarse con claridad para que otro investigador pueda replicar el procedimiento. En el caso **CAITIZEN**, puede redactarse que **PLSpredict/CVPAT** se ejecutó con **17 folds**, **10 repeticiones**, tipo de predicción **Early antecedents** y generador aleatorio **Fixed seed**. Esta información debe aparecer antes de interpretar los resultados, porque forma parte de la trazabilidad metodológica del análisis. En síntesis, la configuración no es un paso meramente técnico: establece las condiciones bajo las cuales se evaluará la capacidad predictiva del modelo. Por ello, debe documentarse antes de iniciar el cálculo y abrir el reporte de resultados (Ringle et al., 2012).

Criterio de evaluación predictiva

La evaluación del poder predictivo fuera de la muestra mediante **PLSpredict** permite determinar si un modelo **PLS-SEM** no solo explica adecuadamente las relaciones entre constructos, sino si también puede **anticipar valores de indicadores endógenos en observaciones no utilizadas directamente para estimar el modelo**. Esta diferencia es central: un modelo puede presentar buen **R^2** , coeficientes significativos y tamaños de efecto relevantes, pero aun así no necesariamente mostrar buen desempeño predictivo fuera de la muestra. Por ello, la evaluación predictiva debe realizarse como una etapa específica y no como una simple extensión del modelo estructural explicativo (Shmueli et al., 2019).

El primer paso consiste en revisar los valores **$Q^2_{predict}$** de todos los indicadores asociados con los constructos endógenos. Este criterio permite identificar si el modelo tiene **relevancia predictiva inicial**. Cuando **$Q^2_{predict} \leq 0$** , la predicción no se confirma, porque el modelo no supera una referencia ingenua basada en valores promedio. En cambio, cuando **$Q^2_{predict} > 0$** , puede afirmarse que el modelo posee relevancia predictiva fuera de muestra, aunque todavía no puede concluirse que tiene poder predictivo alto (Hair et al., 2021).

Una vez confirmada la positividad de **$Q^2_{predict}$** , el segundo paso consiste en revisar la distribución de los errores de predicción. La lógica propuesta por Shmueli et al. plantea una pregunta clave: **si los errores de predicción están distribuidos de manera altamente simétrica**. Esta decisión permite seleccionar el indicador de error más adecuado. Si los errores son aproximadamente simétricos, se recomienda usar

RMSE; si no son simétricos o presentan valores extremos, se recomienda utilizar **MAE**, por ser menos sensible a desviaciones extremas (Shmueli et al., 2019).

El **RMSE** o raíz del error cuadrático medio penaliza con mayor fuerza los errores grandes, por lo que resulta útil cuando se desea detectar desviaciones predictivas pronunciadas. En cambio, el **MAE** expresa el error absoluto medio y ofrece una lectura más estable cuando los errores no presentan distribución simétrica. Esta distinción es relevante porque la elección del criterio de error puede modificar la interpretación del poder predictivo del modelo. Por ello, no debe elegirse RMSE o MAE de manera automática, sino conforme al comportamiento de los errores de predicción (Hair et al., 2019).

El tercer paso corresponde a la comparación entre el modelo **PLS-SEM** y el modelo lineal de referencia **LM**. Esta comparación es indispensable porque un **$Q^2_{predict}$ positivo** solo confirma relevancia predictiva básica, pero no demuestra que el modelo **PLS-SEM** sea mejor que una alternativa lineal simple. La regla de decisión es directa: si **PLS-SEM_RMSE < LM_RMSE** o **PLS-SEM_MAE < LM_MAE**, entonces el modelo **PLS-SEM** predice mejor que el modelo lineal de referencia para ese indicador (Shmueli et al., 2019).

La clasificación del poder predictivo depende del número de indicadores en los que el modelo **PLS-SEM** presenta errores menores que el modelo LM. Si **PLS-SEM** no supera a **LM** en ningún indicador, la relevancia predictiva no queda confirmada frente al **benchmark** lineal. Si lo supera en una **minoría de indicadores**, el poder predictivo se considera **bajo**; si lo supera en la **mayoría**, se interpreta como **medio**; y si lo supera en **todos los indicadores**, puede afirmarse que el modelo presenta **alto poder predictivo** (Shmueli et al., 2019).

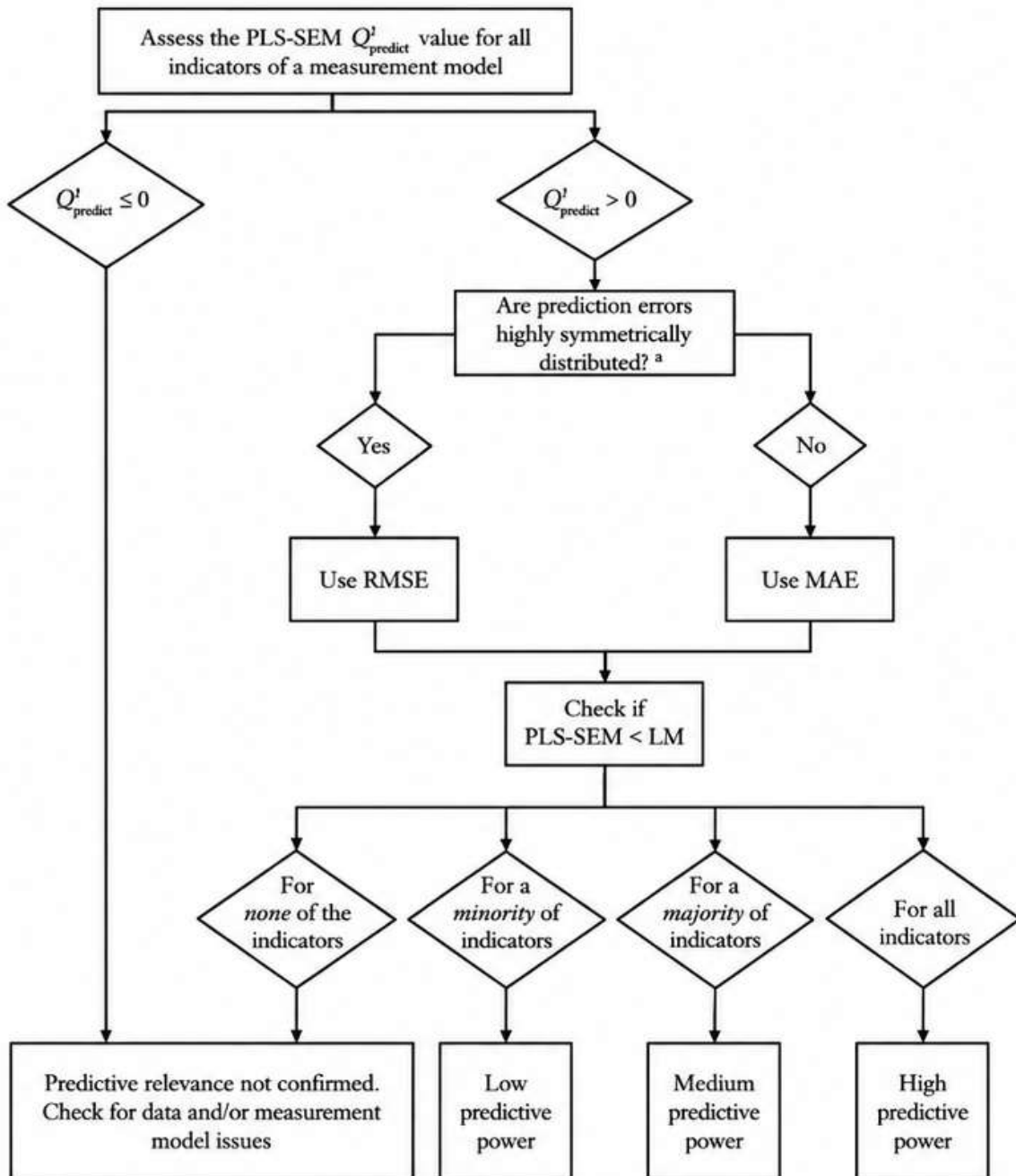
Desde una perspectiva metodológica, este criterio permite separar tres niveles de interpretación. El primero corresponde a la **relevancia predictiva**, evaluada mediante **$Q^2_{predict}$** . El segundo corresponde a la **naturaleza del error**, que orienta el uso de **RMSE** o **MAE**. El tercero corresponde a la **superioridad predictiva relativa**, evaluada mediante la comparación entre **PLS-SEM** y **LM**. Esta secuencia evita conclusiones exageradas y fortalece la transparencia del análisis predictivo (Hair et al., 2021).

En el reporte académico, la interpretación debe comunicarse como una **ruta de decisión**. Primero se indica si los indicadores endógenos tienen **$Q^2_{predict} > 0$** . Después se justifica si se utilizará **RMSE** o **MAE**, de acuerdo con la simetría de los errores. Finalmente, se reporta en cuántos indicadores el modelo **PLS-SEM** presenta menor error que **LM**. Esta última comparación permite asignar el diagnóstico final: poder predictivo no confirmado, bajo, medio o alto (Shmueli et al., 2019).

En consecuencia, la evaluación predictiva mediante **PLSpredict** no debe interpretarse únicamente a partir de **$Q^2_{predict}$** . Aunque un valor positivo confirma relevancia predictiva fuera de muestra, el poder predictivo del modelo debe establecerse comparando los errores del modelo **PLS-SEM** con los errores del modelo

lineal de referencia. Solo cuando el modelo **PLS-SEM** presenta errores menores que **LM** en una proporción suficiente de indicadores puede afirmarse que posee poder predictivo bajo, medio o alto (Shmueli et al., 2019). Ver **Figura 9.8**.

Figura 9.8. Criterio de decisión para valorar el poder predictivo fuera de muestra mediante PLSpredict



Fuente: Adaptado de Shmueli et al. (2019).

Interpretación de Q^2_{predict} , RMSE y MAE

Una vez abierto el reporte completo de **PLSpredict/CVPAT**, la interpretación debe iniciar en la sección **PLSpredict MV summary - Overview**, porque ahí se presentan los resultados para los **indicadores manifiestos** de los constructos endógenos. En el caso **CAITIZEN**, el interés principal se concentra en los indicadores **CAITIZEN1 a CAITIZEN10**, ya que representan las variables observables que el modelo intenta anticipar. Esta lectura es importante porque **PLSpredict** no se limita a evaluar el constructo latente en abstracto, sino el desempeño predictivo sobre sus indicadores. Por ello, la interpretación debe comenzar con Q^2_{predict} , antes de revisar los errores de predicción (Shmueli et al., 2019).

El criterio Q^2_{predict} permite identificar si el modelo tiene **relevancia predictiva** para cada indicador endógeno. En términos generales, un valor Q^2_{predict} **positivo** indica que el modelo genera predicciones mejores que una referencia ingenua basada en valores promedio. En las capturas del modelo **CAITIZEN**, todos los indicadores de los constructos endógenos presentan valores positivos, lo cual respalda la existencia de capacidad predictiva inicial. Esta evidencia es fundamental porque permite afirmar que el modelo no solo explica relaciones estructurales, sino que también ofrece un desempeño anticipatorio sobre los indicadores observables de la ciudadanía sostenible asistida por inteligencia artificial (Hair et al., 2021). Ver **Figura 9.9**.

Figura 9.9. Resumen *PLSpredict MV summary - Overview* para los indicadores CAITIZEN.

	Q^2_{predict}	PLS-SEM_RMSE	PLS-SEM_MAE	LM_RMSE	LM_MAE	W_RMSE	W_MAE
CAITIZEN1	0.409	1.116	0.873	1.213	0.891	1.529	1.236
CAITIZEN10	0.485	0.981	0.734	1.029	0.765	1.367	1.083
CAITIZEN2	0.402	1.081	0.781	1.109	0.816	1.359	1.110
CAITIZEN3	0.337	1.157	0.860	1.240	0.923	1.421	1.151
CAITIZEN4	0.312	1.246	0.919	1.310	0.970	1.501	1.255
CAITIZEN5	0.297	1.208	0.919	1.306	1.006	1.439	1.175
CAITIZEN6	0.247	1.376	1.022	1.474	1.088	1.586	1.304
CAITIZEN7	0.379	1.171	0.855	1.267	0.918	1.475	1.211
CAITIZEN8	0.405	1.061	0.790	1.182	0.812	1.376	1.116
CAITIZEN9	0.406	1.047	0.788	1.115	0.837	1.398	1.122

Fuente: Elaboración propia con SmartPLS.

Después de revisar Q^2_{predict} , debe continuar con los errores de predicción **RMSE** y **MAE**. El **RMSE** o raíz del error cuadrático medio penaliza con mayor fuerza los errores grandes, por lo que resulta útil cuando se desea identificar desviaciones predictivas más pronunciadas. Es posible apoyarse en la opción de copiar los resultados a Excel, con el fin de organizar los datos, revisar los errores predictivos y preparar una tabla interpretativa para el reporte académico. Ver **Figura 9.10**.

Figura 9.10. Resumen *PLSpredict MV summary* copiado a formato Excel

	A	B	C	D	E	F	G	H
1		Q²predict	PLS-SEM_RMSE	PLS-SEM_MAE	LM_RMSE	LM_MAE	IA_RMSE	IA_MAE
2	CAITIZEN1	0.409	1.176	0.873	1.213	0.891	1.529	1.236
3	CAITIZEN10	0.485	0.981	0.734	1.029	0.765	1.367	1.083
4	CAITIZEN2	0.402	1.051	0.781	1.109	0.816	1.359	1.11
5	CAITIZEN3	0.337	1.157	0.86	1.24	0.923	1.421	1.151
6	CAITIZEN4	0.312	1.245	0.919	1.31	0.97	1.501	1.255
7	CAITIZEN5	0.297	1.206	0.918	1.306	1.006	1.439	1.175
8	CAITIZEN6	0.247	1.376	1.022	1.474	1.088	1.586	1.304
9	CAITIZEN7	0.37	1.171	0.858	1.267	0.918	1.475	1.211
10	CAITIZEN8	0.405	1.061	0.79	1.102	0.812	1.376	1.116
11	CAITIZEN9	0.405	1.047	0.785	1.115	0.837	1.358	1.122

Fuente: Elaboración propia con SmartPLS.

La comparación central de *PLSpredict* se realiza entre los errores del modelo **PLS-SEM** y los errores del **modelo lineal de referencia LM**. Esta comparación es necesaria porque **Q²predict positivo** solo indica que el modelo tiene relevancia predictiva básica, pero no basta para valorar si su desempeño es superior al de un modelo alternativo simple. Por ello, se revisan los valores **PLS-SEM_RMSE**, **PLS-SEM_MAE**, **LM_RMSE** y **LM_MAE**. Si los errores del **PLS-SEM** son menores que los errores del LM, el modelo **PLS-SEM** muestra mejor desempeño predictivo para ese indicador. En consecuencia, la clasificación final del poder predictivo debe establecerse según el número de indicadores en los que **PLS-SEM** presenta errores menores que LM, como se muestra en la Tabla 9.1 (Shmueli et al., 2019). Ver **Tabla 9.1**.

Tabla 9.1. Evaluación predictiva *PLSpredict* del modelo CAITIZEN

Indicador	Q²_predict	PLS-SEM_RMSE	PLS-SEM_MAE	LM_RMSE	LM_MAE	IA_RMSE	IA_MAE	ΔRMSE (LM - PLS)	ΔMAE (LM - PLS)	Q² positivo	RMSE PLS<LM	MAE PLS<LM	Decisión interpretativa
CAITIZEN1	0.409	1.176	0.873	1.213	0.891	1.529	1.236	0.037	0.018	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN10	0.485	0.981	0.734	1.029	0.765	1.367	1.083	0.048	0.031	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN2	0.402	1.051	0.781	1.109	0.816	1.359	1.11	0.058	0.035	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN3	0.337	1.157	0.86	1.24	0.923	1.421	1.151	0.083	0.063	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN4	0.312	1.245	0.919	1.31	0.97	1.501	1.255	0.085	0.051	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN5	0.297	1.206	0.918	1.306	1.006	1.439	1.175	0.100	0.088	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN6	0.247	1.376	1.022	1.474	1.088	1.586	1.304	0.098	0.066	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN7	0.37	1.171	0.858	1.267	0.918	1.475	1.211	0.096	0.060	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN8	0.405	1.061	0.79	1.102	0.812	1.376	1.116	0.041	0.022	Si	Si	Si	Relevancia predictiva y desempeño favorable frente
CAITIZEN9	0.405	1.047	0.785	1.115	0.837	1.358	1.122	0.068	0.052	Si	Si	Si	Relevancia predictiva y desempeño favorable frente

Nota. **Q²predict** positivo indica relevancia predictiva fuera de muestra. **ΔRMSE** y **ΔMAE** se calcularon como **LM - PLS-SEM**. Por tanto, valores positivos indican que el modelo **PLS-SEM** produjo errores

menores que el modelo lineal de referencia **LM**; valores negativos indican que **LM** presentó menor error predictivo. La clasificación del poder predictivo se establece según el número de indicadores en los que **PLS-SEM** supera a **LM**.

Fuente: Elaboración propia con Excel de los datos de SmartPLS.

En síntesis, la interpretación de **PLSpredict** debe seguir una secuencia clara: primero se verifica que **$Q^2_{predict}$** sea positivo para los indicadores endógenos; después se revisan **RMSE** y **MAE**; finalmente, se compara el desempeño del modelo **PLS-SEM** contra el modelo lineal de referencia **LM**. Esta comparación permite determinar si el modelo presenta poder predictivo no confirmado, bajo, medio o alto. En el modelo evaluado, la existencia de valores **$Q^2_{predict} > 0$** confirma relevancia predictiva fuera de muestra; no obstante, la clasificación final del poder predictivo depende del número de indicadores en los que **PLS-SEM** presenta errores menores que **LM**. Por tanto, la sección transforma la salida del software en una decisión metodológica fundamentada (Hair et al., 2019; Shmueli et al., 2019).

Síntesis del procedimiento

El capítulo inicia estableciendo una distinción metodológica central: **explicar no equivale necesariamente a predecir**. Después de haber evaluado el modelo de medición, el ajuste global, el **Bootstrapping** y el modelo estructural, **PLSpredict/CVPAT** se incorpora como una etapa adicional para valorar si el modelo **CAITIZEN** puede anticipar indicadores endógenos fuera de muestra. Esta decisión evita asumir que un **R^2** aceptable, coeficientes significativos o efectos f^2 relevantes bastan para sostener desempeño predictivo. Así, el capítulo transita de la explicación estructural hacia una evaluación predictiva específica (Shmueli, 2010).

El procedimiento continúa con la preparación del **modelo final depurado**, identificado como **CAITIZEN_PLS_SIN_CAIL2**. Esta versión debe usarse porque conserva la especificación teórica final, los constructos reflectivos, los indicadores depurados y las relaciones estructurales previamente validadas. El análisis predictivo no debe ejecutarse sobre modelos provisionales ni sobre la **copia FIT**, ya que esta última cumple una función diferente para el ajuste global. Con ello se mantiene la trazabilidad entre teoría, medición, estructura y predicción, evitando evaluar capacidad predictiva sobre una versión no definitiva del modelo (Chin, 1998).

Una vez abierto el modelo final, el capítulo orienta al lector hacia el menú **Calculate** → **PLSpredict/CVPAT** en **SmartPLS4**. Esta acción marca el paso operativo del análisis explicativo al análisis predictivo. Mientras el algoritmo **PLS-SEM** estima relaciones estructurales y el **Bootstrapping** evalúa estabilidad inferencial, **PLSpredict** se concentra en los errores de predicción de los indicadores endógenos. Por ello, esta herramienta no sustituye las fases anteriores, sino que las complementa con una pregunta distinta: si el modelo puede anticipar adecuadamente los valores observables del constructo dependiente (Hair et al., 2019).

La configuración predictiva constituye una parte esencial del procedimiento. Para el modelo **CAITIZEN** se propone utilizar **17 folds**, **10 repeticiones**, tipo de predicción **Early antecedents** y generador aleatorio **Fixed seed**. Esta decisión se justifica a partir de la muestra de **511** casos válidos, dividiendo aproximadamente **511 entre 30** para obtener particiones de validación suficientemente estables. El uso de **30 observaciones** por sección se presenta como criterio didáctico y práctico, no como regla universal. **Con ello se busca reducir la sensibilidad a particiones pequeñas y fortalecer la estabilidad del cálculo** (Green, 1991).

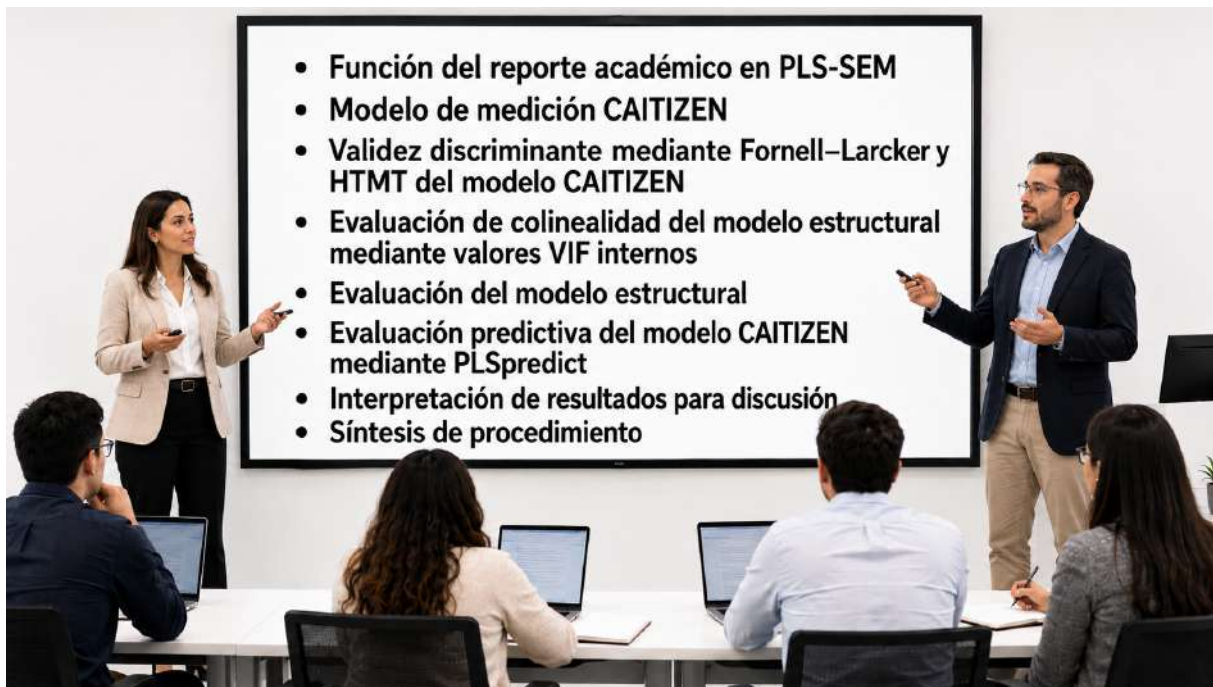
Después de configurar los parámetros, se ejecuta el cálculo mediante **Start calculation**. En esta fase **SmartPLS4** aplica la validación cruzada predictiva conforme a los valores definidos, es decir, **17 folds y 10 repeticiones**. La ejecución debe documentarse porque demuestra que el análisis no se realizó con parámetros predeterminados sin justificación, sino con una configuración metodológicamente razonada para el caso **CAITIZEN**. Al finalizar, se accede al reporte completo mediante **Open full report**, desde donde se revisan los resultados predictivos generados por el software (SmartPLS GmbH, s. f.-d).

La interpretación del reporte inicia con **$Q^2_{predict}$** , porque este indicador permite identificar si existe relevancia predictiva fuera de muestra. Si **$Q^2_{predict}$** es positivo, el modelo predice mejor que una referencia ingenua basada en promedios; si es igual o menor que cero, la predicción no se confirma. En el caso **CAITIZEN**, esta lectura se concentra en los indicadores endógenos, especialmente aquellos asociados al constructo dependiente. Esta fase evita sostener afirmaciones predictivas únicamente con base en R^2 , coeficientes path o significancia estadística del modelo estructural (Hair et al., 2021).

Posteriormente, el capítulo indica revisar los errores de predicción **RMSE** y **MAE**. El **RMSE** permite detectar desviaciones predictivas más pronunciadas porque penaliza con mayor fuerza los errores grandes, mientras que el **MAE** ofrece una lectura más estable cuando los errores no son simétricos o existen valores extremos. Esta distinción permite elegir el criterio de error más adecuado según la distribución de los residuos predictivos. De esta forma, el procedimiento no se limita a confirmar **$Q^2_{predict}$** positivo, sino que analiza la magnitud del error en los indicadores endógenos (Shmueli et al., 2019).

Finalmente, el capítulo cierra con la comparación entre los errores del modelo **PLS-SEM** y los del modelo lineal de referencia **LM**. Esta comparación permite clasificar el poder predictivo como no confirmado, bajo, medio o alto, según el número de indicadores en los que **PLS-SEM** presenta errores menores que LM. Para facilitar la lectura, se recomienda calcular **$\Delta RMSE = LM - PLS-SEM$** y **$\Delta MAE = LM - PLS-SEM$** . Valores positivos indican mejor desempeño del **PLS-SEM**. Así, el capítulo convierte la salida técnica de **SmartPLS4** en una decisión metodológica clara y defendible (Ringle et al., 2012).

CAPÍTULO 10. Integración del modelo PLS-SEM



El **Capítulo 10. Integración del modelo PLS-SEM**, tiene como propósito mostrar cómo los resultados generados en **SmartPLS4** pueden transformarse en un **argumento académico, metodológico y aplicado**. Reportar un modelo **PLS-SEM** no consiste únicamente en copiar tablas, figuras o valores estadísticos; implica explicar **qué se evaluó, con qué procedimiento, qué decisiones metodológicas se tomaron y cómo los hallazgos sostienen o limitan el modelo propuesto**. En el caso **CAITIZEN**, el reporte debe integrar la evidencia sobre **alfabetización crítica en IA, responsabilidad ética, justicia de datos, colaboración humano-IA, transparencia metacognitiva y ciudadanía sostenible asistida por IA** dentro de una narrativa científica coherente (Hair et al., 2019).

La primera función del reporte académico es demostrar que el análisis siguió una **secuencia metodológica correcta**. En **PLS-SEM**, la interpretación del modelo estructural solo es defendible cuando previamente se ha evaluado la calidad del modelo de medición. Por ello, el capítulo organiza el reporte desde la **confiabilidad interna, las cargas externas, la validez convergente y la validez discriminante**, para después avanzar hacia **colinealidad, coeficientes estructurales, significancia estadística, R^2 , f^2 y capacidad predictiva**. Esta secuencia evita aceptar hipótesis sobre constructos que no han demostrado consistencia empírica suficiente (Hair et al., 2021)

Un aporte central del capítulo es distinguir entre **resultado, decisión e interpretación**. Un valor numérico no constituye por sí mismo una conclusión; debe vincularse con un **criterio metodológico** y con una **lectura conceptual**. Por ejemplo, una carga externa adecuada respalda la calidad de un indicador, un **AVE superior a 0.50** apoya la validez convergente y un **HTMT aceptable** confirma la distinción entre constructos. En **CAITIZEN**, la exclusión de **CAIL2** debe reportarse como una decisión metodológica orientada a mejorar el AVE del constructo CAIL, sin perder trazabilidad sobre el instrumento original (Fornell & Larcker, 1981).

El capítulo también enfatiza que la **significancia estadística no debe sustituir la relevancia teórica**. En el modelo estructural, una trayectoria puede resultar significativa y, al mismo tiempo, mostrar un tamaño de efecto pequeño; o bien, puede no ser significativa y conservar valor conceptual dentro del modelo. Por ello, la interpretación debe integrar **coeficientes β , t-values, p-values, intervalos de confianza y f^2** . En el caso **CAITIZEN**, el reporte no debe limitarse a aceptar o rechazar hipótesis, sino explicar qué significa que **HAIC, AFDJ y MTPP** sean predictores directos relevantes, mientras **CAIL y EAR** mantengan valor formativo o conceptual (Cohen, 1988).

Otra función del reporte académico es conectar **explicación y predicción**. El R^2 permite conocer qué proporción de la varianza de **CAITIZEN** es explicada por sus predictores dentro de la muestra; sin embargo, la capacidad predictiva requiere una evaluación adicional mediante **PLSpredict**. Por ello, el capítulo incorpora **$Q^2_{predict}$, RMSE** y la comparación entre el modelo **PLS-SEM** y el modelo lineal de referencia. Esta distinción es relevante porque un modelo puede explicar adecuadamente un fenómeno y, aun así, no mostrar desempeño predictivo suficiente fuera de muestra (Shmueli et al., 2019).

La validez discriminante ocupa un lugar especial en el reporte porque permite demostrar que los constructos del modelo son **empíricamente distinguibles**, aunque estén conceptualmente relacionados. En **CAITIZEN**, dimensiones como **responsabilidad ética, justicia de datos y alfabetización crítica** pertenecen a un mismo campo formativo, pero no deben interpretarse como equivalentes. Por ello, el capítulo integra los criterios **Fornell–Larcker y HTMT** en una matriz conjunta, mostrando la **raíz cuadrada del AVE, las correlaciones entre constructos y las razones HTMT**. Esta integración fortalece la defensa metodológica del modelo ante constructos próximos (Henseler et al., 2015).

Finalmente, el capítulo cumple una función de **trazabilidad científica y editorial**. Cada tabla debe indicar **qué modelo se utilizó, qué salida de SmartPLS4 produjo los datos y qué archivo exportado sirvió para construir el reporte final**. Así, el modelo **CAITIZEN_PLS_sin CAIL2** se reserva para medición, estructura y predicción; mientras que la copia **FIT** se utiliza solo para ajuste global del modelo saturado. En síntesis, este capítulo enseña a convertir salidas estadísticas en un **reporte académico defendible, transparente y reproducible**, donde la evidencia empírica se articula con teoría, método e interpretación aplicada (Ringle et al., 2012).

Juan Mejía Trejo

Función del reporte académico en PLS-SEM

El **reporte académico en PLS-SEM** cumple la función de transformar los **resultados estadísticos** en una explicación **metodológica, teórica y aplicada** del modelo analizado. No basta con presentar salidas de **SmartPLS4**; el investigador debe construir una narrativa que muestre cómo los **datos**, los **indicadores**, los **constructos** y las **relaciones estructurales** sostienen o limitan el modelo propuesto. En el caso **CAITIZEN**, el reporte debe explicar cómo la **alfabetización crítica en inteligencia artificial**, la **responsabilidad ética**, la **justicia de datos**, la **colaboración humano-IA** y la **transparencia metacognitiva** se integran en una estructura explicativa y predictiva de la **ciudadanía sostenible asistida por IA**. Por tanto, reportar significa **interpretar evidencia**, no solo reproducir tablas (Hair et al., 2019).

La primera función del reporte académico es demostrar que el modelo fue evaluado en una **secuencia metodológica correcta**. En **PLS-SEM**, la interpretación de las **rutas estructurales** solo es defendible si previamente se verificó la calidad del **modelo de medición**. Por ello, el reporte debe iniciar con **confiabilidad interna**, **cargas externas**, **validez convergente** y **validez discriminante**. Después debe avanzar hacia **colinealidad**, **coeficientes de trayectoria**, **significancia estadística**, R^2 , f^2 y **capacidad predictiva**. Esta secuencia evita aceptar hipótesis sobre constructos mal medidos o empíricamente indistintos. En el caso **CAITIZEN**, esta lógica permite mostrar que los constructos reflectivos fueron evaluados antes de interpretar los efectos estructurales entre **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** (Hair et al., 2021).

La segunda función del reporte es distinguir entre **resultado**, **decisión e interpretación**. Un valor numérico no habla por sí mismo; necesita ser ubicado dentro de un **criterio metodológico** y dentro de una **lectura sustantiva**. Por ejemplo, una **carga externa aceptable** respalda la calidad de un indicador; un **AVE superior a 0.50** apoya la **validez convergente**; un **HTMT bajo** confirma distinción entre constructos; un **VIF aceptable** permite interpretar rutas sin colinealidad crítica; y un **coeficiente significativo** indica apoyo empírico para una hipótesis. Sin embargo, cada decisión debe explicarse con prudencia. En **CAITIZEN**, eliminar un indicador como **CAIL2** no es una acción mecánica, sino una decisión orientada a mejorar la calidad del constructo sin afectar su significado conceptual (Fornell & Larcker, 1981).

La tercera función del reporte académico es evitar que la **significancia estadística** sustituya la **relevancia teórica**. En **PLS-SEM**, una ruta puede ser significativa y tener un efecto pequeño, o puede no ser significativa y aun así ofrecer información teórica importante. Por ello, el reporte debe integrar β , **t-value**, **p-value**, **intervalos de confianza** y f^2 , pero también debe explicar qué implica cada relación para el fenómeno estudiado. En el modelo **CAITIZEN**, no basta con afirmar que una hipótesis fue aceptada o rechazada; debe discutirse qué significa que **CAIL funcione como antecedente fundacional** y qué implica que algunas capacidades proximales expliquen mejor la **ciudadanía sostenible asistida por IA** que otras (Cohen, 1992).

La cuarta función del reporte es conectar **explicación y predicción**. El R^2 informa la capacidad explicativa del modelo dentro de la muestra, mientras que **PLSpredict** permite evaluar su capacidad predictiva fuera de muestra. Esta diferencia es importante porque un modelo puede explicar adecuadamente la varianza de un constructo endógeno y, aun así, requerir evidencia adicional sobre su desempeño predictivo. En el caso **CAITIZEN**, el reporte debe mostrar no solo qué constructos explican la **ciudadanía sostenible asistida por IA**, sino también si el modelo puede anticipar los indicadores del constructo dependiente con errores menores que un modelo de referencia. Por ello, $Q^2_{predict}$, **RMSE** y **MAE** deben interpretarse como evidencia complementaria de utilidad predictiva (Shmueli et al., 2019).

Finalmente, el reporte académico cumple una función de **trazabilidad científica**. Cada **tabla**, **figura** y **párrafo** debe permitir que otro investigador comprenda qué modelo fue estimado, qué versión de **SmartPLS4** se utilizó, qué indicadores fueron conservados o eliminados, qué criterios se aplicaron y cómo se llegó a las conclusiones. En un libro aplicado, esta función es central porque enseña al lector a pasar del uso operativo del software a la **escritura científica**. En el caso **CAITIZEN**, el reporte debe mostrar con claridad la transición desde la **medición reflectiva** hasta la **validación estructural y predictiva** del modelo. Así, el capítulo no enseña solo a reportar resultados, sino a convertir evidencia estadística en **argumento académico defendible** (Ringle et al., 2012).

Modelo de medición CAITIZEN

La construcción de la tabla final del modelo de medición se realizó a partir del modelo depurado **CAITIZEN_PLS_sin CAIL2**. Este modelo corresponde a la versión **PLS-SEM** en la que el indicador **CAIL2** fue eliminado previamente con el propósito de fortalecer la calidad psicométrica del constructo **CAIL**, particularmente su validez convergente, evaluada mediante el **Average Variance Extracted —AVE—**. En consecuencia, todos los valores reportados en la tabla final corresponden al modelo ajustado y no al modelo inicial que incluía la totalidad de los indicadores.

Los archivos de **Excel** utilizados para construir la tabla proceden directamente de la función de exportación de resultados de **SmartPLS**. Una vez ejecutados los procedimientos de cálculo, el software permite descargar las salidas en formato Excel, lo que facilita su revisión, depuración, organización e integración en tablas académicas. En este caso, los archivos **CONFIABILIDAD.xlsx** y **BS_CARGAS EXTERNAS.xlsx** fueron exportados desde **SmartPLS** después de estimar el modelo **CAITIZEN_PLS_sin CAIL2** en los modos de cálculo correspondientes.

En primer lugar, el modelo **CAITIZEN_PLS_sin CAIL2** se estimó mediante el modo **PLS Algorithm** en **SmartPLS**. A partir de esta corrida se obtuvo la salida **Construct Reliability and Validity**, exportada posteriormente como **CONFIABILIDAD.xlsx**. De este archivo se tomaron los valores de **Cronbach Alpha**, **rho_A**, **rho_C** y **AVE** para cada constructo del modelo: **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN**. En la tabla final, estos valores se incorporaron en el encabezado descriptivo de cada variable

Juan Mejía Trejo

latente, debido a que representan la evaluación global del constructo y no el desempeño individual de cada ítem.

Posteriormente, sobre el mismo modelo depurado **CAITIZEN_PLS_sin CAIL2**, se ejecutó el procedimiento de **Bootstrapping** en **SmartPLS4**. De esta corrida se obtuvo la salida **Outer Loadings**, exportada como **BS_CARGAS EXTERNAS.xlsx**. De dicho archivo se tomaron los valores de **Original Sample**, **T Statistics** y **P Values**. El valor de **Original Sample** se reportó como **outer loading**; el valor de **T Statistics** se colocó entre paréntesis como **t-value**; y el valor de **p-values** se incorporó en la columna **p-value**. Con ello, cada indicador retenido quedó documentado con su carga externa y con su significancia estadística.

Debe precisarse que **CAIL2** no se reporta como indicador retenido, porque no forma parte del modelo final **CAITIZEN_PLS_sin CAIL2**. No obstante, puede conservarse en la tabla mediante una nota metodológica que mantenga la trazabilidad del proceso de depuración. La redacción sugerida es: **“Indicator removed to improve the construct’s AVE”**. Esta nota permite reconocer que el ítem formó parte del instrumento original, pero fue excluido del modelo estimado final por razones metodológicas asociadas con la mejora de la validez convergente del constructo **CAIL**.

Con base en el modelo depurado **CAITIZEN_PLS_sin CAIL2**, se construyó la **Tabla 10.1**, integrando los resultados exportados desde **SmartPLS** en formato **Excel**. Los valores de **Cronbach Alpha**, **rho_A**, **rho_C** y **AVE** fueron obtenidos mediante **PLS Algorithm** desde la salida **Construct Reliability and Validity**, exportada como **CONFIABILIDAD.xlsx**. Por su parte, las cargas externas, los valores **t** y los valores **p** fueron obtenidos mediante **Bootstrapping** desde la salida **Outer Loadings**, exportada como **BS_CARGAS EXTERNAS.xlsx**. Esta integración permite verificar la consistencia interna, la validez convergente y la significancia estadística de los indicadores retenidos antes de avanzar a la interpretación del modelo **CAITIZEN**.

Tabla 10.1. Evaluación del modelo de medición CAITIZEN_PLS_sin CAIL2: consistencia interna, cargas externas y validez convergente

Factor 1	Código	CAIL. Critical Artificial Intelligence Literacy. Students' capacity to understand, evaluate, verify, and responsibly use AI systems in academic, professional, and civic contexts Cronbach's Alpha (≥ 0.70) = 0.884; rho_A (Dijkstra–Henseler's ≥ 0.70) = 0.888; rho_C (Composite Reliability ≥ 0.70) = 0.907; AVE (≥ 0.50) = 0.521.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)
Critical Artificial Intelligence Literacy	CAIL1	I understand how artificial intelligence systems transform data into patterns or predictions.	0.662 (19.219)	0.000	Long & Magerko (2020)
Critical Artificial Intelligence Literacy	CAIL2	I distinguish among the main types of artificial intelligence, such as predictive, generative, and autonomous AI, and their applications.	Indicator removed to improve the construct's AVE		Ng et al. (2021)

Juan Mejía Trejo

CAPÍTULO 10. Integración del modelo PLS-SEM

Critical Artificial Intelligence Literacy	CAIL3	I understand how the clarity and precision of a prompt influence AI-generated responses.	0.593 (15.583)	0.000	Waalder et al. (2025)
Critical Artificial Intelligence Literacy	CAIL4	I identify the technical limitations and margins of error of the AI models I use.	0.757 (31.206)	0.000	Southworth et al. (2023)
Critical Artificial Intelligence Literacy	CAIL5	I detect when an AI-generated output contains bias or unsupported inferences.	0.761 (32.789)	0.000	Wang & Wang (2025)
Critical Artificial Intelligence Literacy	CAIL6	I verify AI-generated information by comparing it with reliable academic or human sources.	0.754 (29.681)	0.000	Córdova-Esparza (2025)
Critical Artificial Intelligence Literacy	CAIL7	I recognize when a decision should be made by a person rather than by an algorithm.	0.742 (24.226)	0.000	Gunasekara et al. (2025)
Critical Artificial Intelligence Literacy	CAIL8	I analyze the ethical implications of using AI in teaching, research, and institutional management.	0.802 (40.435)	0.000	Miao & Cukurova (2024)
Critical Artificial Intelligence Literacy	CAIL9	I identify when an AI application may violate privacy, equity, or digital rights.	0.745 (30.415)	0.000	González-Argote et al. (2025)
Critical Artificial Intelligence Literacy	CAIL10	I promote the responsible, transparent, and traceable use of AI in my professional or academic environment.	0.706 (24.044)	0.000	Gunasekara et al. (2025)
Factor 2	Code	EAR. Ethical Awareness and Responsibility. Students' capacity to recognize ethical risks, assume responsibility, preserve academic integrity, protect privacy, and promote transparent and socially responsible AI use. Cronbach's Alpha (≥ 0.70) = 0.911; rho_A (Dijkstra-Henseler's ≥ 0.70) = 0.917; rho_C (Composite Reliability ≥ 0.70) = 0.925; AVE (≥ 0.50) = 0.554.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)
Ethical Awareness and Responsibility	EAR1	I explicitly declare when and how I use artificial intelligence in my academic work to maintain transparency and personal responsibility.	0.639 (19.185)	0.000	Kong & Zhu (2025)
Ethical Awareness and Responsibility	EAR2	Before using AI, I analyze possible ethical risks, such as bias, privacy, authorship, or misinformation, and how to prevent them.	0.713 (28.189)	0.000	Gunasekara et al. (2025)
Ethical Awareness and Responsibility	EAR3	I respect privacy and data protection when using AI tools by avoiding the disclosure of sensitive information.	0.724 (26.643)	0.000	González-Argote et al. (2025)
Ethical Awareness and Responsibility	EAR4	I avoid applying AI in ways that may cause harm, exclusion, or misinformation in my academic environment.	0.768 (31.596)	0.000	Córdova-Esparza (2025)
Ethical Awareness and Responsibility	EAR5	I encourage my peers to use AI ethically, transparently, and responsibly as a practice of academic integrity.	0.754 (32.491)	0.000	Kong & Zhu (2025)
Ethical Awareness and Responsibility	EAR6	I take responsibility for AI-generated outputs in my academic activities	0.775 (33.451)	0.000	Papagiannidis et al. (2025)

Juan Mejía Trejo

CAPÍTULO 10. Integración del modelo PLS-SEM

		and correct errors derived from their use.			
Ethical Awareness and Responsibility	EAR7	I evaluate the social effects of the AI I use, considering equity, inclusion, and sustainability.	0.762 (31.713)	0.000	Córdova-Esparza (2025)
Ethical Awareness and Responsibility	EAR8	I recognize that academic products must preserve human authorship, using AI only as technical or creative support.	0.755 (29.914)	0.000	Wang & Wang (2025)
Ethical Awareness and Responsibility	EAR9	I reject using AI to plagiarize, falsify, or impersonate my own or others' academic work.	0.751 (26.923)	0.000	Kong & Zhu (2025)
Ethical Awareness and Responsibility	EAR10	I promote transparency and the common good in the use of AI, ensuring that its benefits are shared equitably.	0.794 (42.896)	0.000	Papagiannidis et al. (2025)
Factor 3	Code	AFDJ. Awareness of Fairness and Data Justice. Students' capacity to identify algorithmic bias, recognize data-related inequalities, value transparency, and support fairness, accountability, and justice in AI systems. Cronbach's Alpha (≥ 0.70) = 0.919; rho_A (Dijkstra-Henseler's ≥ 0.70) = 0.922; rho_C (Composite Reliability ≥ 0.70) = 0.932; AVE (≥ 0.50) = 0.581.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)
Awareness of Fairness and Data Justice	AFDJ1	I recognize that limited or homogeneous datasets may generate unfair algorithmic decisions.	0.678 (22.066)	0.000	González-Argote et al. (2025)
Awareness of Fairness and Data Justice	AFDJ2	I consider external auditing of algorithms necessary to detect and correct bias.	0.747 (27.486)	0.000	Decker et al. (2025)
Awareness of Fairness and Data Justice	AFDJ3	I value AI systems making public the criteria they use to make decisions.	0.771 (33.435)	0.000	Demirchyan (2025)
Awareness of Fairness and Data Justice	AFDJ4	I believe that including cultural and gender diversity in development teams improves algorithmic fairness.	0.720 (26.639)	0.000	Pham et al. (2025)
Awareness of Fairness and Data Justice	AFDJ5	I identify the importance of continuously supervising AI models to prevent discrimination.	0.745 (30.743)	0.000	Demirchyan (2025)
Awareness of Fairness and Data Justice	AFDJ6	I consider it essential for educational institutions to establish policies to evaluate algorithmic justice.	0.785 (32.474)	0.000	Decker et al. (2025)
Awareness of Fairness and Data Justice	AFDJ7	I observe that imbalance in data representation can affect a system's accuracy and impartiality.	0.812 (42.974)	0.000	Pham et al. (2025)
Awareness of Fairness and Data Justice	AFDJ8	I defend people's right to know the variables that influence an algorithm's decisions.	0.768 (27.375)	0.000	Decker et al. (2025)
Awareness of Fairness and Data Justice	AFDJ9	I consider transparency in the origin and processing of data to be key for equity in AI.	0.803 (40.303)	0.000	González-Argote et al. (2025)

Juan Mejía Trejo

CAPÍTULO 10. Integración del modelo PLS-SEM

Awareness of Fairness and Data Justice	AFDJ10	I support the publication of institutional reports on detected biases and improvements implemented in educational algorithms.	0.781 (36.238)	0.000	Pham et al. (2025)
Factor 4	Code	HAIC. Human–AI Creative Collaboration. Students' capacity to use AI as a creative partner for ideation, problem-solving, experimentation, reflective improvement, and innovation-oriented knowledge production. Cronbach's Alpha (≥ 0.70) = 0.942; rho_A (Dijkstra–Henseler's ≥ 0.70) = 0.943; rho_C (Composite Reliability ≥ 0.70) = 0.950; AVE (≥ 0.50) = 0.656.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)
Human–AI Creative Collaboration	HAIC1	I use AI as a starting point to generate ideas that I later develop personally.	0.737 (27.509)	0.000	Georgieva & Georgiev (2025)
Human–AI Creative Collaboration	HAIC2	I combine my own contributions with AI suggestions to create more innovative outcomes.	0.809 (38.818)	0.000	Rafner et al. (2025)
Human–AI Creative Collaboration	HAIC3	I analyze how AI responses influence my creative process.	0.793 (33.949)	0.000	Wang et al. (2025)
Human–AI Creative Collaboration	HAIC4	I learn new ways of solving problems when interacting with AI.	0.827 (41.902)	0.000	Wang et al. (2025)
Human–AI Creative Collaboration	HAIC5	I consider AI a creative partner that complements my human abilities.	0.802 (36.362)	0.000	Salma et al. (2025)
Human–AI Creative Collaboration	HAIC6	Collaboration with AI inspires new perspectives or approaches that I would not have had alone.	0.787 (32.799)	0.000	Rafner et al. (2025)
Human–AI Creative Collaboration	HAIC7	I use AI to experiment with different creative styles or solutions.	0.842 (48.799)	0.000	Georgieva & Georgiev (2025)
Human–AI Creative Collaboration	HAIC8	I reflect on how my work improves through interaction with AI.	0.820 (41.686)	0.000	Wang et al. (2025)
Human–AI Creative Collaboration	HAIC9	I adjust AI proposals to align them with my own creative vision or purpose.	0.836 (50.158)	0.000	Salma et al. (2025)
Human–AI Creative Collaboration	HAIC10	In creative projects, I combine my human decisions with AI's generative capabilities to achieve original outcomes.	0.841 (50.992)	0.000	Rafner et al. (2025)
Factor 5	Code	MTPP. Metacognitive Transparency in Prompting Practices. Students' capacity to plan, document, monitor, adjust, and reflect on prompting practices to improve learning, transparency, and responsible AI interaction. Cronbach's Alpha (≥ 0.70) = 0.918; rho_A (Dijkstra–Henseler's ≥ 0.70) = 0.924; rho_C (Composite Reliability ≥ 0.70) = 0.931; AVE (≥ 0.50) = 0.577.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)

Juan Mejía Trejo

CAPÍTULO 10. Integración del modelo PLS-SEM

Metacognitive Transparency in Prompting Practices	MTPP1	I systematically record the prompts I use in order to learn from my own process.	0.676 (20.319)	0.000	Haidar et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP2	I adjust my prompts when I detect that AI responses do not meet my learning objectives.	0.768 (30.939)	0.000	Haidar et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP3	I analyze changes in AI responses to understand how my requests affect the results.	0.774 (31.853)	0.000	Tsakeni et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP4	I reflect on the clarity and precision of my prompts before submitting them.	0.796 (40.070)	0.000	Haidar et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP5	I document the reasoning I follow when modifying a prompt or interpreting a response.	0.711 (21.682)	0.000	Tsakeni et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP6	I compare different versions of the same prompt to identify which one produces more useful results.	0.701 (19.845)	0.000	Waalder et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP7	I evaluate whether my prompting strategies contribute to a deeper understanding of the topic.	0.815 (43.871)	0.000	Tsakeni et al. (2025)
Metacognitive Transparency in Prompting Practices	MTPP8	I plan the structure and purpose of the prompt before beginning my interaction with AI.	0.822 (46.923)	0.000	Miao & Cukurova (2024)
Metacognitive Transparency in Prompting Practices	MTPP9	I review the coherence between my initial objectives and the results obtained from AI.	0.777 (35.146)	0.000	Wang & Wang (2025)
Metacognitive Transparency in Prompting Practices	MTPP10	At the end of the activity, I reflect on what I learned from the process of interacting with AI.	0.741 (28.741)	0.000	Haidar et al. (2025)
Factor 6	Code	CAITIZEN. AI-assisted Sustainable Citizenship. Students' capacity to participate responsibly in AI-mediated societies by integrating critical literacy, ethics, fairness, creativity, metacognition, innovation, and social sustainability. Cronbach's Alpha (≥ 0.70) = 0.910; rho_A (Dijkstra–Henseler's ≥ 0.70) = 0.912; rho_C (Composite Reliability ≥ 0.70) = 0.925; AVE (≥ 0.50) = 0.553.	Outer loading (t-value)	p-Value	Associated references according to Mejía-Trejo (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 1	I identify myself as part of a society that coexists daily with artificial intelligence.	0.712 (25.121)	0.000	Mejía-Trejo (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 2	I consider it important to understand how AI influences my rights and responsibilities as a citizen.	0.740 (25.553)	0.000	Mejía-Trejo (2025)

Juan Mejía Trejo



AI-assisted Sustainable Citizenship	CAITIZEN 3	I believe that everyone should participate in public decisions related to the use of AI.	0.719 (24.357)	0.000	Decker et al. (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 4	I am interested in AI being used transparently by governments and institutions.	0.742 (27.366)	0.000	Papagiannidis et al. (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 5	I think that access to AI knowledge is a form of citizen empowerment.	0.731 (29.067)	0.000	Southworth et al. (2023)
AI-assisted Sustainable Citizenship	CAITIZEN 6	I am concerned that AI may reinforce inequalities if inclusive policies are not implemented.	0.658 (18.382)	0.000	González-Argote et al. (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 7	I trust that informed citizenship can guide the ethical development of AI.	0.779 (34.150)	0.000	Gunasekara et al. (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 8	I am willing to learn continuously in order to adapt to an AI-driven society.	0.750 (24.017)	0.000	Córdova-Esparza (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 9	I believe that being a citizen in the AI era implies actively collaborating in its regulation and oversight.	0.795 (37.658)	0.000	Papagiannidis et al. (2025)
AI-assisted Sustainable Citizenship	CAITIZEN 10	I value the importance of combining technological innovation with collective responsibility and social sustainability.	0.795 (39.911)	0.000	Mejía-Trejo (2025)

Nota. La tabla presenta los resultados del modelo de medición reflectivo para los seis constructos. Las cargas externas, los valores *t* y los valores *p* corresponden a los indicadores conservados después del refinamiento del modelo. **CAIL2** fue excluido porque su eliminación mejoró el **AVE** del constructo **CAIL**. Las referencias asociadas indican el sustento teórico utilizado para la adaptación o formulación de cada indicador.

Fuente: Elaboración propia con **SmartPLS4**.

Así, podemos resumir la siguiente secuencia de proceso en la **Tabla 10.2**

**Tabla 10.2. Resumen de proceso de la valuación del modelo de medición
CAITIZEN_PLS_sin CAIL2**

Fuente	Procedencia	Aporta
Instrumento base	Tabla de operacionalización del modelo CAITIZEN	Factor, código, ítem y referencias
CONFIABILIDAD.xlsx	SmartPLS → PLS Algorithm → Construct Reliability and Validity → Export Excel	Cronbach Alpha, rho_A, rho_C y AVE
BS_CARGAS EXTERNAS.xlsx	SmartPLS → Bootstrapping → Outer Loadings → Export Excel	Outer loading, t-value y p-value

Fuente: Elaboración propia con **SmartPLS4**.

Validez discriminante mediante Fornell–Larcker y HTMT del modelo CAITIZEN

Después de evaluar la consistencia interna, las cargas externas y la validez convergente del modelo **CAITIZEN_PLS_sin CAIL2**, se procedió a analizar la **validez discriminante**. Esta etapa permite verificar si los constructos **AFDJ**, **CAIL**, **CAITIZEN**, **EAR**, **HAIC** y **MTTP** son empíricamente distintos entre sí. Para ello, se utilizaron dos

salidas de **SmartPLS** exportadas en formato Excel: **FORNELL_LARCKER.xlsx** y **HTMT.xlsx**.

El archivo **FORNELL_LARCKER.xlsx** proviene del modo **PLS Algorithm**, específicamente de la salida **Fornell–Larcker Criterion**. De este archivo se tomaron los valores de la diagonal, que representan la **raíz cuadrada del AVE**, y las correlaciones entre constructos, ubicadas por debajo de la diagonal. Bajo este criterio, cada valor diagonal debe ser mayor que las correlaciones del constructo con las demás variables latentes.

El archivo **HTMT.xlsx** también proviene del modo **PLS Algorithm**, desde la salida **HTMT** de **SmartPLS**. De este archivo se tomaron los valores de la razón **Heterotrait–Monotrait**, los cuales se colocaron por encima de la diagonal en la tabla final. El criterio **HTMT** permite evaluar si los constructos son suficientemente diferentes entre sí. Valores menores a **0.85** indican validez discriminante adecuada bajo un criterio conservador, mientras que valores menores a **0.90** pueden aceptarse cuando los constructos están conceptualmente relacionados.

Con base en ambos archivos, se construyó la **Tabla 10.3**. La diagonal y la parte inferior de la tabla provienen de **FORNELL_LARCKER.xlsx**, mientras que la parte superior proviene de **HTMT.xlsx**. Esta integración permite presentar en una sola matriz los dos criterios principales de validez discriminante del modelo **CAITIZEN_PLS_sin CAIL2**.

Tabla 10.3. Evaluación de la validez discriminante mediante el criterio de Fornell–Larcker y la razón HTMT

Criterio de Fornell–Larcker: raíz cuadrada del AVE en la diagonal para evaluar la validez discriminante						
Criterio HTMT: razón HTMT ≤ 0.85 o ≤ 0.90 para evaluar la validez discriminante						
Construct	AFDJ	CAIL	CAITIZEN	EAR	HAIC	MTTP
AFDJ	0.762	0.794	0.824	0.827	0.788	0.763
CAIL	0.718	0.722	0.679	0.834	0.694	0.732
CAITIZEN	0.759	0.614	0.743	0.712	0.820	0.744
EAR	0.766	0.759	0.661	0.745	0.712	0.716
HAIC	0.735	0.637	0.765	0.673	0.810	0.752
MTTP	0.708	0.665	0.695	0.669	0.713	0.760

Nota. Los valores de la diagonal representan la raíz cuadrada del **AVE** de cada constructo, conforme al criterio de Fornell–Larcker. Los valores ubicados por debajo de la diagonal corresponden a las correlaciones entre constructos, mientras que los valores ubicados por encima de la diagonal corresponden a las razones **HTMT**. El **HTMT** evalúa la validez discriminante al determinar si los constructos son empíricamente distintos. Valores inferiores a 0.85 indican validez discriminante adecuada bajo un criterio conservador, mientras que valores inferiores a 0.90 pueden ser aceptables para constructos conceptualmente relacionados (Henseler et al., 2015). Además, el **Bootstrapping** puede utilizarse para verificar si los intervalos de confianza del **HTMT** incluyen el valor de 1.0; cuando dichos intervalos no incluyen 1.0, se obtiene evidencia estadística adicional de validez discriminante. **Franke y Sarstedt (2019)** respaldan el **HTMT** como un enfoque más confiable que los criterios tradicionales.

Fuente: Elaboración propia con **SmartPLS4**.

Así, podemos resumir la siguiente secuencia de proceso en la **Tabla 10.4**.

Tabla 10.4. Resumen de proceso de validación discriminante

Paso	Objetivo	Modo de cálculo en SmartPLS	Salida exportada	Archivo Excel	Datos utilizados	Ubicación en la tabla final
1	Partir del modelo depurado	CAITIZEN_PLS_sin CAIL2	Modelo ajustado	—	Constructos retenidos: AFDJ, CAIL, CAITIZEN , EAR, HAIC y MTPP	Base general de la tabla
2	Obtener la raíz cuadrada del AVE	PLS Algorithm	Fornell–Larcker Criterion	FORNELL_LARCKER.xlsx	Valores de la diagonal	Diagonal principal
3	Incorporar correlaciones entre constructos	PLS Algorithm	Fornell–Larcker Criterion	FORNELL_LARCKER.xlsx	Correlaciones entre variables latentes	Debajo de la diagonal
4	Obtener las razones HTMT	PLS Algorithm	HTMT	HTMT.xlsx	Valores HTMT entre constructos	Encima de la diagonal
5	Integrar la matriz final	Trabajo editorial en Excel	Matriz combinada	Tabla final elaborada	Diagonal, correlaciones y HTMT	Tabla X completa
6	Evaluar Fornell–Larcker	Revisión interpretativa	Fornell–Larcker Criterion	FORNELL_LARCKER.xlsx	La diagonal debe superar las correlaciones del constructo	Interpretación posterior
7	Evaluar HTMT	Revisión interpretativa	HTMT	HTMT.xlsx	HTMT ≤ 0.85 o ≤ 0.90	Interpretación posterior

Fuente: Elaboración propia con **SmartPLS4**.

Evaluación de colinealidad del modelo estructural mediante valores VIF internos

Después de evaluar la validez discriminante del modelo **CAITIZEN_PLS_sin CAIL2**, se revisó la posible colinealidad entre los constructos predictores del modelo estructural. Esta etapa permite determinar si las variables independientes **AFDJ**, **CAIL**, **EAR**, **HAIC** y **MTPP** presentan niveles excesivos de asociación al explicar la variable dependiente **CAITIZEN**. Para ello, se utilizó el **Variance Inflation Factor —VIF—**, indicador recomendado en **PLS-SEM** para identificar posibles problemas de redundancia estadística entre predictores antes de interpretar los coeficientes estructurales (Hair et al., 2019; Hair et al., 2022).

Los datos de la **Tabla 10.5** provienen del archivo **VIF.xlsx**, exportado desde la función de resultados en formato Excel de **SmartPLS4**, después de ejecutar el modelo **CAITIZEN_PLS_sin CAIL2** mediante el modo **PLS Algorithm**. En **SmartPLS**, los

Juan Mejía Trejo

valores **VIF** internos se localizan en la sección **Collinearity Statistics / Inner VIF Values**, correspondiente a la evaluación preliminar del modelo estructural. Por tanto, esta tabla no se calcula mediante **Bootstrapping** ni mediante **Model Fit**: el primero se reserva para la significancia de los coeficientes estructurales, mientras que el segundo corresponde a otros indicadores de ajuste global o aproximado del modelo.

Tabla 10.5. Evaluación de colinealidad del modelo estructural mediante valores VIF internos

Relación Estructural	Factor de Inflación de la Varianza (VIF)	Evaluación
AFDJ → CAITIZEN	3.398	Aceptable bajo el criterio VIF < 5; por encima del criterio estricto VIF < 3.
CAIL → CAITIZEN	2.785	Sin colinealidad crítica.
EAR → CAITIZEN	3.213	Aceptable bajo el criterio VIF < 5; por encima del criterio estricto VIF < 3.
HAIC → CAITIZEN	2.671	Sin colinealidad crítica.
MTPP → CAITIZEN	2.583	Sin colinealidad crítica.

Nota. VIF = Factor de Inflación de la Varianza. Los valores **VIF** internos fueron obtenidos del algoritmo **PLS-SEM** en **SmartPLS 4** para evaluar la colinealidad entre los constructos predictores de **CAITIZEN**. Valores inferiores a 3.00 indican ausencia de colinealidad crítica bajo un criterio conservador, mientras que valores inferiores a 5.00 suelen considerarse aceptables en **PLS-SEM**. Los resultados muestran que todas las relaciones predictoras se mantienen por debajo del umbral máximo aceptable de 5.00; sin embargo, **AFDJ → CAITIZEN** y **EAR → CAITIZEN** superan el criterio más estricto de 3.00, por lo que deben interpretarse con cautela.

Fuente: Elaboración propia con **SmartPLS4**.

Los resultados muestran que todos los valores **VIF** se encuentran por debajo del umbral crítico de **5.00**, lo que indica ausencia de colinealidad severa entre los predictores del constructo **CAITIZEN**. En particular, se reportaron los siguientes valores: **AFDJ → CAITIZEN = 3.398**, **CAIL → CAITIZEN = 2.785**, **EAR → CAITIZEN = 3.213**, **HAIC → CAITIZEN = 2.671** y **MTPP → CAITIZEN = 2.583**. En consecuencia, el modelo puede avanzar hacia la evaluación de los coeficientes estructurales sin evidencia de redundancia crítica entre predictores.

No obstante, bajo el criterio más conservador de **VIF < 3.00**, las relaciones **AFDJ → CAITIZEN** y **EAR → CAITIZEN** deben considerarse aceptables, aunque ubicadas en una zona de atención metodológica. Esto no invalida el modelo, debido a que ambas permanecen por debajo del umbral de **5.00**; sin embargo, sugiere una mayor cercanía conceptual con otros predictores. En el modelo **CAITIZEN**, esta proximidad es teóricamente razonable, ya que la justicia algorítmica y la responsabilidad ética forman parte de una arquitectura normativa común vinculada con la ciudadanía sostenible asistida por inteligencia artificial.

Evaluación del modelo estructural

Una vez confirmada la calidad del **modelo de medición reflectivo**, mediante la revisión de cargas externas, confiabilidad interna, validez convergente, validez discriminante y ausencia de colinealidad crítica, se procedió a evaluar el **modelo estructural CAITIZEN_PLS_sin CAIL2**. Esta etapa permite determinar si las trayectorias hipotéticas propuestas entre **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTPP** y **CAITIZEN** cuentan con respaldo empírico. Mientras la evaluación del modelo de medición verifica si los indicadores representan adecuadamente a sus constructos, la evaluación estructural analiza la dirección, magnitud, significancia estadística y relevancia práctica de las relaciones entre variables latentes. Esta secuencia responde al procedimiento recomendado en **PLS-SEM**, donde primero se valida el modelo de medición y posteriormente se interpreta el modelo estructural (Hair et al., 2019; Hair et al., 2022).

Para construir la **Tabla 10.6**, se utilizaron dos salidas principales de **SmartPLS4**. En primer lugar, se empleó el **Bootstrapping** ejecutado sobre el modelo normal final depurado **CAITIZEN_PLS_sin CAIL2**. De esta salida se obtuvieron los coeficientes estructurales β , los **t-values**, los **p-values** y los intervalos de confianza al 95%. Estos resultados permiten valorar la significancia estadística de las rutas estructurales y decidir si las hipótesis reciben o no apoyo empírico. En segundo lugar, se utilizó el **PLS-SEM Algorithm**, aplicado al mismo modelo, para obtener los tamaños del efecto f^2 desde la salida **f-square**. El modelo **FIT** no se utilizó para construir esta tabla, porque su función se limita exclusivamente a la evaluación del ajuste global del modelo saturado mediante **SRMR**, **d_ULS** y **d_G**.

Los coeficientes estructurales β indican la dirección y magnitud de cada relación entre constructos. Un coeficiente positivo señala que, al aumentar el predictor, también tiende a incrementarse **CAITIZEN**; en cambio, un coeficiente negativo indica una relación inversa. Sin embargo, la interpretación del coeficiente β no debe realizarse de manera aislada, ya que su relevancia depende de su significancia estadística y de la estabilidad del intervalo de confianza. Por ello, para aceptar o rechazar una hipótesis se revisan conjuntamente el **t-value**, el **p-value** y el intervalo de confianza al 95%. En este análisis, una hipótesis se acepta cuando $p < 0.05$ y el intervalo de confianza no incluye el valor cero; si el intervalo cruza cero o el valor p es mayor a 0.05, la hipótesis se rechaza (Hair et al., 2019; Hair et al., 2022).

El tamaño del efecto f^2 complementa la interpretación de los coeficientes estructurales, porque permite valorar la relevancia práctica de cada predictor dentro del modelo. Mientras el coeficiente β muestra la fuerza y dirección de la relación, el f^2 indica cuánto contribuye cada constructo predictor a explicar el constructo endógeno **CAITIZEN**. En esta investigación, los valores f^2 se interpretaron con base en los criterios propuestos por Cohen (1988): valores menores a **0.02** indican ausencia de efecto o efecto insignificante; valores entre **0.02 y 0.15** indican efecto pequeño; valores entre **0.15 y 0.35** indican efecto medio; y valores iguales o superiores a **0.35** indican efecto grande. Estos criterios son utilizados ampliamente en **PLS-SEM** para

complementar la interpretación estadística con la relevancia práctica del modelo (Cohen, 1988; Hair et al., 2019).

Con base en este procedimiento, la **Tabla 10.6** integra la evaluación de las cinco hipótesis del modelo directo **CAITIZEN**: **H1: CAIL** → **CAITIZEN**, **H2: EAR** → **CAITIZEN**, **H3: AFDJ** → **CAITIZEN**, **H4: HAIC** → **CAITIZEN** y **H5: MTPP** → **CAITIZEN**. Esta tabla permite observar simultáneamente qué relaciones son estadísticamente significativas, cuáles no reciben apoyo empírico y qué relevancia práctica tiene cada predictor en la explicación de la ciudadanía sostenible asistida por inteligencia artificial. En consecuencia, la **Tabla 10.6** funciona como la evidencia central para valorar el grado de soporte empírico del modelo estructural propuesto.

Tabla 10.6. Evaluación de hipótesis, coeficientes estructurales y tamaño del efecto f^2 del modelo CAITIZEN_PLS_sin CAIL2

Hipótesis	Trayectorias β	Coeficiente β [t-value; p-value]	IC 95% [2.5%; 97.5%]	Resultado	Tamaño del Efecto f^2	Interpretación del Efecto
H1: CAIL positively and significantly predicts CAITIZEN .	CAIL → CAITIZEN	-0.019 [0.443; 0.658]	[-0.102; 0.064]	Rechazado	0.000	Sin Efecto
H2: EAR positively and significantly predicts CAITIZEN .	EAR → CAITIZEN	0.054 [0.987; 0.324]	[-0.051; 0.162]	Rechazado	0.003	Efecto Insignificante
H3: AFDJ positively and significantly predicts CAITIZEN .	AFDJ → CAITIZEN	0.340 [5.550; 0.000]	[0.218; 0.455]	Aceptada	0.107	Pequeño
H4: HAIC positively and significantly predicts CAITIZEN .	HAIC → CAITIZEN	0.374 [7.065; 0.000]	[0.272; 0.480]	Aceptada	0.164	Medio
H5: MTPP positively and significantly predicts CAITIZEN .	MTPP → CAITIZEN	0.164 [3.641; 0.000]	[0.074; 0.252]	Aceptada	0.033	Pequeño

Nota. Los coeficientes β , t-values, p-values e intervalos de confianza provienen del **Bootstrapping** aplicado al modelo normal final depurado **CAITIZEN_PLS_sin CAIL2**. Los tamaños del efecto f^2 provienen del **PLS-SEM Algorithm**, salida **f-square**. La interpretación del tamaño del efecto se realizó con base en los criterios de Cohen (1988): $f^2 < 0.02$ = sin efecto o efecto insignificante, $0.02 \leq f^2 < 0.15$ = efecto pequeño, $0.15 \leq f^2 < 0.35$ = efecto medio y $f^2 \geq 0.35$ = efecto grande. Estos criterios son también utilizados en la literatura **PLS-SEM** para interpretar la relevancia práctica de los predictores estructurales (Hair et al., 2019; Hair et al., 2022). El modelo FIT no se utiliza para esta tabla; solo

corresponde a ajuste global mediante **SRMR**, **d_ULS** y **d_G**. Fuente: Elaboración propia con **SmartPLS4**. Así, podemos resumir la siguiente secuencia de proceso en la **Tabla 10.7**.

Tabla 10.7. Secuencia de proceso

Paso	Qué se hace	Modo en SmartPLS	Archivo / salida	Dato que aporta
1	Partir del modelo depurado	CAITIZEN_PLS_sin CAIL2	Modelo ajustado	Relaciones estructurales H1–H5
2	Ejecutar el modelo estructural con remuestreo	Bootstrapping	PATH_COEFF.xlsx / T-PVALUES.xlsx	Coefficiente β , <i>t-value</i> y p-value
3	Tomar el coeficiente β	Bootstrapping	Path Coefficients	Original sample (O)
4	Tomar significancia estadística	Bootstrapping	Path Coefficients	T statistics y P values
5	Tomar intervalos de confianza	Bootstrapping	Confidence Intervals	IC 95% [2.5%; 97.5%]
6	Ejecutar o revisar el algoritmo PLS	PLS Algorithm	fsquare.xlsx	Tamaño del efecto f^2
7	Decidir resultado de hipótesis	Interpretación del investigador	Tabla integrada	Aceptada / Rechazada
8	Interpretar tamaño del efecto	Criterios metodológicos	Valor f^2	Sin efecto, pequeño, medio o grande
9	Integrar todo en la tabla final	Trabajo editorial	Tabla 10.6	Tabla completa del modelo estructural

Evaluación predictiva del modelo CAITIZEN mediante PLSpredict

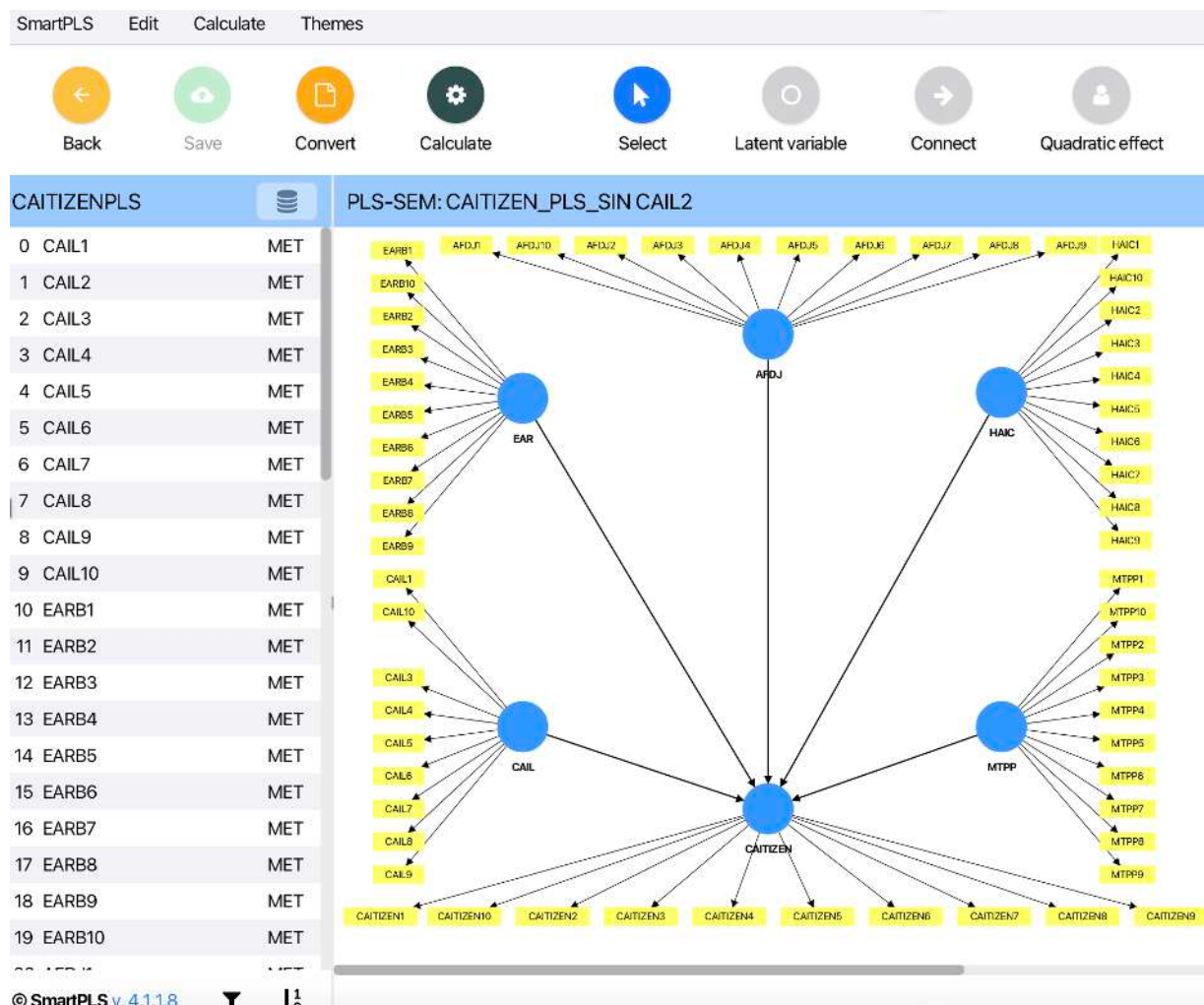
Una vez evaluado el **modelo de medición**, la **validez discriminante**, la **colinealidad interna** y el **modelo estructural**, se procedió a examinar la **capacidad predictiva fuera de muestra** del modelo **CAITIZEN_PLS_sin CAIL2**. Esta fase complementa la evaluación explicativa, porque permite valorar si el modelo no solo explica la varianza de **CAITIZEN** dentro de la muestra, sino si también puede anticipar adecuadamente los valores de sus indicadores observables. En **PLS-SEM**, esta etapa se reporta mediante **PLSpredict**, procedimiento que permite comparar el desempeño predictivo del modelo **PLS-SEM** frente a un modelo lineal de referencia, a partir de indicadores como **$Q^2_{predict}$** y **RMSE** (Shmueli et al., 2019). Esta lógica es consistente con la secuencia metodológica del capítulo, donde la capacidad predictiva se reporta después de la medición, validez discriminante, colinealidad, coeficientes estructurales, significancia, **R^2** y **f^2** .

Para iniciar el procedimiento, se mantuvo abierto en **SmartPLS4** el modelo normal final depurado **CAITIZEN_PLS_sin CAIL2**. Este es el mismo modelo utilizado para reportar cargas externas, confiabilidad, **AVE**, **HTMT**, **VIF**, coeficientes estructurales,

Juan Mejía Trejo

hipótesis, R^2 y f^2 . No debe utilizarse el modelo **FIT**, ya que este se reserva exclusivamente para el ajuste global del modelo saturado mediante **SRMR**, **d_ULS** y **d_G**. En la **Figura 10.1**, se muestra el modelo directo en el que **CAIL**, **EAR**, **AFDJ**, **HAIC** y **MTTP** predicen al constructo endógeno **CAITIZEN**.

Figura 10.1. Modelo CAITIZEN_PLS_sin CAIL2 previo a la ejecución de PLSpredict.

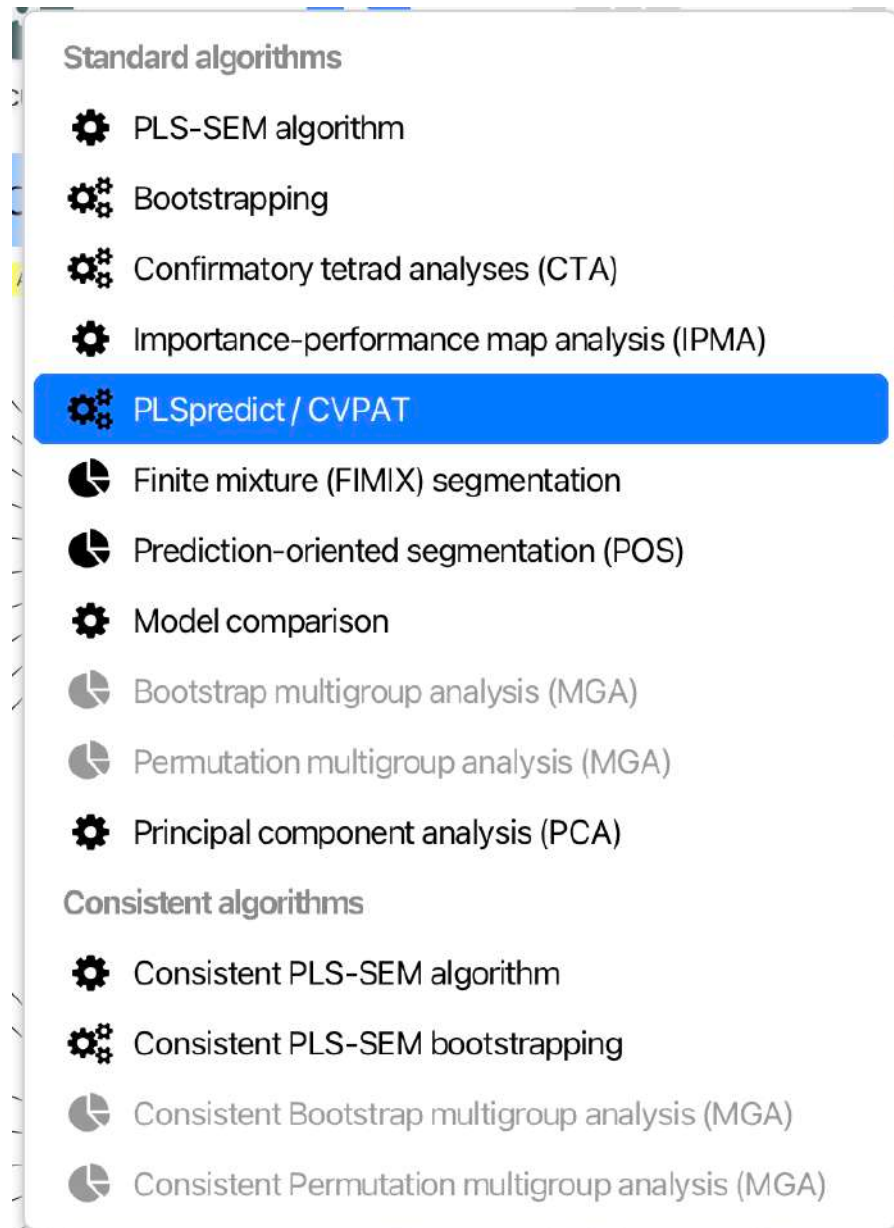


Fuente: Elaboración propia con **SmartPLS4**

Posteriormente, desde el menú **Calculate**, se seleccionó la opción **PLSpredict / CVPAT**. Este procedimiento es distinto del **PLS-SEM Algorithm** y del **Bootstrapping**. El **PLS-SEM Algorithm** se utiliza para obtener, entre otros resultados, confiabilidad, **AVE**, **VIF**, R^2 y f^2 ; el **Bootstrapping** se emplea para valorar la significancia estadística de cargas externas y trayectorias estructurales; en cambio, **PLSpredict / CVPAT** se utiliza para evaluar el desempeño predictivo del modelo fuera de muestra. La selección del procedimiento se muestra en la **Figura 10.2**.

Juan Mejía Tréjo

Figura 10.2. Selección del procedimiento *PLSpredict* / *CVPAT* en SmartPLS4



Fuente: Elaboración propia con **SmartPLS4**.

Después de seleccionar ***PLSpredict* / *CVPAT***, **SmartPLS4** despliega la ventana de configuración del análisis. Para este caso se utilizaron **17 folds**, **10 repeticiones**, tipo de predicción ***Early antecedents*** y generador aleatorio con ***Fixed seed***. Esta configuración permite dividir los datos en subconjuntos de entrenamiento y validación, repetir el procedimiento y obtener predicciones fuera de muestra para los indicadores del constructo endógeno. En este análisis, la evaluación se centró en los indicadores **CAITIZEN1** a **CAITIZEN10**, porque representan las manifestaciones observables de la ciudadanía asistida por inteligencia artificial. Ver **Figura 10.3**.

Juan Mejía Tréjo

Figura 10.3. Configuración de *PLSpredict* / CVPAT para el modelo CAITIZEN_PLS_sin CAIL2

The PLSpredict algorithm has been developed by Shmueli et al. (2016, 2019). The method uses training and holdout samples to generate and evaluate predictions from PLS path model estimations. The results also include the outcomes of the cross-validated predictive ability test (CVPAT; Sharma et al., 2023), which also applies an out-of-sample prediction approach to assess a model's predictive capabilities.

[More info](#)

PLSpredict / CVPAT **PLS setup**

[Number of folds](#) 17

[No. of repetitions](#) 10

[Prediction type](#) Early antecedents

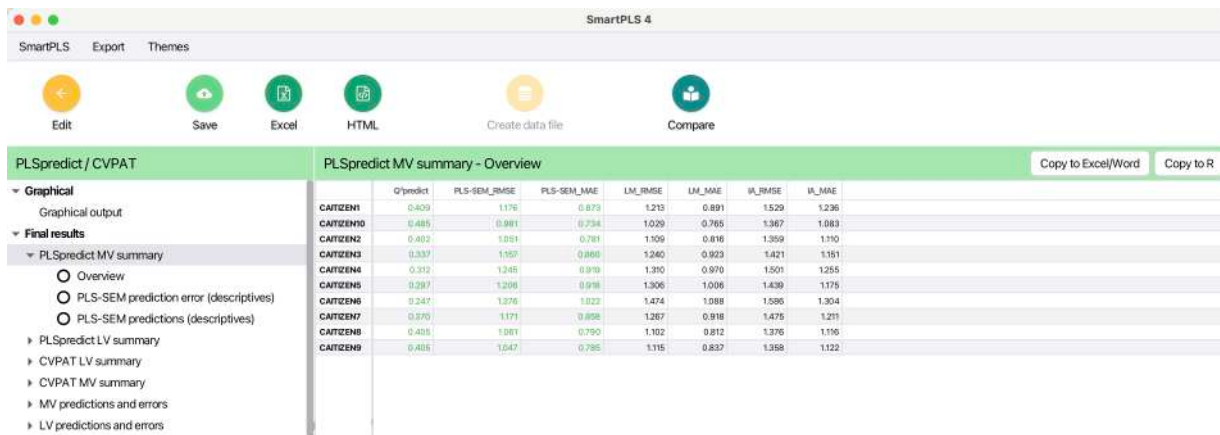
[Random number generator](#) Fixed seed

[Default settings](#) ☐ Open report

Fuente: Elaboración propia con **SmartPLS4**.

Una vez ejecutado el cálculo, **SmartPLS4** genera la salida ***PLSpredict MV summary – Overview***. De esta salida se tomaron los valores **$Q^2_{predict}$** , **PLS-SEM_RMSE** y **LM_RMSE**. El valor **$Q^2_{predict}$** permite valorar la relevancia predictiva fuera de muestra; cuando es superior a cero, indica que el modelo posee capacidad predictiva para el indicador evaluado. Por su parte, **RMSE PLS-SEM** representa el error de predicción del modelo **PLS-SEM**, mientras que **RMSE LM** corresponde al error del modelo lineal de referencia. La salida utilizada para construir la tabla se muestra en la **Figura 10.4**.

Figura 10.4. Resultados PLSpredict MV summary – Overview para los indicadores CAITIZEN



Fuente: Elaboración propia con SmartPLS4.

Con base en la salida anterior, se elaboró la **Tabla 10.8**. En la primera columna se colocaron los indicadores del constructo endógeno **CAITIZEN**. En la segunda columna se registraron los valores **Q²_predict**. En la tercera columna se reportó el **RMSE PLS-SEM**. En la cuarta columna se colocó el **RMSE LM**. Finalmente, la columna **RMSE PLS-LM** se calculó restando el error del modelo lineal al error del modelo **PLS-SEM**, es decir: **RMSE PLS-SEM – RMSE LM**. Cuando esta diferencia es negativa, significa que el modelo **PLS-SEM** produjo menor error de predicción que el modelo lineal de referencia.

Tabla 10.8. Evaluación PLSpredict para CAITIZEN

Indicador	Q ² _predict	RMSE PLS-SEM	RMSE LM	RMSE PLS-LM	Evaluación
CAITIZEN1	0.409	1.176	1.213	-0.037	PLS mejor
CAITIZEN2	0.402	1.051	1.109	-0.058	PLS mejor
CAITIZEN3	0.337	1.157	1.240	-0.083	PLS mejor
CAITIZEN4	0.312	1.245	1.310	-0.065	PLS mejor
CAITIZEN5	0.297	1.206	1.306	-0.100	PLS mejor
CAITIZEN6	0.247	1.376	1.474	-0.098	PLS mejor
CAITIZEN7	0.370	1.171	1.267	-0.096	PLS mejor
CAITIZEN8	0.405	1.061	1.102	-0.041	PLS mejor
CAITIZEN9	0.405	1.047	1.115	-0.068	PLS mejor
CAITIZEN10	0.485	0.981	1.029	-0.048	PLS mejor

Nota. **CAITIZEN** = Ciudadanía asistida por inteligencia artificial para una formación sostenible, ética y en red. **Q²_predict > 0** indica relevancia predictiva. **RMSE** = raíz del error cuadrático medio; **LM** = modelo lineal de referencia. Los valores negativos en la columna **RMSE PLS-LM** indican que el modelo **PLS-SEM** produjo errores de predicción menores que el modelo lineal de referencia.

El modelo **CAITIZEN** presenta **relevancia predictiva fuera de muestra**, ya que todos los indicadores muestran valores **Q²_predict superiores a cero**. Asimismo, en los diez indicadores, el **RMSE del modelo PLS-SEM** es menor que el **RMSE del modelo lineal de referencia LM**. Por ello, las diferencias **PLS-LM** son negativas en todos los casos, lo cual indica que el modelo **PLS-SEM** predice mejor que el

modelo lineal comparativo. En conjunto, estos resultados respaldan una **capacidad predictiva favorable del modelo CAITIZEN**.

Fuente: Elaboración propia con **SmartPLS4**.

Interpretación de resultados para discusión

La **interpretación de resultados** constituye el puente entre el reporte estadístico y la discusión académica del modelo **CAITIZEN** (Mejía-Trejo, 2025). Una vez evaluados el modelo de medición, la validez discriminante, la colinealidad estructural, las hipótesis, los tamaños del efecto y la capacidad predictiva, es necesario traducir los resultados numéricos en una lectura sustantiva del fenómeno estudiado. En esta etapa, el objetivo no es repetir los valores de las tablas, sino explicar qué significan en términos teóricos, metodológicos y aplicados para la ciudadanía sostenible asistida por inteligencia artificial. Esta lógica coincide con la función del reporte académico en **PLS-SEM**: transformar los resultados de **SmartPLS4** en una explicación metodológica, teórica y aplicada del modelo analizado.

En primer lugar, los resultados del **modelo de medición reflectivo** muestran que los constructos **CAIL**, **EAR**, **AFDJ**, **HAIC**, **MTTP** y **CAITIZEN** presentan condiciones aceptables de confiabilidad interna y validez convergente después de la depuración del indicador **CAIL2**. Esta decisión fortaleció la calidad psicométrica del constructo **CAIL**, particularmente en términos de **AVE**, sin eliminar la lógica conceptual del modelo. Por tanto, la estructura de medición puede considerarse defendible, ya que los indicadores retenidos representan adecuadamente las variables latentes propuestas. Esta interpretación es consistente con la secuencia recomendada en **PLS-SEM**, donde la interpretación estructural solo es válida si previamente se verifica la calidad del modelo de medición (Hair et al., 2019; Hair et al., 2022).

En segundo lugar, la evaluación de la **validez discriminante** permite sostener que los constructos del modelo, aunque conceptualmente relacionados, no son equivalentes entre sí. Esto es relevante porque **alfabetización crítica en IA**, **responsabilidad ética**, **justicia de datos**, **colaboración humano-IA** y **transparencia metacognitiva en prompting** pertenecen a un mismo campo de formación ciudadana, pero no miden exactamente lo mismo. La integración de los criterios **Fornell-Larcker** y **HTMT** permite argumentar que cada dimensión conserva identidad empírica propia. En particular, el uso de **HTMT** fortalece la evaluación, ya que se considera un criterio más sensible para detectar problemas de validez discriminante entre constructos conceptualmente cercanos (Henseler et al., 2015; Franke & Sarstedt, 2019).

En tercer lugar, los valores **VIF internos** muestran que no existe colinealidad crítica entre los predictores de **CAITIZEN**. Aunque algunas relaciones se ubican por encima del criterio conservador de **VIF < 3.00**, permanecen por debajo del umbral máximo aceptable de **VIF < 5.00**, por lo que el modelo puede interpretarse sin evidencia de redundancia severa entre predictores. Esta lectura es metodológicamente importante porque, antes de interpretar los coeficientes estructurales, se debe descartar que las

Juan Mejía Trejo

relaciones estimadas estén distorsionadas por asociaciones excesivas entre variables independientes (Hair et al., 2019; Hair et al., 2022). En términos sustantivos, la cercanía entre variables como **AFDJ** y **EAR** es esperable, porque ambas pertenecen a una arquitectura normativa vinculada con ética, justicia, equidad y responsabilidad en el uso de IA.

En cuarto lugar, la evaluación del **modelo estructural** muestra que el modelo **CAITIZEN_PLS_sin CAIL2** obtiene apoyo empírico parcial. Las hipótesis **H3**, **H4** y **H5** fueron aceptadas, mientras que **H1** y **H2** fueron rechazadas. Esto significa que **AFDJ**, **HAIC** y **MTPP** son predictores directos significativos de **CAITIZEN**, mientras que **CAIL** y **EAR** no muestran efectos directos estadísticamente significativos en esta especificación del modelo. Esta interpretación debe realizarse considerando de manera conjunta los coeficientes β , los valores **t**, los valores **p**, los intervalos de confianza y los tamaños del efecto f^2 , ya que la significancia estadística no debe sustituir la relevancia teórica ni la lectura sustantiva del fenómeno (Cohen, 1988; Hair et al., 2019).

El resultado más relevante corresponde a la relación **HAIC** → **CAITIZEN**, que presenta el mayor coeficiente estructural y el mayor tamaño del efecto. Esto sugiere que la **colaboración creativa humano-IA** ocupa un papel central en la configuración de la ciudadanía sostenible asistida por IA. En términos aplicados, el hallazgo indica que no basta con que los estudiantes comprendan la IA o reconozcan sus riesgos; también deben aprender a interactuar creativamente con ella, combinar juicio humano con capacidades generativas y producir soluciones con sentido académico, ético y social. Por tanto, **HAIC** aparece como una dimensión clave para pasar del conocimiento sobre IA al uso colaborativo, productivo y socialmente orientado de la IA.

La relación **AFDJ** → **CAITIZEN** también resulta significativa, lo que confirma la importancia de la **justicia de datos**, la equidad algorítmica y la conciencia sobre sesgos en la formación ciudadana asistida por IA. Este resultado indica que la ciudadanía en contextos mediados por inteligencia artificial requiere reconocer cómo los datos, los algoritmos y las decisiones automatizadas pueden producir exclusión, desigualdad o falta de transparencia. En consecuencia, **AFDJ** funciona como una dimensión normativa y crítica del modelo, porque vincula la ciudadanía asistida por IA con el derecho a comprender, cuestionar y supervisar los sistemas algorítmicos que inciden en la vida académica, institucional y social.

La relación **MTPP** → **CAITIZEN** muestra un efecto positivo y significativo, aunque de menor magnitud. Esto significa que la **transparencia metacognitiva en las prácticas de prompting** contribuye a la ciudadanía asistida por IA, especialmente porque promueve reflexión, documentación, ajuste y evaluación consciente de la interacción con sistemas inteligentes. En términos formativos, este resultado permite argumentar que el prompting no debe entenderse solo como habilidad técnica, sino como una práctica reflexiva, trazable y responsable. Así, **MTPP** aporta al modelo una dimensión procedimental: no solo importa usar IA, sino saber explicar cómo se interactuó con ella, qué decisiones se tomaron y cómo se evaluaron sus resultados.

En contraste, la no significancia de **CAIL** → **CAITIZEN** puede interpretarse como evidencia de que la **alfabetización crítica en IA**, por sí sola, no garantiza ciudadanía sostenible asistida por **IA** cuando se analiza simultáneamente con predictores más aplicados. La comprensión de los sistemas de **IA** es necesaria, pero puede no ser suficiente para explicar comportamientos ciudadanos avanzados. Su papel podría ser más basal, formativo o indirecto, funcionando como condición previa para desarrollar justicia de datos, colaboración humano-**IA** o transparencia metacognitiva. Esta interpretación permite conservar la relevancia conceptual de **CAIL**, aunque su efecto directo no haya sido estadísticamente significativo.

De manera similar, la no significancia de **EAR** → **CAITIZEN** sugiere que la **conciencia ética y responsabilidad** no necesariamente se traduce en ciudadanía asistida por **IA** si no se articula con prácticas concretas de justicia, colaboración y reflexión. Esto no invalida la importancia de la ética; más bien indica que la ética declarativa puede requerir mediaciones prácticas para convertirse en acción ciudadana. En el modelo **CAITIZEN**, la ética no debe limitarse al reconocimiento de principios, sino vincularse con conductas observables como auditar sesgos, proteger datos, declarar usos de **IA**, documentar prompts y asumir responsabilidad por los resultados generados con apoyo algorítmico.

Finalmente, los resultados de **PLSpredict** fortalecen la interpretación del modelo porque muestran que **CAITIZEN_PLS_sin CAIL2** no solo posee utilidad explicativa, sino también capacidad predictiva fuera de muestra. Todos los indicadores del constructo **CAITIZEN** presentan $Q^2_{predict} > 0$, y en todos los casos el error **RMSE PLS-SEM** es menor que el error del modelo lineal de referencia. Esta evidencia permite sostener que el modelo tiene desempeño predictivo favorable y que sus constructos antecedentes pueden anticipar adecuadamente las manifestaciones observables de la ciudadanía asistida por inteligencia artificial. En **PLS-SEM**, esta evaluación es relevante porque conecta la explicación del modelo con su utilidad predictiva fuera de muestra (Shmueli et al., 2019).

En conjunto, la discusión debe presentar al modelo **CAITIZEN** como una propuesta con soporte empírico parcial, pero metodológicamente defendible. Los resultados sugieren que la ciudadanía sostenible asistida por **IA** se explica principalmente por capacidades aplicadas y reflexivas: **justicia de datos, colaboración creativo-humano-IA y transparencia metacognitiva en prompting**. En cambio, la **alfabetización crítica en IA** y la **responsabilidad ética** parecen operar como fundamentos conceptuales necesarios, aunque no como predictores directos significativos en el modelo estructural evaluado.

En síntesis, la interpretación de resultados permite afirmar que el modelo **CAITIZEN** aporta evidencia para comprender la ciudadanía asistida por inteligencia artificial como una competencia compleja, integrada por dimensiones críticas, éticas, justas, creativas y metacognitivas. Sin embargo, los resultados muestran que las dimensiones con mayor poder explicativo directo son aquellas vinculadas con la acción aplicada: colaborar creativamente con, reconocer la justicia de datos y transparentar el proceso

de prompting. Esta lectura debe orientar la discusión final, las implicaciones educativas y las recomendaciones para futuras investigaciones.

Síntesis de procedimiento

La **síntesis del procedimiento** permite cerrar el capítulo mostrando la secuencia completa seguida para reportar el modelo **CAITIZEN_PLS_sin CAIL2**. Su función es ordenar, en un solo apartado, qué se evaluó, con qué procedimiento de **SmartPLS4**, qué datos se extrajeron y en qué tabla quedaron integrados. Con ello se garantiza la trazabilidad metodológica del reporte y se evita confundir los distintos usos del **PLS-SEM Algorithm**, el **Bootstrapping**, **PLSpredict** y el modelo **FIT**. Esta organización responde al objetivo central del capítulo: transformar los resultados estadísticos en una explicación metodológica, teórica y aplicada del modelo **CAITIZEN**.

El procedimiento inició con el **modelo normal final depurado**, denominado **CAITIZEN_PLS_sin CAIL2**. Sobre este modelo se evaluó primero el **modelo de medición reflectivo**, mediante confiabilidad interna, validez convergente, cargas externas y significancia de los indicadores. Después se revisó la **validez discriminante**, mediante los criterios **Fornell-Larcker** y **HTMT**. Posteriormente se evaluó la **colinealidad estructural** mediante los valores **VIF internos**. Una vez cumplidas estas etapas, se procedió a la evaluación del **modelo estructural**, incorporando coeficientes β , **t-values**, **p-values**, intervalos de confianza, aceptación o rechazo de hipótesis y tamaños del efecto f^2 .

Después de la evaluación estructural, se examinó la **capacidad predictiva fuera de muestra** mediante **PLSpredict / CVPAT**, tomando como referencia los valores **Q²_predict**, **RMSE PLS-SEM**, **RMSE LM** y la diferencia **RMSE PLS-LM**. Finalmente, el modelo **FIT** se mantuvo como una copia separada, utilizada únicamente para el análisis de ajuste global del modelo saturado mediante **SRMR**, **d_ULS**, **d_G**, **HI95** y **HI99**. Por tanto, el modelo **FIT** no debe mezclarse con las tablas de medición, hipótesis, R^2 , f^2 ni **PLSpredict**. Ver **Tabla 10.9**.

Tabla 10.9. Síntesis del procedimiento de reporte PLS-SEM del modelo CAITIZEN

Etapas	Qué se evalúa	Modelo usado	Modo / salida en SmartPLS4	Datos que aporta	Producto en el capítulo
1	Depuración del modelo	CAITIZEN_PLS_sin CAIL2	Revisión del modelo ajustado	Indicadores retenidos y exclusión de CAIL2	Modelo final depurado
2	Confiabilidad interna y validez convergente	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → Construct Reliability and Validity	Cronbach Alpha, rho_A, rho_C y AVE	Tabla 10.1

Juan Mejía Trejo

CAPÍTULO 10. Integración del modelo PLS-SEM

Etapas	Qué se evalúa	Modelo usado	Modo / salida en SmartPLS4	Datos que aporta	Producto en el capítulo
3	Cargas externas y significancia de indicadores	CAITIZEN_PLS_sin CAIL2	Bootstrapping → Outer Loadings	Outer loading, t-value y p-value	Tabla 10.1
4	Resumen del modelo de medición	CAITIZEN_PLS_sin CAIL2	Integración editorial	Fuente, procedencia y datos utilizados	Tabla 10.2
5	Validez discriminante Fornell–Larcker	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → Fornell–Larcker Criterion	Raíz cuadrada del AVE y correlaciones	Tabla 10.3
6	Validez discriminante HTMT	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → HTMT	Razones HTMT entre constructos	Tabla 10.3
7	Resumen de validez discriminante	CAITIZEN_PLS_sin CAIL2	Integración editorial	Procedencia de Fornell–Larcker y HTMT	Tabla 10.4
8	Colinealidad estructural	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → Inner VIF Values	VIF de cada predictor hacia CAITIZEN	Tabla 10.5
9	Capacidad explicativa	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → R-square	R ² y R ² ajustado de CAITIZEN	Tabla de R ² / R ² ajustado
10	Coefficientes estructurales e hipótesis	CAITIZEN_PLS_sin CAIL2	Bootstrapping → Path Coefficients	β, t-value, p-value e intervalos de confianza	Tabla 10.6
11	Tamaño del efecto	CAITIZEN_PLS_sin CAIL2	PLS-SEM Algorithm → f-square	f ² de cada predictor sobre CAITIZEN	Tabla 10.6
12	Resumen del modelo estructural	CAITIZEN_PLS_sin CAIL2	Integración editorial	Secuencia para β, t, p, IC, hipótesis y f ²	Tabla 10.7
13	Capacidad predictiva fuera de muestra	CAITIZEN_PLS_sin CAIL2	PLSpredict / CVPAT → PLSpredict MV summary	Q ² _predict, RMSE PLS-SEM , RMSE LM y RMSE PLS-LM	Tabla 10.8
14	Ajuste global del modelo saturado	CAITIZEN_PLS_sin CAIL2_FIT	Model Fit / Bootstrapping FIT	SRMR, d_ULS, d_G, HI95 y HI99	Tabla específica de ajuste global
15	Interpretación de resultados	CAITIZEN_PLS_sin CAIL2 y FIT por separado	Revisión metodológica y sustantiva	Lectura integrada de medición, estructura, predicción y ajuste	Apartado de discusión
16	Cierre del reporte	Integración final	Síntesis editorial	Coherencia del procedimiento completo	Conclusión del capítulo

Juan Mejía Trejo

Nota. El **PLS-SEM Algorithm** se utilizó para obtener confiabilidad, AVE, Fornell–Larcker, HTMT, VIF, R^2 , R^2 ajustado y f^2 . El **Bootstrapping** se utilizó para obtener la significancia estadística de las cargas externas y de las trayectorias estructurales. **PLSpredict / CVPAT** se utilizó para evaluar la capacidad predictiva fuera de muestra. El modelo **CAITIZEN_PLS_sin CAIL2_FIT** se utilizó únicamente para evaluar el ajuste global del modelo saturado mediante **SRMR**, **d_ULS**, **d_G**, **HI95** y **HI99**.

Fuente: Elaboración propia con **SmartPLS4**.

En síntesis, el reporte siguió una secuencia progresiva: **modelo de medición** → **validez discriminante** → **colinealidad** → **capacidad explicativa** → **hipótesis** → **tamaño del efecto** → **capacidad predictiva** → **ajuste global** → **interpretación para discusión**. Esta organización permite presentar el modelo **CAITIZEN** con claridad metodológica, evitando mezclar salidas que cumplen funciones distintas. El modelo normal **CAITIZEN_PLS_sin CAIL2** sostiene la evaluación de medición, estructura y predicción; mientras que el modelo **FIT** permanece separado para el análisis de ajuste global.

REFERENCIAS



- Becker, J.-M., Klein, K., & Wetzels, M. (2012). Hierarchical latent variable models in **PLS-SEM**: Guidelines for using reflective-formative type models. *Long Range Planning*, 45(5–6), 359–394. <https://doi.org/10.1016/j.lrp.2012.10.001>
- Cepeda-Carrion, G., Nitzl, C., & Roldán, J. L. (2017). Mediation analyses in partial least squares structural equation modeling: Guidelines and empirical examples. En H. Latan & R. Noonan (Eds.), *Partial least squares path modeling: Basic concepts, methodological issues and applications* (pp. 173–195). Springer. https://doi.org/10.1007/978-3-319-64069-3_8
- Chin, W. W. (1998). The partial least squares approach to structural equation modeling. En G. A. Marcoulides (Ed.), *Modern methods for business research* (pp. 295–336). Lawrence Erlbaum. https://www.researchgate.net/publication/311766005_The_Partial_Least_Squares_Approach_to_Structural_Equation_Modeling
- Churchill, G. A. (1979). A paradigm for developing better measures of marketing constructs. *Journal of Marketing Research*, 16(1), 64–73. <https://nguyenthphanhphuong.com/wp-content/uploads/2026/02/Churchill-1979.pdf>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Lawrence Erlbaum Associates. <https://utstat.toronto.edu/~brunner/oldclass/378f16/readings/CohenPower.pdf>

Juan Mejía Tréjo

- Cohen, J. (1992). A power primer. *Psychological Bulletin*, 112(1), 155–159.
<https://www2.psych.ubc.ca/~schaller/528Readings/Cohen1992.pdf>
- Diamantopoulos, A., & Winklhofer, H. M. (2001). Index construction with formative indicators: An alternative to scale development. *Journal of Marketing Research*, 38(2), 269–277
<https://doi.org/10.1509/jmkr.38.2.269.188>
- Dijkstra, T. K., & Henseler, J. (2015). Consistent partial least squares *path* modeling. *MIS Quarterly*, 39(2), 297–316. <https://doi.org/10.25300/MISQ/2015/39.2.02>
- Franke, G., & Sarstedt, M. (2019). Heuristics versus statistics in discriminant validity testing: A comparison of four procedures. *Internet Research*, 29(3), 430–447.
<https://doi.org/10.1108/IntR-12-2017-0515>
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50. <https://doi.org/10.1177/002224378101800104>
- Gefen, D., Straub, D. W., & Boudreau, M.-C. (2000). Structural equation modeling and regression: Guidelines for research practice. *Communications of the Association for Information Systems*, 4, Article 7.
<https://cits.tamui.edu/kock/nedwebarticles/gefenetal2000.pdf>
- Green, S. B. (1991). How many subjects does it take to do a regression analysis? *Multivariate Behavioral Research*, 26(3), 499–510.
<https://parsmodir.com/wp-content/uploads/2013/03/green1991.pdf>
- Gregor, S. (2006). The nature of theory in information systems. *MIS Quarterly*, 30(3), 611–642. <https://doi.org/10.2307/25148742>
- Gudergan, S. P., Ringle, C. M., Wende, S., & Will, A. (2008). Confirmatory tetrad analysis in PLS *path* modeling. *Journal of Business Research*, 61(12), 1238–1249.
<https://doi.org/10.1016/j.jbusres.2008.01.012>
- Guenther, P., Guenther, M., Ringle, C. M., Zaefarian, G., & Cartwright, S. (2023). Improving **PLS-SEM** use for business marketing research. *Industrial Marketing Management*, 111, 127–142. PDF: <https://eprints.whiterose.ac.uk/id/eprint/200118/>
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of **PLS-SEM**. *European Business Review*, 31(1), 2–24.
[https://htw.oil.com.tw/web/cycuim/phd/asset/phd09/Romi_2019_Hairetal_EBR PL S-SEM.pdf](https://htw.oil.com.tw/web/cycuim/phd/asset/phd09/Romi_2019_Hairetal_EBR_PL_S-SEM.pdf)
- Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., & Ray, S. (2021). *Partial least squares structural equation modeling (PLS-SEM) using R*. Springer.
<https://link.springer.com/book/10.1007/978-3-030-80519-7>
- 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 Publications.
https://www.PLS-SEM.net/PLS-SEM-books/a-primer-on-PLS-SEM-3rd-ed/?utm_source=chatgpt.com
- Henseler, J. (2017). Bridging design and behavioral research with variance-based structural equation modeling. *Journal of Advertising*, 46(1), 178–192. PDF: https://ris.utwente.nl/ws/files/6995998/2017_JA_Henseler.pdf
- 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, 115–135.:
<https://run.unl.pt/server/api/core/bitstreams/d124ea66-0ba4-4c13-a658-af460fbce714/content>
- Hu, L. T., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1–55.
<https://doi.org/10.1080/10705519909540118>
- Jarvis, C. B., MacKenzie, S. B., & Podsakoff, P. M. (2003). A critical review of construct indicators and measurement model misspecification in marketing and consumer research. *Journal of Consumer Research*, 30(2), 199–218.
https://statmath.wu.ac.at/courses/r-se-mbr/dl/M1c_Jarvis_Mackenzie_Podsakoff_2003.pdf
- Kock, N. (2015). Common method bias in **PLS-SEM**: A full collinearity assessment approach. *International Journal of e-Collaboration*, 11(4), 1–10.
https://cits.tamtu.edu/kock/pubs/journals/2015JournalJeC_CommMethBias/Kock_2015_IJeC_CommonMethodBiasPLS.pdf
- Kock, N., & Hadaya, P. (2018). Minimum sample size estimation in **PLS-SEM**: The inverse square root and gamma-exponential methods. *Information Systems Journal*, 28(1), 227–261.
https://cits.tamtu.edu/kock/pubs/journals/2018/Kock_Hadaya_2018_ISJ_SampleSizeePLS.pdf
- Kock, N., & Lynn, G. S. (2012). Lateral collinearity and misleading results in variance-based SEM: An illustration and recommendations. *Journal of the Association for Information Systems*, 13(7), 546–580. <https://doi.org/10.17705/1jais.00302>
- Mejía-Trejo, J. (2025). Innovating sustainable artificial intelligence citizenship: A qualitative study of the **CAITIZEN** model using ATLAS.ti. *Scientia et PRAXIS*, 5(10), 126–154. <https://doi.org/10.55965/setp.5.10.a5>
- Mejía-Trejo, J. (2026). *De la idea al modelo estructural: Investigación conceptual en ciencias de la administración*. AMIDI Editorial.
<https://doi.org/10.55965/abib.9789709606133>
- Reinartz, W., Haenlein, M., & Henseler, J. (2009). An empirical comparison of the efficacy of covariance-based and variance-based SEM. *International Journal of Research in Marketing*, 26(4), 332–344. PDF:
https://flora.insead.edu/fichiersti_wp/inseadwp2009/2009-44.pdf
- Ringle, C. M., Sarstedt, M., & Straub, D. W. (2012). Editor's comments: A critical look at the use of **PLS-SEM** in MIS Quarterly. *MIS Quarterly*, 36(1), iii–xiv.
<https://doi.org/10.2307/41410402>
- Sarstedt, M., Hair, J. F., Jr., 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.
<https://doi.org/10.1016/j.ausmj.2019.05.003>
- Schuberth, F., Rademaker, M. E., & Henseler, J. (2020). Estimating and assessing second-order constructs using PLS-PM: The case of composites of composites. *Industrial Management & Data Systems*, 120(12), 2211–2241.
<https://doi.org/10.1108/IMDS-12-2019-0642>

REFERENCIAS

- Shmueli, G. (2010). To explain or to predict? *Statistical Science*, 25(3), 289–310.
<https://doi.org/10.1214/10-STS330> PDF:
<https://www.stat.berkeley.edu/~aldous/157/Papers/shmueli.pdf>
- Shmueli, G., Sarstedt, M., Hair, J. F., Cheah, J.-H., Ting, H., Vaithilingam, S., & Ringle, C. M. (2019). Predictive model assessment in **PLS-SEM**: Guidelines for using PLSpredict. *European Journal of Marketing*, 53(11), 2322–2347.
<https://doi.org/10.1108/EJM-02-2019-0189>
- SmartPLS GmbH. (s. f.-a). *Algorithms and techniques*. SmartPLS.
<https://www.smartpls.com/documentation/algorithms-and-techniques/>
- SmartPLS GmbH. (s. f.-b). **Bootstrapping**. SmartPLS.
<https://www.smartpls.com/documentation/algorithms-and-techniques/Bootstrapping/>
- SmartPLS GmbH. (s. f.-c). **PLS-SEM** algorithm. SmartPLS.
<https://www.smartpls.com/documentation/algorithms-and-techniques/pls/>
- SmartPLS GmbH. (s. f.-d). *SmartPLS documentation: Resources*. SmartPLS.
<https://www.smartpls.com/documentation/>
- SmartPLS GmbH. (s. f.-e). *Your first PLS path model*. SmartPLS.
<https://www.smartpls.com/documentation/tutorials/first-pls-path-model/>
- Williams, L. J., Vandenberg, R. J., & Edwards, J. R. (2009). Structural equation modeling in management research: A guide for improved analysis. *The Academy of Management Annals*, 3(1), 543–604.
<https://doi.org/10.5465/19416520903065683>.

Análisis estructural aplicado. Uso de SmartPLS en la investigación de las ciencias de la administración

© 2026 Academia Mexicana de Investigación y Docencia en Innovación (AMIDI)

Maquetación, diseño y distribución digital:

Academia Mexicana de Investigación y Docencia en Innovación (**AMIDI**).
Responsable del registro DOI, la gestión de metadatos y la publicación en EL
repositorio de AMIDI. Biblioteca perteneciente a Editorial AMIDI

Av. Paseo de los Virreyes 920,
Col. Virreyes Residencial
C.P. 45110
Zapopan, Jalisco, México

eBook hecho y editado en México / eBook made and edited in Mexico
Se terminó de editar en **Julio de 2026**

Juan Mejía Tréjo



Guadalajara, Jal. a 10 de junio de 2026

DICTAMEN DE OBRA ACADÉMICA COMO PRODUCTO DE INVESTIGACIÓN

La **Curaduría Editorial de AMIDI.Biblioteca**, en coordinación con el **Comité Científico**, el **Comité de Redacción y Revisión Inicial** y el proceso de **evaluación académica por pares externos**, documentó el proceso editorial, técnico y académico de la obra:

Análisis estructural aplicado:
Uso de SmartPLS en la investigación de las ciencias de la administración

AUTORES, COORDINADORES O RESPONSABLES ACADÉMICOS

Nombre completo	Institución de adscripción	Rol
Juan Mejía Trejo	Universidad de Guadalajara (CUCEA-UdeG)	Autor

MODALIDAD EDITORIAL Y DE PROCEDENCIA ACADÉMICA DE LA OBRA

X Edición AMIDI	Obra editada, evaluada y publicada bajo responsabilidad editorial de AMIDI Editorial, conforme a sus políticas académicas y editoriales.
☐ Coedición AMIDI / producto de investigación conjunta interinstitucional	Obra derivada de colaboración académica entre autores, grupos de investigación o instituciones participantes, publicada mediante un proceso formal de coedición con AMIDI Editorial.
☐ Obra académica asociada a proceso de coedición AMIDI	Obra académica vinculada a un proceso de coedición, colaboración institucional o articulación editorial con AMIDI Editorial que incorpora a AMIDI.Biblioteca por su pertinencia académica, trazabilidad institucional, valor científico, educativo o formativo, y afinidad con las líneas académicas y editoriales del sello AMIDI. Esta modalidad no modifica la titularidad, el ISBN, la responsabilidad editorial de origen ni los derechos previamente establecidos, salvo que un instrumento formal disponga lo contrario.

INSTANCIAS EDITORIALES, DE EVALUACIÓN ACADÉMICA Y TRAZABILIDAD

- **Curaduría Editorial:** coordinó la revisión editorial inicial; verificó la pertinencia de la obra, la documentación legal, editorial y técnica, la modalidad editorial declarada, y dio seguimiento al dictamen, edición, maquetación, metadatos, preservación digital y publicación.
- **Comité Científico:** asesoró sobre la pertinencia académica general de la obra, revisó de forma preliminar y colegiada la relación entre contenido, enfoque, capítulos y tema central, y sugirió perfiles especializados para la evaluación académica por pares externos.

- **Comité de Redacción y Revisión Inicial:** realizó la revisión preliminar de forma, estructura, coherencia argumentativa, aparato crítico, referencias, cumplimiento de normas editoriales y pertinencia para avanzar a evaluación académica por pares externos.
- **Personas dictaminadoras externas:** evaluaron bajo modalidad doble ciego la calidad, originalidad, pertinencia, consistencia metodológica, claridad argumentativa, actualidad bibliográfica, estructura académica, aportación al campo de conocimiento y posible impacto académico, científico, tecnológico, social, humanístico o aplicado de la obra.

ACCIONES REALIZADAS DURANTE EL PROCESO EDITORIAL, TÉCNICO Y DE VISIBILIDAD ACADÉMICA

1. Verificación de consistencia entre portada, página legal, ficha pública, metadatos, DOI, ISBN, licencia, archivo digital, modalidad editorial declarada y documentación de respaldo.
2. Verificación del cumplimiento de políticas editoriales, éticas, de acceso abierto, licenciamiento, originalidad, antiplagio, revisión académica, preservación digital, interoperabilidad y trazabilidad editorial.
3. Integración del expediente editorial con recepción, revisión inicial, evaluación académica, observaciones, respuesta de autores, autorización editorial y evidencias de trazabilidad.
4. Comunicación de dictámenes y verificación de atención de recomendaciones, ajustes académicos, mejoras editoriales y reducción de similitudes, cuando resultó aplicable.
5. Gestión técnica de metadatos, identificadores persistentes, ficha pública, preservación digital, acceso abierto e interoperabilidad para facilitar la recuperación, cosecha, consulta y trazabilidad académica de la obra.
6. Gestión de visibilidad, distribución y seguimiento académico en plataformas, catálogos, repositorios, buscadores, sistemas bibliográficos, cosechadores, entornos biométricos y canales de difusión académica pertinentes, sin garantía de inclusión, permanencia, indexación, citación o impacto en sistemas externos.
7. Conservación de evidencias del proceso editorial, académico, técnico, documental y de visibilidad como parte de la trazabilidad editorial de la obra, así como de su publicación continua.

ESTRUCTURA DE LA OBRA EVALUADA

La obra evaluada se organiza académicamente de la siguiente manera:

Capítulo	Título
Introducción	
Capítulo 1	Bases del análisis estructural aplicado



Capítulo 2	Bases operativas del modelado PLS-SEM en SmartPLS4
Capítulo 3	Construcción del modelo en SmartPLS4
Capítulo 4	PLS-SEM Algorithm y colinealidad
Capítulo 5	Evaluación del modelo de medición reflectivo
Capítulo 6	Ajuste global del modelo saturado
Capítulo 7	Bootstrapping e inferencia del modelo
Capítulo 8	Evaluación del modelo estructural
Capítulo 9	De la explicación estructural a la predicción del modelo
Capítulo 10	Integración del modelo PLS-SEM
Referencias	

Para efectos del presente dictamen, la estructura se consigna hasta el nivel de capítulo y subtema principal. El índice completo, con apartados y subapartados específicos, forma parte del expediente editorial y de la obra publicada.

CLASIFICACIÓN ACADÉMICA DE LA OBRA COMO PRODUCTO ACADÉMICO DE INVESTIGACIÓN REGISTRADO COMO LIBRO ELECTRÓNICO

Con base en el proceso editorial, técnico y académico documentado por **AMIDI Editorial / AMIDI.Biblioteca**, la obra se clasifica como **producto académico de investigación registrado como libro electrónico**, conforme a su fundamentación, metodología, estructura, evaluación académica, contribución al campo de conocimiento y trazabilidad editorial.

Categoría aplicable	Criterio breve de clasificación
☐ Investigación teórica-conceptual	Desarrolla, analiza, sistematiza o propone conceptos, categorías, modelos, enfoques o relaciones explicativas.
☐ Investigación documental-integrativa	Se sustenta en revisión, análisis, comparación, síntesis o integración de literatura académica y documentos especializados.
☐ Investigación empírica	Se apoya en datos primarios o secundarios mediante métodos cuantitativos, cualitativos o mixtos.
X Investigación aplicada	Traduce conocimiento académico en modelos, métodos, diagnósticos, instrumentos, procesos, soluciones o propuestas para contextos específicos.
☐ Investigación aplicada con orientación formativa	Obra con finalidad didáctica, profesionalizante o de apoyo a procesos de enseñanza-aprendizaje, sustentada en revisión especializada, sistematización conceptual, diseño metodológico, casos, ejercicios, instrumentos, rutas de aprendizaje o propuestas aplicadas académicamente



	fundamentadas.
☐ Investigación metodológica o instrumental	Desarrolla, explica, valida o sistematiza métodos, técnicas, escalas, instrumentos, procedimientos o rutas analíticas.
☐ Innovación, desarrollo tecnológico o transferencia	Presenta modelos, procesos, herramientas, metodologías, aplicaciones o propuestas vinculadas con innovación o solución de problemas.

DETERMINACIÓN ACADÉMICA Y EDITORIAL

Una vez emitidas las observaciones y recomendaciones de mejora por parte de los evaluadores, y una vez atendidas por los autores, se determina que la obra reúne los elementos necesarios para ser reconocida como **libro académico de investigación registrado en modalidad electrónica**, al presentar fundamentación especializada, metodología explícita, coherencia entre título, objeto de estudio, estructura capitular, desarrollo temático, conclusiones, referencias y contribución al campo de conocimiento. La orientación aplicada y formativa de algunos capítulos no reduce el carácter investigativo de la obra; por el contrario, expresa la transferencia y aplicación del conocimiento académico en procesos de investigación cuantitativa, modelación estructural, análisis metodológico y desarrollo de capacidades de investigación.

De esta forma, el resultado del dictamen de aceptación de la obra fue:

FAVORABLE PARA SU PUBLICACIÓN, PRESERVACIÓN, DIFUSIÓN E INCORPORACIÓN COMO LIBRO ACADÉMICO DE INVESTIGACIÓN REGISTRADO EN MODALIDAD ELECTRÓNICA

Sirva el presente dictamen como constancia para los fines académicos, editoriales, institucionales, curriculares, de evaluación y trazabilidad documental que a los interesados convengan.

Atentamente



Dr. Carlos Omar Aguilar Navarro
Curador AMIDI Editorial